

Variational Bayes Updates for Normal Linear Regression

Santiago Montoya-Blandón

ECON 5120: Bayesian Data Analysis
Adam Smith Business School, University of Glasgow

Setup

Recall that the Bayesian linear regression framework requires a likelihood and a set of priors. We have an outcome $\mathbf{y} = (y_1, \dots, y_n)$ that we aim to predict using a set of covariates or features $\mathbf{X}$ arranged in an $n \times p$ matrix such that

$$\mathbf{X} = \begin{bmatrix} x_{1,1} & \cdots & x_{1,p} \\ \vdots & \ddots & \vdots \\ x_{n,1} & \cdots & x_{n,p} \end{bmatrix}.$$

We do not make restrictions on the size of p (the amount of features). Indeed, variational techniques were introduced into learning procedures as a way to obtain reliable and scalable inference for models with large p , potentially with $p > n$. As the likelihood for this model we will use a classical normal linear regression model:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u}, \quad \mathbf{u} \sim \mathcal{N}_n(\mathbf{0}_n, \sigma^2 \mathbf{I}_n) \implies \mathbf{y} \mid \boldsymbol{\beta}, \sigma^2; \mathbf{X} \sim \mathcal{N}_n(\mathbf{X}\boldsymbol{\beta}, \sigma^2 \mathbf{I}_n), \quad (1)$$

where $\mathbf{0}_n$ is a n -dimensional vector of zeros and $\mathbf{I}_n$ is the $n \times n$ identity matrix. We will use standard independent priors for $\boldsymbol{\beta}$ and σ^2 , such that the joint prior is $p(\boldsymbol{\beta}, \sigma^2) = p(\boldsymbol{\beta})p(\sigma^2)$ with

$$\boldsymbol{\beta} \sim \mathcal{N}_p(\underline{\boldsymbol{\beta}}, \underline{\mathbf{B}}) \quad \text{and} \quad \sigma^2 \sim \text{Inv-Gamma}(\underline{\alpha}/2, \underline{\delta}/2) \quad (2)$$

Using Bayes' rule (posterior proportional to likelihood times prior), we can write the joint posterior $p(\boldsymbol{\beta}, \sigma^2 \mid \mathbf{y}, \mathbf{X}) \propto p(\mathbf{y} \mid \boldsymbol{\beta}, \sigma^2; \mathbf{X})p(\boldsymbol{\beta}, \sigma^2)$ without proportionality constants as

$$\begin{aligned} p(\boldsymbol{\beta}, \sigma^2 \mid \mathbf{y}, \mathbf{X}) &\propto (\sigma^2)^{-n/2} \exp \left\{ -\frac{1}{2\sigma^2} (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^\top (\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) \right\} \\ &\times \exp \left\{ -\frac{1}{2} (\boldsymbol{\beta} - \underline{\boldsymbol{\beta}})^\top \underline{\mathbf{B}}^{-1} (\boldsymbol{\beta} - \underline{\boldsymbol{\beta}}) \right\} \\ &\times (\sigma^2)^{-\underline{\alpha}/2-1} \exp \left\{ -\frac{\underline{\delta}}{2\sigma^2} \right\} \end{aligned} \quad (3)$$

We will leave this expression as is, and use the terms we need to build the respective conditionals that we require for the optimal variational approximations. Before proceeding,

I want to remind you of a very useful trick known as “completing the square”. For any p -dimensional vectors $\mathbf{x}$, $\mathbf{b}$, and $p \times p$ matrix $\mathbf{B}$, we have

$$\begin{aligned} (\mathbf{x} - \mathbf{b})^\top \mathbf{B}^{-1} (\mathbf{x} - \mathbf{b}) &= \mathbf{x}^\top \mathbf{B}^{-1} \mathbf{x} - \mathbf{b}^\top \mathbf{B}^{-1} \mathbf{x} - \mathbf{x}^\top \mathbf{B}^{-1} \mathbf{b} + \mathbf{b}^\top \mathbf{B}^{-1} \mathbf{b} \\ &= \mathbf{x}^\top \mathbf{B}^{-1} \mathbf{x} - 2\mathbf{b}^\top \mathbf{B}^{-1} \mathbf{x} + \mathbf{b}^\top \mathbf{B}^{-1} \mathbf{b} \end{aligned}$$

This means we can complete a quadratic form as follows:

$$\begin{aligned} Q &:= \mathbf{x}^\top \mathbf{B}^{-1} \mathbf{x} - 2\mathbf{b}^\top \mathbf{B}^{-1} \mathbf{x} \\ \implies Q &= (\mathbf{x} - \mathbf{b})^\top \mathbf{B}^{-1} (\mathbf{x} - \mathbf{b}) - \mathbf{b}^\top \mathbf{B}^{-1} \mathbf{b} \end{aligned} \tag{4}$$

Variational Approximations

General case with mean-field approximations

Recall that to implement variational Bayes we require a class of approximations $\mathcal{Q}$. We usually use *mean-field approximations*, which for this problem means we can define

$$\mathcal{Q} := \{q(\boldsymbol{\beta}, \sigma^2) : q(\boldsymbol{\beta}, \sigma^2) = q(\boldsymbol{\beta}) \cdot q(\sigma^2)\}$$

Note that we are not constraining the forms of $q(\boldsymbol{\beta})$ or $q(\sigma^2)$ at this stage, we are only proposing that they will be independent. This means that rather than assuming particular functional forms for these distributions, the distributional forms for $q(\boldsymbol{\beta})$ and $q(\sigma^2)$ will be obtained as a consequence of the optimal approximation. Recall the optimal approximation in the general case is given as

$$\begin{aligned} q_j^*(\theta_j) &\propto \exp \{E_{-j}[\log p(\theta_j \mid \boldsymbol{\theta}_{-j}; \mathbf{y})]\} \\ &\exp \left\{ \int \log p(\theta_j \mid \boldsymbol{\theta}_{-j}; \mathbf{y}) \prod_{\ell \neq j} q_\ell^*(\theta_\ell) d\boldsymbol{\theta}_{-j} \right\} \end{aligned} \tag{5}$$

where $p(\theta_j \mid \boldsymbol{\theta}_{-j}; \mathbf{y})$ is the conditional posterior of θ_j conditional on all other parameters excluding θ_j , denoted as $\boldsymbol{\theta}_{-j}$, across $j = 1, \dots, p$. Note that the expectation in (5) is obtained with respect to the optimal approximations of the other parameters $q_1^*(\theta_1), \dots, q_{j-1}^*(\theta_{j-1}), q_{j+1}^*(\theta_{j+1}), \dots, q_p^*(\theta_p)$. In addition, notice that if we can write $E_{-j}[\log p(\theta_j \mid \boldsymbol{\theta}_{-j}; \mathbf{y})]$ as a function of θ_j , say $f(\theta_j)$, and a constant with respect to θ_j , then (5) implies we would have

$$q_j^*(\theta_j) \propto \exp \{f(\theta_j) + \text{const.}\} \propto \exp \{f(\theta_j)\}$$

This means we can ignore all constant terms that appear in the $E_{-j}[\log p(\theta_j \mid \boldsymbol{\theta}_{-j}; \mathbf{y})]$ terms.

Variational approximation for Bayesian linear regression with independent priors

In our normal linear regression case, the optimal approximations take the shape

$$\begin{aligned} q^*(\boldsymbol{\beta}) &\propto \exp \{E_{q^*(\sigma^2)}[\log p(\boldsymbol{\beta} \mid \sigma^2; \mathbf{y}, \mathbf{X})]\} \\ q^*(\sigma^2) &\propto \exp \{E_{q^*(\boldsymbol{\beta})}[\log p(\sigma^2 \mid \boldsymbol{\beta}; \mathbf{y}, \mathbf{X})]\} \end{aligned} \tag{6}$$

We will show the independence assumption of the mean-field approximation is all we need to obtain the specific distributional forms that solve (6). Start with $q^*(\boldsymbol{\beta})$, which requires the conditional posterior $p(\boldsymbol{\beta} \mid \sigma^2; \mathbf{y}, \mathbf{X})$. To obtain this posterior, we can simply keep all terms from (3) that contain $\boldsymbol{\beta}$:

$$p(\boldsymbol{\beta}, \sigma^2 \mid \mathbf{y}, \mathbf{X}) \propto \exp \left\{ -\frac{1}{2} \left[\frac{1}{\sigma^2} (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^\top (\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) + (\boldsymbol{\beta} - \underline{\boldsymbol{\beta}})^\top \underline{\mathbf{B}}^{-1} (\boldsymbol{\beta} - \underline{\boldsymbol{\beta}}) \right] \right\} \quad (7)$$

Recall the notation $f(x) \propto g(x)$ means $f(x) = c \cdot g(x)$ for any functions f and g and some scalar proportionality constant c that is independent of x . This means $\log f(x) = \log(c) + \log g(x) = \log g(x) + \text{const.}$ follows as a consequence. Taking logarithms and expectations with respect to $q^*(\sigma^2)$ yields

$$\begin{aligned} \mathbb{E}_{q^*(\sigma^2)}[\log p(\boldsymbol{\beta} \mid \sigma^2; \mathbf{y}, \mathbf{X})] &= -\frac{1}{2} \mathbb{E}_{q^*(\sigma^2)} \left[\frac{1}{\sigma^2} (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^\top (\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) \right. \\ &\quad \left. + (\boldsymbol{\beta} - \underline{\boldsymbol{\beta}})^\top \underline{\mathbf{B}}^{-1} (\boldsymbol{\beta} - \underline{\boldsymbol{\beta}}) \right] + \text{const.} \\ &= -\frac{1}{2} \mathbb{E}_{q^*(\sigma^2)} \left[\frac{1}{\sigma^2} \right] (\mathbf{y}^\top \mathbf{y} - 2\boldsymbol{\beta}^\top \mathbf{X}^\top \mathbf{y} + \boldsymbol{\beta}^\top \mathbf{X}^\top \mathbf{X} \boldsymbol{\beta}) \\ &\quad - \frac{1}{2} (\boldsymbol{\beta}^\top \underline{\mathbf{B}}^{-1} \boldsymbol{\beta} - 2\boldsymbol{\beta}^\top \underline{\mathbf{B}}^{-1} \underline{\boldsymbol{\beta}} + \underline{\boldsymbol{\beta}}^\top \underline{\mathbf{B}}^{-1} \underline{\boldsymbol{\beta}}) + \text{const.} \\ &= -\frac{1}{2} \left[\boldsymbol{\beta}^\top \left(\mathbb{E}_{q^*(\sigma^2)} \left[\frac{1}{\sigma^2} \right] \mathbf{X}^\top \mathbf{X} + \underline{\mathbf{B}}^{-1} \right) \boldsymbol{\beta} \right. \\ &\quad \left. - 2\boldsymbol{\beta}^\top \left(\mathbb{E}_{q^*(\sigma^2)} \left[\frac{1}{\sigma^2} \right] \mathbf{X}^\top \mathbf{y} + \underline{\mathbf{B}}^{-1} \underline{\boldsymbol{\beta}} \right) \right] + \text{const.} \\ &= -\frac{1}{2} (\boldsymbol{\beta}^\top \bar{\mathbf{B}}^{-1} \boldsymbol{\beta} - 2\boldsymbol{\beta}^\top \bar{\mathbf{B}}^{-1} \bar{\boldsymbol{\beta}}) + \text{const.} \end{aligned} \quad (8)$$

where we defined

$$\bar{\mathbf{B}} := \left(\mathbb{E}_{q^*(\sigma^2)} \left[\frac{1}{\sigma^2} \right] \mathbf{X}^\top \mathbf{X} + \underline{\mathbf{B}}^{-1} \right)^{-1} \quad \text{and} \quad \bar{\boldsymbol{\beta}} := \bar{\mathbf{B}} \left(\mathbb{E}_{q^*(\sigma^2)} \left[\frac{1}{\sigma^2} \right] \mathbf{X}^\top \mathbf{y} + \underline{\mathbf{B}}^{-1} \underline{\boldsymbol{\beta}} \right) \quad (9)$$

Note that while we do not know how to solve for the term $\mathbb{E}_{q^*(\sigma^2)}[1/\sigma^2]$ yet, we do know it will not depend on $\boldsymbol{\beta}$. This means that both $\bar{\mathbf{B}}$ and $\bar{\boldsymbol{\beta}}$ are constants with respect to $\boldsymbol{\beta}$. We can then complete the square on $\boldsymbol{\beta}$ in equation (8) using the result in (4) to obtain¹

$$\mathbb{E}_{q^*(\sigma^2)}[\log p(\boldsymbol{\beta} \mid \sigma^2; \mathbf{y}, \mathbf{X})] = -\frac{1}{2} (\boldsymbol{\beta} - \bar{\boldsymbol{\beta}})^\top \bar{\mathbf{B}}^{-1} (\boldsymbol{\beta} - \bar{\boldsymbol{\beta}}) + \text{const.}$$

As we can ignore all constant factors in these expectations, we can then write

$$q^*(\boldsymbol{\beta}) \propto \exp \left\{ -\frac{1}{2} (\boldsymbol{\beta} - \bar{\boldsymbol{\beta}})^\top \bar{\mathbf{B}}^{-1} (\boldsymbol{\beta} - \bar{\boldsymbol{\beta}}) \right\}$$

¹We set $\mathbf{x} = \boldsymbol{\beta}$, $\mathbf{b} = \bar{\boldsymbol{\beta}}$, and $\mathbf{B} = \bar{\mathbf{B}}$ when completing the square, using the fact that $\bar{\boldsymbol{\beta}} = \bar{\mathbf{B}}^{-1} \bar{\mathbf{B}} \bar{\boldsymbol{\beta}}$.

Even though we do not yet have a closed-form expression for $E_{q^*(\sigma^2)}[1/\sigma^2]$, we can recognize the previous equation as the kernel (or quasi-density) of a multivariate normal distribution on β with mean vector $\bar{\beta}$ and variance-covariance matrix $\bar{B}$. That is, we have found

$$q^*(\beta) = \mathcal{N}_p(\beta \mid \bar{\beta}, \bar{B}) \implies \beta \mid y, X \sim_{q^*} \mathcal{N}_p(\bar{\beta}, \bar{B}) \quad (10)$$

where the notation “ $\sim_{q^*}$ ” emphasizes that this is not the true exact posterior, but rather the posterior under the optimal variational approximation. Let us proceed with finding $q^*(\sigma^2)$ and we will go back and calculate the elements in (9) once this is done. As before, we start from the conditional posterior $p(\sigma^2 \mid \beta; \mathbf{y}, \mathbf{X})$, found by keeping all terms from (3) that contain σ^2 :

$$p(\sigma^2 \mid \beta; \mathbf{y}, \mathbf{X}) \propto (\sigma^2)^{-(n+\underline{\alpha})/2-1} \times \exp \left\{ -\frac{1}{2\sigma^2} [(\mathbf{y} - \mathbf{X}\beta)^\top (\mathbf{y} - \mathbf{X}\beta) + \underline{\delta}] \right\} \quad (11)$$

Taking logarithms and expectations with respect to $q^*(\beta)$ we obtain

$$\begin{aligned} E_{q^*(\beta)}[\log p(\sigma^2 \mid \beta; \mathbf{y}, \mathbf{X})] &= -\frac{1}{2\sigma^2} \{ E_{q^*(\beta)} [(\mathbf{y} - \mathbf{X}\beta)^\top (\mathbf{y} - \mathbf{X}\beta)] + \underline{\delta} \} \\ &\quad + \left(-\frac{n + \underline{\alpha}}{2} - 1 \right) \log(\sigma^2) + \text{const.} \\ &= -\frac{\bar{\delta}}{2\sigma^2} + \left(-\frac{\bar{\alpha}}{2} - 1 \right) \log(\sigma^2) + \text{const.} \end{aligned} \quad (12)$$

where we have defined

$$\bar{\alpha} := \underline{\alpha} + n \quad \text{and} \quad \bar{\delta} := \underline{\delta} + E_{q^*(\beta)} [(\mathbf{y} - \mathbf{X}\beta)^\top (\mathbf{y} - \mathbf{X}\beta)] \quad (13)$$

As before, we do not have closed-form expressions for these posterior hyperparameters yet, but we do know they are independent of σ^2 . Ignoring constant terms, we now obtain

$$q^*(\sigma^2) \propto (\sigma^2)^{-\bar{\alpha}/2-1} \exp \left\{ -\frac{\bar{\delta}}{2\sigma^2} \right\}$$

We recognize this as the kernel of another inverse Gamma distribution, meaning we obtained

$$q^*(\sigma^2) = \text{Inv-Gamma} \left(\sigma^2 \mid \frac{\bar{\alpha}}{2}, \frac{\bar{\delta}}{2} \right) \implies \sigma^2 \mid y, X \sim_{q^*} \text{Inv-Gamma} \left(\frac{\bar{\alpha}}{2}, \frac{\bar{\delta}}{2} \right) \quad (14)$$

Combining (10) and (14), we have the full variational approximation

$$q^*(\beta, \sigma^2) = q^*(\beta) \cdot q^*(\sigma^2) \quad (15)$$

In order to make this solution operational, we need to solve for the expectations $E_{q^*(\sigma^2)}[1/\sigma^2]$ and $E_{q^*(\beta)} [(\mathbf{y} - \mathbf{X}\beta)^\top (\mathbf{y} - \mathbf{X}\beta)]$. We now know the form of the individual optimal approximations, so this can be accomplished analytically. For $E_{q^*(\sigma^2)}[1/\sigma^2]$, note that as σ^2 is distributed as inverse Gamma under the variational approximation, $1/\sigma^2$ is distributed as

Gamma with shape $\bar{\alpha}/2$ and *rate* $\bar{\delta}/2$. Using known properties of the Gamma distribution, this means

$$\mathbb{E}_{q^*(\sigma^2)} \left[\frac{1}{\sigma^2} \right] = \frac{\bar{\alpha}}{\bar{\delta}} \quad (16)$$

For $\mathbb{E}_{q^*(\beta)} [(\mathbf{y} - \mathbf{X}\beta)^\top (\mathbf{y} - \mathbf{X}\beta)]$, note that $\beta \mid \mathbf{y}, \mathbf{X} \sim_{q^*} \mathcal{N}_p(\bar{\beta}, \bar{\mathbf{B}})$ implies $\varepsilon \mid \mathbf{y}, \mathbf{X} \sim_{q^*} \mathcal{N}_n(\mathbf{y} - \mathbf{X}\bar{\beta}, \mathbf{X}\bar{\mathbf{B}}\mathbf{X}^\top)$, where $\varepsilon := \mathbf{y} - \mathbf{X}\beta$. We again invoke known results to find²

$$\mathbb{E}_{q^*(\beta)} [\varepsilon^\top \varepsilon] = \text{tr}(\bar{\mathbf{B}}\mathbf{X}^\top \mathbf{X}) + (\mathbf{y} - \mathbf{X}\bar{\beta})^\top (\mathbf{y} - \mathbf{X}\bar{\beta}) \quad (17)$$

where $\text{tr}(\mathbf{A})$ is the trace operator that equals the sum of the elements of the diagonal for any matrix $\mathbf{A}$.

Evidence Lower Bound

Finally, we note how to evaluate the quality of the variational approximation. In the general case, recall we are using the *evidence lower bound* (ELBO) as our objective function to evaluate any approximation $q \in \mathcal{Q}$:

$$\text{elbo}(q) = \mathbb{E}_q \left[\log \left(\frac{p(\mathbf{y} \mid \boldsymbol{\theta}) p(\boldsymbol{\theta})}{q(\boldsymbol{\theta})} \right) \right] \quad (18)$$

For general approximations, if we have a way to sample from the approximated posterior $q(\boldsymbol{\theta})$, we can estimate (18) using these posterior draws. Under the independence of the mean-field approximation, we can further decompose our objective as

$$\text{elbo}(q) = \mathbb{E}_q [\log p(\mathbf{y} \mid \boldsymbol{\theta})] + \mathbb{E}_q [\log p(\boldsymbol{\theta})] - \sum_{j=1}^p \mathbb{E}_{q_j} [\log q_j(\theta_j)] \quad (19)$$

Replacing the likelihood (1) and prior structure (2) into (19) provides the ELBO formula for Bayesian linear regression. Rather than evaluating all approximations $q \in \mathcal{Q}$, we can focus on computing the ELBO for the *optimal variational approximation* $q^*(\boldsymbol{\beta}, \sigma^2)$:

$$\begin{aligned} \text{elbo}(q^*) &= \mathbb{E}_{q^*} [\log p(\mathbf{y} \mid \boldsymbol{\beta}, \sigma^2; \mathbf{X})] + \mathbb{E}_{q^*(\beta)} [\log p(\boldsymbol{\beta})] + \mathbb{E}_{q^*(\sigma^2)} [\log p(\sigma^2)] \\ &\quad - \mathbb{E}_{q^*(\beta)} [\log q^*(\boldsymbol{\beta})] - \mathbb{E}_{q^*(\sigma^2)} [\log q^*(\sigma^2)] \end{aligned} \quad (20)$$

Now that we know the optimal approximations and that they come from simple families, we can analytically compute every term in (20). These are given in the following expressions,

²See [https://en.wikipedia.org/wiki/Quadratic_form_\(statistics\)#Expectation](https://en.wikipedia.org/wiki/Quadratic_form_(statistics)#Expectation) for a full derivation.

obtained without proof as they use more known results.³

$$\begin{aligned}
E_{q^*} [\log p(\mathbf{y} \mid \boldsymbol{\beta}, \sigma^2; \mathbf{X})] &= \frac{n}{2} \left[\psi \left(\frac{\bar{\alpha}}{2} \right) - \log \left(\frac{\bar{\delta}}{2} \right) \right] - \frac{n}{2} \log(2\pi) \\
&\quad - \frac{\bar{\alpha}}{2\bar{\delta}} [\text{tr}(\bar{\mathbf{B}} \mathbf{X}^\top \mathbf{X}) + (\mathbf{y} - \mathbf{X} \bar{\boldsymbol{\beta}})^\top (\mathbf{y} - \mathbf{X} \bar{\boldsymbol{\beta}})] \\
E_{q^*(\boldsymbol{\beta})} [\log p(\boldsymbol{\beta})] &= -\frac{p}{2} \log(2\pi) - \frac{1}{2} \log(|\underline{\mathbf{B}}|) \\
&\quad - \frac{1}{2} [\text{tr}(\bar{\mathbf{B}} \underline{\mathbf{B}}^{-1}) + (\bar{\boldsymbol{\beta}} - \underline{\boldsymbol{\beta}})^\top \underline{\mathbf{B}}^{-1} (\bar{\boldsymbol{\beta}} - \underline{\boldsymbol{\beta}})] \\
E_{q^*(\sigma^2)} [\log p(\sigma^2)] &= \frac{\alpha}{2} \log \left(\frac{\delta}{2} \right) - \log \left[\Gamma \left(\frac{\alpha}{2} \right) \right] \\
&\quad + \left(\frac{\alpha}{2} + 1 \right) \left[\psi \left(\frac{\bar{\alpha}}{2} \right) - \log \left(\frac{\bar{\delta}}{2} \right) \right] - \frac{\delta \cdot \bar{\alpha}}{2\bar{\delta}} \\
E_{q^*(\boldsymbol{\beta})} [\log q^*(\boldsymbol{\beta})] &= -\frac{p}{2} [1 + \log(2\pi)] - \frac{1}{2} \log(|\bar{\mathbf{B}}|) \\
E_{q^*(\sigma^2)} [\log q^*(\sigma^2)] &= \left(\frac{\bar{\alpha}}{2} + 1 \right) \psi \left(\frac{\bar{\alpha}}{2} \right) - \frac{\bar{\alpha}}{2} - \log \left[\Gamma \left(\frac{\bar{\alpha}}{2} \right) \right] - \log \left(\frac{\bar{\delta}}{2} \right)
\end{aligned} \tag{21}$$

Note that for computing these expectations we do require the proportionality constants, as they ensure the likelihood, prior and variational distributions are all in the same scale. We need this property as we hope to *compare across models* with different posteriors; i.e., different posterior hyperparameters.

While these equations might look complicated, they are all easy to calculate using a computer. This means that once we have obtained a set of values $(\bar{\boldsymbol{\beta}}, \bar{\mathbf{B}}, \bar{\alpha}, \bar{\delta})$, we can easily compute $\text{elbo}(q^*)$ associated to the current optimal variational approximation given by those parameter values. We then select the best fitting model (hyperparameter values) as the one that maximizes the ELBO.

Implementation and Comparison to Gibbs Sampling

We now explore how to implement these equations in practice. Note the interdependent nature of (16) and (17), as computing one requires knowledge of the other. Additionally, note that while we started with completely general approximations, the families of the optimal variational approximations $q^*(\boldsymbol{\beta})$ (normal) and $q^*(\sigma^2)$ (inverse Gamma) were found analytically. All that is left to determine are their posterior hyperparameter values.

In practice, this means we choose initial values for either $\bar{\alpha}$ and $\bar{\delta}$ or $\bar{\boldsymbol{\beta}}$ and $\bar{\mathbf{B}}$, and iteratively update each set of parameters values, alternating between using (9) and (13). This allows us to first compute say $\bar{\boldsymbol{\beta}}$ and $\bar{\mathbf{B}}$, then compute $\bar{\alpha}$ and $\bar{\delta}$, and repeat until convergence to obtain $(\boldsymbol{\beta}^*, \mathbf{B}^*, \alpha^*, \delta^*)$. Finally, we use these pre-optimized values for the

³See https://en.wikipedia.org/wiki/Multivariate_normal_distribution and https://en.wikipedia.org/wiki/Inverse-gamma_distribution for sources and discussions of these properties. Here, $\psi(\cdot)$ represents the digamma function.

posterior hyperparameters to quickly sample many draws from the full variational posterior approximation:

$$q^*(\boldsymbol{\beta}, \sigma^2) = \mathcal{N}_p(\boldsymbol{\beta} \mid \boldsymbol{\beta}^*, \mathbf{B}^*) \cdot \text{Inv-Gamma} \left(\sigma^2 \mid \frac{\alpha^*}{2}, \frac{\delta^*}{2} \right) \quad (22)$$

A procedure that implements the coordinate ascent mean-field variational Bayes is given in Algorithm 1. It is instructive to compare this algorithm to Algorithm 2. This procedure implements a Gibbs sampler that can sample from the exact joint posterior (3) using MCMC methods. Indeed, using the same steps as those used for deriving (8) and (12), we could directly obtain the necessary conditional posteriors (**you should verify this!**):

$$\boldsymbol{\beta} \mid \sigma^2; \mathbf{y}, \mathbf{X} \sim \mathcal{N}_p(\boldsymbol{\beta} \mid \tilde{\boldsymbol{\beta}}, \tilde{\mathbf{B}}) \quad \text{with} \quad (23)$$

$$\tilde{\mathbf{B}} := \left(\frac{\mathbf{X}^\top \mathbf{X}}{\sigma^2} + \underline{\mathbf{B}}^{-1} \right)^{-1} \quad \text{and} \quad \tilde{\boldsymbol{\beta}} := \tilde{\mathbf{B}} \left(\frac{\mathbf{X}^\top \mathbf{y}}{\sigma^2} + \underline{\mathbf{B}}^{-1} \underline{\boldsymbol{\beta}} \right)$$

$$\sigma^2 \mid \boldsymbol{\beta}; \mathbf{y}, \mathbf{X} \sim \text{Inv-Gamma} \left(\sigma^2 \mid \frac{\tilde{\alpha}}{2}, \frac{\tilde{\delta}}{2} \right) \quad \text{with} \quad (24)$$

$$\tilde{\alpha} = \bar{\alpha} \quad \text{and} \quad \tilde{\delta} := \underline{\delta} + (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^\top (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})$$

Notice that Gibbs uses the realizations of $\boldsymbol{\beta}$ and σ^2 to update the posterior hyperparameters within each simulation loop. In contrast, variational Bayes replaces the realized value with the posterior mean of $\boldsymbol{\beta}$ and σ^2 (under the optimal variational approximation) when computing the posterior hyperparameters through (9) and (13). By moving the optimization of the posterior hyperparameters from the inner loop to the outer loop of the code, we achieve dramatic speed gains in fitting and sampling posteriors from Bayesian approaches.

Algorithm 1 Coordinate Ascent Mean-Field Variational Bayes for Linear Regression with Independent Priors

Require: Outcome $\mathbf{y}$, covariates or features $\mathbf{X}$, prior hyperparameters $(\underline{\beta}, \underline{\mathbf{B}}, \underline{\alpha}, \underline{\delta})$, initial values α_0 and δ_0 , burn-in iterations S_0 , keep iterations S_1 , and tolerance ϵ

- 1: $S \leftarrow S_0 + S_1$
- 2: Initialize $\bar{\alpha} \leftarrow \alpha_0$ and $\bar{\delta} \leftarrow \delta_0$ ▷ Initial value of parameters
- 3: Initialize $\text{elbo}_0 \leftarrow -\infty$ and $\text{elbo}_1 \leftarrow 0$
- 4: $\tau \leftarrow |\text{elbo}_1 - \text{elbo}_0|$ ▷ Convergence test quantity
- 5: **while** $\tau > \epsilon$ **do** ▷ Fitting loop goes first
- 6: Update $\tilde{\beta}$ and $\tilde{\mathbf{B}}$ according to (9)
- 7: Update $\bar{\alpha}$ and $\bar{\delta}$ according to (13)
- 8: $\text{elbo}_1 \leftarrow \text{elbo}(q^*)$ using (20) and (21)
- 9: $\tau \leftarrow |\text{elbo}_1 - \text{elbo}_0|$ ▷ Update test quantity
- 10: $\text{elbo}_0 \leftarrow \text{elbo}_1$
- 11: **end while**
- 12: $\beta^* \leftarrow \tilde{\beta}$, $\mathbf{B}^* \leftarrow \tilde{\mathbf{B}}$, $\alpha^* \leftarrow \bar{\alpha}$, and $\delta^* \leftarrow \bar{\delta}$ ▷ Save pre-optimized values
- 13: **for** $s = 1, \dots, S$ **do** ▷ Main simulation loop
- 14: Draw $\beta^{(s)}$ from $\mathcal{N}_p(\beta^*, \mathbf{B}^*)$
- 15: Draw $\sigma^{2(s)}$ from $\text{Inv-Gamma}(\alpha^*/2, \delta^*/2)$
- 16: **end for**
- 17: **return** $(\beta^{(S_0+1)}, \dots, \beta^{(S)})$ and $(\sigma^{2(S_0+1)}, \dots, \sigma^{2(S)})$

Algorithm 2 Gibbs Sampler for Bayesian Linear Regression with Independent Priors

Require: Outcome $\mathbf{y}$, covariates or features $\mathbf{X}$, prior hyperparameters $(\underline{\beta}, \underline{\mathbf{B}}, \underline{\alpha}, \underline{\delta})$, initial value σ_0^2 , burn-in iterations S_0 , and keep iterations S_1

- 1: $S \leftarrow S_0 + S_1$
- 2: Initialize $\sigma^{2(0)} \leftarrow \sigma_0^2$ ▷ Initial draw
- 3: **for** $s = 1, \dots, S$ **do** ▷ Main simulation loop
- 4: Calculate $\tilde{\mathbf{B}}$ and $\tilde{\beta}$ using (23) with $\sigma^2 = \sigma^{2(s-1)}$
- 5: Draw $\beta^{(s)}$ from $\mathcal{N}_p(\tilde{\beta}, \tilde{\mathbf{B}})$ ▷ Update β
- 6: Calculate $\tilde{\alpha}$ and $\tilde{\delta}$ using (24) with $\beta = \beta^{(s)}$
- 7: Draw $\sigma^{2(s)}$ from $\text{Inv-Gamma}(\tilde{\alpha}/2, \tilde{\delta}/2)$ ▷ Update σ^2
- 8: **end for**
- 9: **return** $(\beta^{(S_0+1)}, \dots, \beta^{(S)})$ and $(\sigma^{2(S_0+1)}, \dots, \sigma^{2(S)})$
