

Bayesian State Space Model

Abstract

This methodological framework examines the Bayesian State Space Model (BSSM), which is particularly suitable for estimating model parameters with a high degree of confidence when working with short time series. The BSSM is implemented through several stages, including the Variational Bayesian M-step (maximisation), Variational Bayesian E-step (expectation), calculation of the lower bound of the logarithmic marginal likelihood, updating of hyperparameters, and comparison of the lower bound of the logarithmic marginal likelihood. The objective of this article is to expand the range of modelling approaches used to improve the accuracy of estimates of the fundamental value of housing prices in Uzbekistan's real estate market.

Keywords: covariance matrix, marginal probability, error, prior distribution, posterior distribution, time series.

Bayesian State Space Model

Bayesian state space model (BSSM) is considered more efficient than the state space model (SSM)¹ for estimating parameters with a high degree of confidence in short time series. In particular, BSSM provides opportunities for managing uncertainty and adapting parameters to observed data. It allows for rapid updating of posterior distributions², enabling better representation of dynamic changes in data over short time series.

BSSM treats all uncertain quantities as random variables and uses the laws of probability. These models are widely used in the fields of signal filtering and prediction. To quantify uncertainties in the model, posterior distributions over all the parameters are required. In particular, the model consists of an observable or measurement equation representing the level of interdependence and a transition equation³ representing dynamic state variables:

$$\begin{aligned}y_t &= Cx_t + Du_t + v_t \\x_t &= Ax_{t-1} + Bu_t + w_t.\end{aligned}$$

Where,

y_t – vector of observable endogenous variables;

x_t – vector of hidden state variables

u_t – vector of observable exogenous variables;

A – $k \times k$ state dynamics matrix;

C – $p \times k$ observation matrix;

B – $k \times d$ input-to-state matrix;

D – $p \times d$ input-to-observation matrix;

$v_t \sim N(0, R)$ – vector of measurement errors for the observation equation;

$w_t \sim N(0, Q)$ – vector of measurement errors for the state equation;

A, B, C and D matrices – matrices of estimated coefficients;

Q and R – covariance matrices for errors.

The displacements in the observation space can be learnt as parameters of the input-to-state and input-to-observation matrices. Since the hidden state is unobserved, state

¹ State space model (SSM) methodology is presented in the Methodological Framework of Central Bank.

² Posterior distribution is the updated distribution of a parameter after observing the data.

³ Beal, M. (2003). Variational Algorithms for Approximate Bayesian Inference. University of Cambridge.

noise covariance matrix Q can be incorporated into the state dynamics matrix A . As a result, hidden state rescaled according to this change. Covariance matrix, R , of output noise can not be incorporated to observation matrix C since the data is observed. In particular, the output noise covariance matrix R , which is assumed diagonal, is defined through the precision vector ρ as follows:

$$R^{-1} = \text{diag}(\rho)$$

Here, R denotes the covariance matrix of output noise, $\text{diag}(p)$ denotes a matrix with all offdiagonal elements equal to zero.

For conjugacy⁴, each dimension of ρ is assumed to be gamma distributed with hyperparameters a and b . Prior distribution of ρ is defined as follows:

$$p(\rho|a, b) = \prod_{s=1}^p \frac{b^a}{\Gamma(a)} \rho_s^{a-1} \exp \{-b\rho_s\}.$$

The parameters and hyperparameters are linked through following probabilities:

$$\begin{aligned} p(a_{(j)}|\alpha) &= N(a_{(j)}|0, \text{diag}(\alpha)^{-1}) \\ p(b_{(j)}|\beta) &= N(b_{(j)}|0, \text{diag}(\beta)^{-1}) \\ &\quad j = 1, 2, \dots, k \\ p(c_{(s)}|\rho_s, \gamma) &= N(c_{(s)}|0, \rho_s^{-1} \text{diag}(\gamma)^{-1}) \\ p(d_{(s)}|\rho_s, \delta) &= N(d_{(s)}|0, \rho_s^{-1} \text{diag}(\delta)^{-1}) \\ p(\rho_s|a, b) &= Ga(\rho_s|a, b) \\ &\quad s = 1, 2, \dots, p. \end{aligned}$$

Where,

$a_{(j)}$ – j th column of matrix A ;

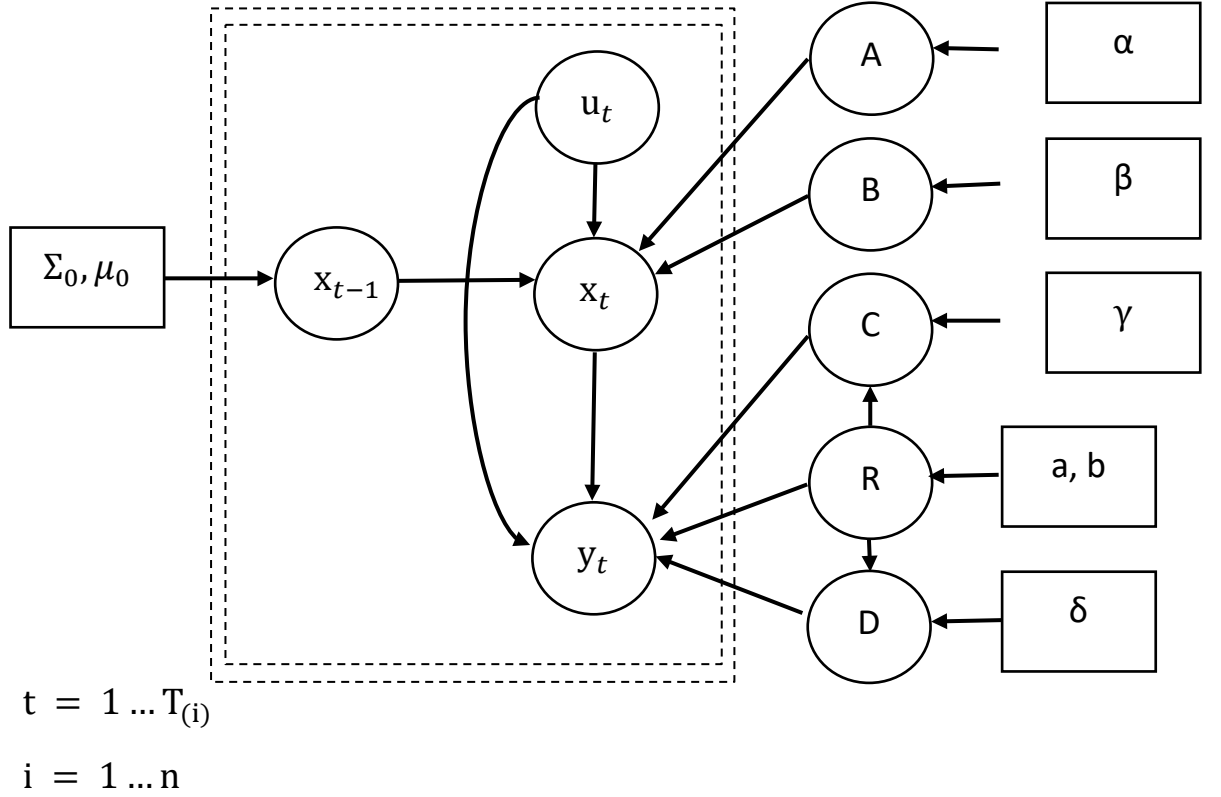
$b_{(j)}$ – j th column of matrix B ;

$c_{(s)}$ – s th column of matrix C ;

$d_{(s)}$ – s th column of matrix D .

⁴ Conjugation is the method of expressing the prior and posterior distributions in the same form to simplify calculations.

Figure 1. BSSM Overview



Source: Beal, M. (2003). Variational Algorithms for Approximate Bayesian Inference. University of Cambridge.

In addition, for each row of matrix A , zero mean Gaussian prior has precision⁵ equal to $diag(\alpha)$. Each row of C matrix is given a zero mean Gaussian prior with precision equal to $diag(\rho s \gamma)$. The dependence of the precision of $c(s)$ on the noise output precision ρs is motivated by conjugacy. Two more hyperparameter vectors β and δ are placed on the rows of the input-related matrices B and D . The covariances across the rows of A and B matrices remain unchanged. However, the covariances within the rows of matrices C and D may change as a result of ρs . If the j th column of A become zero, the j th column at time $t - 1$ is not involved in generating the hidden state at time t . However, j th hidden dimension is used in producing covariance structure in the data at each time step.

BSSM is implemented in six steps. The first step is initialization, in which the auxiliary hidden state is set, prior distribution⁶ of parameters is specified, and the log

⁵ In the distributions, the precision is equal to the inverse of the variance, meaning that the product of the precision and the variance is equal to one.

⁶ Prior distribution is the distribution of the parameters before any data are observed.

marginal probability⁷ is evaluated using Jensen's inequality. The second step is Variational Bayesian M-step (Maximisation). At this step, the optimal forms of the posterior distributions of the parameters are determined and the natural parameter vector is calculated. Next, in the Variational Bayesian E-step (Expectation), the distributions of hidden state variables vector are determined. In the fourth step, the lower bound on the marginal probability, the value of F is computed. The fifth is the hyperparameter updating step, in which hyperparameters are adjusted to maximize lower bound of the marginal probability. In the final step, the lower bound estimate of the log marginal probability is compared with its previous value, and if the lower bound has increased, the Variational Bayesian M-step is repeated. These computational steps are continued iteratively until the lower bound estimate no longer exceeds its previous value.

Initialization

At first, for time $t = 0$, auxiliary hidden state x_0 is initialized. In this context, the auxiliary hidden state x_0 follows Gaussian distribution with mean μ_0 , covariance Σ_0 , and prior distribution is defined as follows:

$$\begin{aligned} p(x_0 | \mu_0, \Sigma_0) &= N(x_0 | \mu_0, \Sigma_0) \\ \mu_0 &\sim N(\mu_0 | 0, b_{\mu_0} I) \\ \Sigma_0 &\sim \prod_{j=1}^k Ga(\Sigma_{0jj}^{-1} | a_{\Sigma_0}, b_{\Sigma_0}) \end{aligned}$$

Where,

x_0 – auxiliary hidden state;

μ_0 – mean of auxiliary hidden state;

Σ_0 – covariance of auxiliary hidden state.

In this case, the prior of x_1 via the state dynamics process is denoted by the following:

$$\begin{aligned} p(x_1 | \mu_0, \Sigma_0, \theta) &= \int dx_0 p(x_0 | \mu_0, \Sigma_0) p(x_1 | x_0, \theta) \\ &= N(x_1 | A\mu_0 + Bu_1, A^T \Sigma_0 A + Q) . \end{aligned}$$

⁷ Marginal probability is the likelihood of an event occurring independently of other events.

For simplicity, the parameters are integrated into a single parameter vector $\theta = (A, B, C, D, R)$. In addition, the dependence of A, B, C and D matrices with hyperparameters is given by the following hyperprior distribution:

$$\begin{aligned}\alpha &\sim \prod_{j=1}^k Ga(\alpha_{(j)} | a_{(\alpha)}, b_{(\alpha)}) \\ \beta &\sim \prod_{c=1}^d Ga(\beta_{(c)} | a_{(\beta)}, b_{(\beta)}) \\ \gamma &\sim \prod_{j=1}^k Ga(\gamma_{(j)} | a_{(\gamma)}, b_{(\gamma)}) \\ \delta &\sim \prod_{c=1}^d Ga(\delta_{(c)} | a_{(\delta)}, b_{(\delta)})\end{aligned}$$

The marginal probability takes the following form:

$$p(y_{1:T}) = \int dA dB dC dD d\rho dx_{0:T} p(A, B, C, D, \rho, x_{0:T}, y_{1:T})$$

Due to some conditional independencies between the parameters of the model, it is possible to see how posterior distributions for the dynamics and output processes decompose. The full joint probability for parameters, hidden variables and observed data, taking into account the inputs, is given as follows:

$$\begin{aligned}p(A, B, C, D, \rho, x_{0:T}, y_{1:T} | u_{1:T}) = \\ p(A | \alpha) p(B | \beta) p(\rho | a, b) p(C | \rho, \gamma) p(D | \rho, \delta) p(x_0 | \mu_0, \Sigma_0) \cdot \\ \prod_{t=1}^T p(x_t | x_{t-1}, A, B, u_t) p(y_t | x_t, C, D, \rho, u_t)\end{aligned}$$

Next, the dependence on the input sequence $u_{1:T}$ is dropped and inputs is left as implicit. A distribution $q(\theta, x)$ over the parameters and hidden variables is introduced. By applying Jensen's inequality lower bound the marginal probability is estimated:

$$\begin{aligned}\ln p(y_{1:T}) = \ln \left(\int dA dB dC dD d\rho dx_{0:T} p(A, B, C, D, \rho, x_{0:T}, y_{1:T}) \right) \geq \\ \int dA dB dC dD d\rho dx_{0:T} q(A, B, C, D, \rho, x_{0:T}) \ln \frac{p(A, B, C, D, \rho, x_{0:T}, y_{1:T})}{q(A, B, C, D, \rho, x_{0:T})} = F\end{aligned}$$

The next step in the variational approximation is to assume some approximate form for the distribution $q(\cdot)$. First, hidden variables and parameters are separated based on conditional independency:

$$\begin{aligned}q(A, B, C, D, \rho, x_{0:T}) &= q_\theta(A, B, C, D, \rho) q_x(x_{0:T}) \\ q(A, B, C, D, \rho, x_{0:T}) &= q_{AB}(A, B) q_{CD\rho}(C, D, \rho) q_x(x_{0:T}).\end{aligned}$$

To simplify the estimation of the lower bound of the log marginal probability, the following terms of conditional probability are used:

$$q_{AB}(A, B) = q_B(B) q_A(A|B)$$

$$q_{CD\rho}(C, D, \rho) = q_\rho(\rho) q_D(D|\rho) q_C(C|D, \rho)$$

Using the above terms, the lower bound of the log marginal probability, F is expressed as the following sum:

$$\begin{aligned} F = & \int dB q_B(B) \ln \frac{p(B|\beta)}{q_B(B)} + \int dB q_B(B) \int dA q_A(A|B) \ln \frac{p(A|\alpha)}{q_A(A|B)} + \\ & \int d\rho q_\rho(\rho) \ln \frac{p(\rho|a, b)}{q_\rho(\rho)} + \int d\rho q_\rho(\rho) \int dD q_D(D|\rho) \ln \frac{p(D|\rho, \delta)}{q_D(D|\rho)} + \\ & \int d\rho q_\rho(\rho) \int dD q_D(D|\rho) \int dC q_C(C|\rho, D) \ln \frac{p(C|\rho, \gamma)}{q_C(C|\rho, D)} - \\ & \int dx_{0:T} q_x(x_{0:T}) \ln q_x(x_{0:T}) + \int dB q_B(B) \int dA q_A(A|B) \int d\rho q_\rho(\rho) \cdot \\ & \int dD q_D(D|\rho) \int dC q_C(C|\rho, D) \int dx_{0:T} q_x(x_{0:T}) \ln p(x_{0:T}, y_{1:T} | A, B, C, D, \rho) = \\ & F(q_x(x_{0:T}), q_B(B), q_A(A|B), q_\rho(\rho), q_D(D|\rho), q_C(C|\rho, D)) \\ & H(q_x(x_{0:T})) = - \int dx_{0:T} q_x(x_{0:T}) \ln q_x(x_{0:T}) \end{aligned}$$

The lower bound of the log marginal probability depends on the hyperparameters. The optimum forms of these approximate posteriors can be found by setting functional derivatives of F equal to zero with respect to each distribution over parameters and hidden variable sequences.

Variational Bayesian M-step (Maximisation), VBM

In this step, the variational posterior distributions over the parameters are determined by applying Bayesian theorem and from these the expected natural parameter vector is computed. The posterior distribution is factorized in the following form:

$$q_\theta(A, B, C, D, \rho) = \prod_{j=1}^k q(b_{(j)}) q(a_{(j)} | b_{(j)}) \prod_{s=1}^p q(\rho_s) \cdot q(d_{(s)} | \rho_s) q(c_{(s)} | \rho_s, d_{(s)})$$

Where,

$b_{(j)}$ – j th row of B matrix;

$a_{(j)}$ – j th row of A matrix;

$d_{(s)}$ – s th row of D matrix;

$c_{(s)}$ – s th row of C matrix;

Some statistics of the input and output variable vectors, required for the calculations, is defined as follows:

$$\begin{aligned}\ddot{U} &= \sum_{t=1}^T u_t u_t^T \\ U_y &= \sum_{t=1}^T u_t y_t^T \\ \ddot{Y} &= \sum_{t=1}^T y_t y_t^T\end{aligned}$$

The sufficient statistics used to compute the posterior distributions of the parameters and the natural parameter vector are obtained using the following equations:

$$\begin{aligned}W_A &= \sum_{t=1}^T \langle x_{t-1} x_{t-1}^T \rangle = \sum_{t=1}^T Y_{t-1,t-1} + \omega_{t-1} \omega_{t-1}^T \\ G_A &= \sum_{t=1}^T \langle x_{t-1} \rangle u_t^T = \sum_{t=1}^T \omega_{t-1} u_t^T \\ \tilde{M} &= \sum_{t=1}^T u_t \langle x_t^T \rangle = \sum_{t=1}^T u_t \omega_t^T \\ S_A &= \sum_{t=1}^T \langle x_{t-1} x_t^T \rangle = \sum_{t=1}^T Y_{t-1,t} + \omega_{t-1} \omega_t^T \\ W_C &= \sum_{t=1}^T \langle x_t x_t^T \rangle = \sum_{t=1}^T Y_{t,t} + \omega_t \omega_t^T \\ G_C &= \sum_{t=1}^T \langle x_t \rangle u_t^T = \sum_{t=1}^T \omega_t u_t^T \\ S_C &= \sum_{t=1}^T \langle x_t \rangle y_t^T = \sum_{t=1}^T \omega_t y_t^T.\end{aligned}$$

The posterior distributions for A and B are defined by the following equations:

$$\begin{aligned}q_B(B) &= \prod_{j=1}^k N(b_{(j)} | \Sigma_B \bar{b}_{(j)}, \Sigma_B) \\ q_A(A|B) &= \prod_{j=1}^k N(a_{(j)} | \Sigma_A [S_{A(j)} - G_A b_{(j)}], \Sigma_A) \\ \Sigma_A^{-1} &= \text{diag}(\alpha) + W_A \\ \Sigma_B^{-1} &= \text{diag}(\beta) + \ddot{U} - G_A^T \Sigma_A G_A \\ \bar{B} &= \hat{M}^T - S_A^T \Sigma_A G_A \\ q_A(A) &= \prod_{j=1}^k N(a_{(j)} | \Sigma_A [S_{A(j)} - G_A \Sigma_B \bar{b}_{(j)}], \tilde{\Sigma}_A) \\ \tilde{\Sigma}_A &= \Sigma_A + \Sigma_A G_A \Sigma_B G_A^T \Sigma_A\end{aligned}$$

Here, $\bar{b}_{(j)}$ denotes j th row of $\bar{B}$ matrix, while $S_{A(j)}$ defines j th column of S_A matrix.

Additionally, the posterior distributions over ρ , C and D are determined by the following:

$$\begin{aligned}
q_\rho(\rho) &= \prod_{s=1}^p Ga(\rho_s | a + \frac{T}{2}, b + \frac{1}{2}G_{ss}) \\
q_D(D|\rho) &= \prod_{s=1}^p N(d_{(s)} | \Sigma_D \bar{d}_{(s)}, \rho_s^{-1} \Sigma_D) \\
q_C(C|D, \rho) &= \prod_{s=1}^p N(c_{(s)} | \Sigma_C [S_{C(s)} - G_C d_{(s)}], \rho_s^{-1} \Sigma_C) \\
\Sigma_C^{-1} &= diag(\gamma) + W_C \\
\Sigma_D^{-1} &= diag(\delta) + \ddot{U} - G_C^T \Sigma_C G_C \\
G &= \ddot{Y} - S_C^T \Sigma_C S_C - \bar{D} \Sigma_D \bar{D}^T \\
\bar{D} &= U_Y^T - S_C^T \Sigma_C G_C \\
q_C(C|\rho) &= \prod_{s=1}^p N(c_{(s)} | \Sigma_C [S_{C(s)} - G_C \Sigma_D \bar{d}_{(s)}], \rho_s^{-1} \hat{\Sigma}_C) \\
\hat{\Sigma}_C &= \Sigma_C + \Sigma_C G_C \Sigma_D G_C^T \Sigma_C
\end{aligned}$$

Here, $\bar{d}_{(s)}$ and $S_{C(s)}$ denote sth row of $\bar{D}$ and sth column of S_C matrices, respectively.

In the Variational Bayesian M-step (Maximisation), $\varphi(\theta)$ is considered the vector of expected natural parameters. These will then be used in the Variational Bayesian E-step (Expectation) which infers the distribution $q_x(x_{0:T})$ over hidden states in the system. The vector of relevant natural parameters is given by the following form:

$$\begin{aligned}
\varphi(\theta) &= \varphi(A, B, C, D, R) = [A, A^T A, B, A^T B, C^T R^{-1} C, \\
&R^{-1} C, C^T R^{-1} D, B^T B, R^{-1}, \ln|R^{-1}|, D^T R^{-1} D, R^{-1} D] . \\
\langle A \rangle &= [S_A - G_A \Sigma_B \bar{B}^T]^T \Sigma_A \\
\langle A^T A \rangle &= \langle A \rangle^T \langle A \rangle + k [\Sigma_A + \Sigma_A G_A \Sigma_B G_A^T \Sigma_A] \\
\langle B \rangle &= \bar{B} \Sigma_B \\
\langle A^T B \rangle &= \Sigma_A [S_A \langle B \rangle - G_A \{ \langle B \rangle^T \langle B \rangle + k \Sigma_B \}] \\
\langle B^T B \rangle &= \langle B \rangle^T \langle B \rangle + k \Sigma_B \\
\langle \rho_s \rangle &= \bar{\rho}_s = \frac{a_\rho + T/2}{b_\rho + G_{ss}/2} \\
\langle \ln \rho_s \rangle &= \overline{\ln \rho_s} = \Psi\left(a_\rho + \frac{T}{2}\right) - \ln(b_\rho + \frac{G_{ss}}{2}) \\
\langle R^{-1} \rangle &= diag(\bar{\rho}) \\
\langle C \rangle &= [S_C - G_C \Sigma_D \bar{D}^T]^T \Sigma_C \\
\langle D \rangle &= \bar{D} \Sigma_D \\
\langle C^T R^{-1} C \rangle &= \langle C \rangle^T diag(\bar{\rho}) \langle C \rangle + p [\Sigma_C + \Sigma_C G_C \Sigma_D G_C^T \Sigma_C] \\
\langle R^{-1} C \rangle &= diag(\bar{\rho}) \langle C \rangle
\end{aligned}$$

$$\begin{aligned}
\langle C^T R^{-1} D \rangle &= \Sigma_C [S_C \text{diag}(\bar{\rho}) \langle D \rangle - G_C \langle D \rangle^T \text{diag}(\bar{\rho}) \langle D \rangle - p G_C \Sigma_D] \\
\langle R^{-1} D \rangle &= \text{diag}(\bar{\rho}) \langle D \rangle \\
\langle D^T R^{-1} D \rangle &= \langle D \rangle^T \text{diag}(\bar{\rho}) \langle D \rangle + p \Sigma_D.
\end{aligned}$$

Here, $\langle \cdot \rangle$ denotes expectation with respect to the variational posterior.

Variational Bayesian E-step (Expectation), VBE

The posterior distribution of the hidden state variable vector $x_{0:T}$ is jointly Gaussian over the time steps, and the logarithm of the posterior distribution of the hidden state vector is expressed as follows:

$$\begin{aligned}
\ln q_x(x_{0:T}) &= -\ln Z + \langle \ln p(A, B, C, D, \rho, x_{0:T}, y_{1:T}) \rangle \\
&= -\ln Z' + \langle \ln p(x_{0:T}, y_{1:T} | A, B, C, D, \rho) \rangle \\
Z' &= \int dx_{0:T} \exp \langle \ln p(x_{0:T}, y_{1:T} | A, B, C, D, \rho) \rangle. \\
\ln Z' &= \sum_{t=1}^T \ln \zeta'_t(y_t) \\
\ln \zeta'_t(y_t) &= -\frac{1}{2} [\langle \ln |2\pi R| \rangle - \ln |\Sigma_{t-1}^{-1} \Sigma_{t-1}^* \Sigma_t| + \mu_{t-1}^T \Sigma_{t-1}^{-1} \mu_{t-1} - \mu_t^T \Sigma_t^{-1} \mu_t \\
&\quad + y_t^T \langle R^{-1} \rangle y_t - 2 y_t^T \langle R^{-1} D \rangle u_t + u_t^T \langle D^T R^{-1} D \rangle u_t \\
&\quad - (\Sigma_{t-1}^{-1} \mu_{t-1} - \langle A^T B \rangle u_t)^T \Sigma_{t-1}^* (\Sigma_{t-1}^{-1} \mu_{t-1} - \langle A^T B \rangle u_t)].
\end{aligned}$$

$\alpha_t(x_t)$ is posterior over the hidden state at time t given observed data up to time t ($t = \{1, 2, \dots, T\}$). At the same time, filtered mean μ_t and covariance Σ_t of hidden state $\alpha_t(x_t)$ are estimated:

$$\alpha_t(x_t) \equiv p(x_t | y_{1:t}).$$

In addition, recursion of $\alpha_t(x_t)$ with previous $\alpha_{t-1}(x_{t-1})$ is formed as following form:

$$\begin{aligned}
\alpha_t(x_t) &= \int dx_{t-1} \frac{p(x_{t-1} | y_{1:t-1}) p(x_t | x_{t-1}) p(y_t | x_t)}{p(y_t | y_{1:t-1})} \\
&= \frac{1}{\zeta_t(y_t)} \int dx_{t-1} \alpha_{t-1}(x_{t-1}) p(x_t | x_{t-1}) p(y_t | x_t) \\
&= \frac{1}{\zeta_t(y_t)} \int dx_{t-1} N(x_{t-1} | \mu_{t-1}, \Sigma_{t-1}) N(x_t | A x_{t-1}, I) N(y_t | C x_t, R) \\
&= N(x_t | \mu_t, \Sigma_t)
\end{aligned}$$

Here, $\zeta_t(y_t) \equiv p(y_t | y_{1:t-1})$ is the filtered output probability.

$$\begin{aligned}\Sigma_{t-1}^* &= (\Sigma_{t-1}^{-1} + A^T A)^{-1} \\ x_{t-1}^* &= \Sigma_{t-1}^* [\Sigma_{t-1}^{-1} \mu_{t-1} + A^T x_t] \\ \Sigma_t &= [I + C^T R^{-1} C - A \Sigma_{t-1}^* A^T]^{-1} \\ \mu_t &= \Sigma_t [C^T R^{-1} y_t + A \Sigma_{t-1}^* \Sigma_{t-1}^{-1} \mu_{t-1}].\end{aligned}$$

At each step the normalising constant $\zeta_t(y_t)$ obtained as the denominator of $\alpha_t(x_t)$, contributes to the calculation of the probability of the data. In this context, the probability of the data can be written in the following form:

$$\begin{aligned}p(y_{1:T}) &= p(y_1) p(y_2 | y_1) \dots p(y_t | y_{1:t-1}) \dots p(y_T | y_{1:T-1}) = \\ &= p(y_1) \prod_{t=2}^T p(y_t | y_{1:t-1}) = \prod_{t=1}^T \zeta_t(y_t) \\ \zeta_t(y_t) &= N(y_t | \varpi_t, \xi_t) \\ \xi_t &= (R^{-1} - R^{-1} C \Sigma_t C^T R^{-1})^{-1} \\ \varpi_t &= \xi_t R^{-1} C \Sigma_t A \Sigma_{t-1}^* \Sigma_{t-1}^{-1} \mu_{t-1}.\end{aligned}$$

With these distributions, the probability of each observation y_t given the previous observations in the sequence is computed, and predictive mean and variance are assigned to the data at each time step as it arrives. However, this predictive distribution will change once the hidden state sequence has been smoothed. To obtain the posterior over the hidden state given all the data, the smoothed estimate γ of the hidden state is determined. In particular, $\gamma_t(x_t) = p(x_t | y_{1:T})$, the smoothed estimate of the hidden state, is expressed as follows:

$$\begin{aligned}\gamma_t(x_t) &= p(x_t | y_{1:T}) = \int dx_{t+1} p(x_t, x_{t+1} | y_{1:T}) \\ &= \int dx_{t+1} p(x_t | x_{t+1}, y_{1:T}) p(x_{t+1} | y_{1:T}) \\ &= \int dx_{t+1} p(x_t | x_{t+1}, y_{1:t}) p(x_{t+1} | y_{1:T}) \\ &= \int dx_{t+1} \frac{p(x_t | y_{1:t}) p(x_{t+1} | x_t)}{\int dx'_t p(x'_t | y_{1:t}) p(x_{t+1} | x'_t)} p(x_{t+1} | y_{1:T}) \\ &= \int dx_{t+1} \left[\frac{\alpha_t(x_t) p(x_{t+1} | x_t)}{\int dx'_t \alpha_t(x'_t) p(x_{t+1} | x'_t)} \right] \gamma_{t+1}(x_{t+1}) \\ \gamma_t(x_t) &= N(x_t | \omega_t, Y_{tt}) \\ K_t &= (Y_{t+1,t+1}^{-1} + A \Sigma_t^* A^T)^{-1}\end{aligned}$$

$$\begin{aligned} Y_{tt} &= [\Sigma_t^{*-1} - A^T K_t A]^{-1} \\ \omega_t &= Y_{tt} [\Sigma_t^{-1} \mu_t + A^T K_t (Y_{t+1,t+1}^{-1} \omega_{t+1} - A \Sigma_t^* \Sigma_t^{-1} \mu_t)] . \end{aligned}$$

In this case, once the data at time t , y_t has been filtered, $\alpha_t(x_t)$ is used instead of hidden variables.

To estimate the states at different time steps, parallel recursion $\beta_t(x_t) = p(y_{t+1:T} | x_t)$ is calculated:

$$\begin{aligned} \beta_{t-1}(x_{t-1}) &= \int dx_t p(x_t | x_{t-1}) p(y_t | x_t) p(y_{t+1:T} | x_t) = \\ &\int dx_t p(x_t | x_{t-1}) p(y_t | x_t) \beta_t(x_t) \propto N(x_{t-1} | \eta_{t-1}, \Psi_{t-1}) . \end{aligned}$$

$\beta_T(x_T) = 1$ is defined in $\Psi_T^{-1} = 0$ form to satisfy the end condition. The recursion proceeds from $t = T$ to $t = 1$. At the last step of this recursion, the probability of all data on the auxiliary hidden state variable x_0 is determined:

$$\begin{aligned} \Psi_t^* &= (I + C^T R^{-1} C + \Psi_t^{-1})^{-1} \\ \Psi_{t-1} &= [A^T A - A^T \Psi_t^* A]^{-1} \\ \eta_{t-1} &= \Psi_{t-1} A^T \Psi_t^* [C^T R^{-1} y_t + \Psi_t^{-1} \eta_t] . \end{aligned}$$

Using the formulas presented above, $\{\eta_t, \Psi_t\}_{t=0}^T$ is computed over the interval from $t = 0$ to $t = T$.

In the Variational Bayesian scenario the marginals cannot be obtained easily with backward sequential pass and instead they are computed by combining α and β messages described as follows. In particular, ω_t and $Y_{t,t}$ are computed for $t = \{0, 1, 2, \dots, T-1\}$ in the following form:

$$\begin{aligned} p(x_t | y_{1:T}) &\propto p(x_t | y_{1:t}) p(y_{t+1:T} | x_t) = \\ \alpha_t(x_t) \beta_t(x_t) &= N(x_t | \omega_t, Y_{tt}) \\ Y_{t,t} &= [\Sigma_t^{-1} + \Psi_t^{-1}]^{-1} \\ \omega_t &= Y_{t,t} [\Sigma_t^{-1} \mu_t + \Psi_t^{-1} \eta_t] . \end{aligned}$$

For $t = \{0, 1, 2, \dots, T-1\}$ taking into account all observations, the cross-covariance of the hidden variable vector, $Y_{t,t+1}$, is defined as follows:

$$Y_{t,t+1} = \Sigma_t^* A^T Y_{t+1,t+1}.$$

Calculation of the lower bound on the log marginal probability

Due to the complexity of directly computing the log marginal probability, its lower bound is evaluated and denoted by F as the estimate of the lower bound of the log marginal probability. As this estimate of the lower bound of the log marginal probability is itself a complex function, it is derived through several steps and ultimately expressed as a simplified summ.

KL (Kullback-Leibler) is divergence between two normal or two gamma distributions over the same variables and computed using the following equations for a pair of variables J and K :

$$\begin{aligned} KL(J) &= \int dJ q(J) \ln \frac{q(J)}{p(J)} \\ KL(J|K) &= \int dJ q(J|K) \ln \frac{q(J|K)}{p(J|K)} \\ < KL(J|K) >_{q(K)} &= \int dK q(K) KL(J|K) \end{aligned}$$

Where,

$KL(J)$ – KL divergence,

$KL(J|K)$ – conditional KL ,

$< KL(J|K) >_{q(K)}$ – expected conditional KL .

By applying these equations, F can be redefined in the following form:

$$\begin{aligned} F &= -KL(B) - < KL(A|B) >_{q(B)} - KL(\rho) - < KL(D|\rho) >_{q(\rho)} - \\ &\quad < KL(C|\rho, D) >_{q(\rho, D)} + H(q_x(x_{0:T})) + \\ &\quad < \ln p(x_{1:T}, y_{1:T} | A, B, C, D, \rho) >_{q(A, B, C, D, \rho)q(x_{1:T})} \cdot \\ H(q_x(x_{0:T})) &= - \int dx_{0:T} q_x(x_{0:T}) \ln q_x(x_{0:T}) = \\ - \int dx_{0:T} q_x(x_{0:T}) &[- \ln Z' + < \ln p(x_{0:T}, y_{1:T} | A, B, C, D, \rho, \mu_0, \Sigma_0) >_{q_\theta(A, B, C, D, \rho)}] \\ &= \ln Z' - < \ln p(x_{0:T}, y_{1:T} | A, B, C, D, \rho, \mu_0, \Sigma_0) >_{q_\theta(A, B, C, D, \rho)q_x(x_{0:T})} \end{aligned}$$

Using $H(q_x(x_{0:T}))$, the equation for F is transformed into the following simplified form, which facilitates its computation:

$$F = -KL(B) - \langle KL(A|B) \rangle_{q(B)} - KL(\rho) - \langle KL(D|\rho) \rangle_{q(\rho)} - \langle KL(C|\rho, D) \rangle_{q(\rho, D)} + \ln Z'.$$

Updating hyperparameters

The hyperparameters α , β , γ , δ , a , b , and the prior parameters Σ_0 and μ_0 , are updated to maximise the lower bound on the marginal probability. In this context, hyperparameters is updated as follows:

$$\begin{aligned} \alpha_j^{-1} &\leftarrow \frac{1}{k} [k\Sigma_A + \Sigma_A [S_A S_A^T - 2G_A \langle B \rangle^T S_A^T + \\ &\quad G_A \{k\Sigma_B + \langle B \rangle^T \langle B \rangle\} G_A^T] \Sigma_A]_{jj} \\ \beta_j^{-1} &\leftarrow \frac{1}{k} [k\Sigma_B + \langle B \rangle^T \langle B \rangle]_{jj} \\ \gamma_j^{-1} &\leftarrow \frac{1}{p} [p\Sigma_C + \Sigma_C [S_C \text{diag}(\bar{\rho}) S_C^T - 2S_C \text{diag}(\bar{\rho}) \langle D \rangle G_C^T + pG_C \Sigma_D G_C' + \\ &\quad G_C \langle D \rangle^T \text{diag}(\bar{\rho}) \langle D \rangle G_C^T] \Sigma_C]_{jj} \\ \delta_j^{-1} &\leftarrow \frac{1}{p} [p\Sigma_D + \langle D \rangle^T \text{diag}(\bar{\rho}) \langle D \rangle]_{jj} \end{aligned}$$

Here, $[\cdot]_{jj}$ denotes (j, j) th element of matrix $[\cdot]$.

In order to maximise the probability of the hidden state sequence under the prior, the hyperparameters of the prior over the auxiliary hidden state are set according to the distribution of the smoothed estimate of x_0 :

$$\begin{aligned} \Sigma_0 &\leftarrow Y_{0,0} \\ \mu_0 &\leftarrow \omega_0 \end{aligned}$$

The hyperparameters a and b governing the prior distribution over the output noise, $R = \text{diag}(\rho)$, are set to the fixed point of the equations:

$$\begin{aligned}\Psi(a) &= \ln b + \frac{1}{p} \sum_{s=1}^p \overline{\ln \rho_s} \\ \frac{1}{b} &= \frac{1}{pa} \sum_{s=1}^p \bar{\rho}_s\end{aligned}$$

Comparison of the lower bound estimates on the log marginal probability

As the estimate of the lower bound of the log marginal probability is determined, in the final step, this estimate is compared with its value from the previous iteration. If the lower bound has increased relative to its value in the preceding iteration, the computations are repeated starting from the Variational Bayesian M-step. This iterative process continues until the lower bound no longer increases. When the estimate of the lower bound on the log marginal probability remains unchanged from its previous value, it is considered to have reached its optimal value. At this stage, the parameters obtained through the hyperparameters also achieve their optimal values. Based on the identified parameters and the hidden variables, the observable endogenous variable Y is estimated.