

Chapter Six

State Space Models and Filtering Techniques

Tingguo Zheng (zhengtg@gmail.com)

School of Economics & WISE & Chow, Xiamen University

(For Chow & SOE & WISE Msc and PhD Program, 2026-2027, Autumn Semester)

- **Reference:**

- “Analysis of Financial Time Series”, 2010, Tsay, Chapter 11.
- Hamilton, 1994 “State-Space Models”, Handbook of Econometrics, Volume 4, chapter 50.

- **Outline:**

- State space representation
- Kalman filtering for linear Gaussian SS models
- Approximate filtering for nonlinear SS models
- Mixed-frequency model and applications
- Sequential Monte Carlo filtering*
- Score-driven models

The main applications of state space models in macroeconomics and finance have dealt with the following topics.

- The extraction of signals such as trends and cycles in macroeconomic time series: see Watson (1986), Clark (1987), Harvey and Jaeger (1993), Hodrick and Prescott (1997), Morley, Nelson and Zivot (2003), Proietti (2006), Luati and Proietti (2011).
- Dynamic factor models, for the extraction of a single index of coincident indicators, see Stock and Watson (1989), Frale et al. (2011), and for large dimensional systems (Jungbacker, Koopman and van der Wel, 2011).

- Stochastic volatility models: see Shephard (2005) and Stock and Watson (2007) for applications to US inflation.
- The class of dynamic stochastic general equilibrium (DSGE) models: Sargent (1989), Fernandez-Villaverde and Rubio-Ramirez (2005), Smets and Wouters (2003), Fernandez-Villaverde (2010).
- Term structure of interest rate: Diebold and Li (2006)
- Mixed-frequency data models: mixed-frequency dynamic factor model by Mariano and Murasawa (2003), and mixed-frequency VAR model by Mariano and Murasawa (2010).

1 Linear Gaussian State Space Models

1.1 State space expression

The general representation of linear Gaussian state space models (Harvey, 1991) can be formulated as:

$$\underset{n \times 1}{\mathbf{y}_t} = \underset{n \times r}{\mathbf{H}_t} \underset{r \times 1}{\boldsymbol{\xi}_t} + \underset{n \times 1}{\mathbf{d}_t} + \underset{n \times 1}{\mathbf{u}_t} \quad (1)$$

$$\underset{r \times 1}{\boldsymbol{\xi}_t} = \underset{r \times r}{\mathbf{F}_t} \underset{r \times 1}{\boldsymbol{\xi}_{t-1}} + \underset{r \times 1}{\mathbf{c}_t} + \underset{r \times 1}{\mathbf{v}_t} \quad (2)$$

where the disturbances $\mathbf{u}_t$ and $\mathbf{v}_t$ are such that:

$$\underset{n \times n}{\mathbf{u}_t} \sim N(\mathbf{0}, \underset{n \times n}{\mathbf{R}_t}), \quad \underset{r \times r}{\mathbf{v}_t} \sim N(\mathbf{0}, \underset{r \times r}{\mathbf{Q}_t}), \quad \text{Cov}(\mathbf{u}_t \mathbf{v}_t) = \mathbf{0}. \quad (3)$$

Equation (1) is the *observation equation* or *measurement equation*; and equation (2) is the *state equation* or *transition equation*.

It should be noticed that:

1. The covariance matrices R_t and Q_t may be not full rank and thus singular;
2. It is often assumed that the innovations u_t and v_t are not correlated. If they are correlated, the reader are referred to another state space form advocated by Koopman, Shephard and Doornik (1999, Statistical algorithm for models in state space using SsfPack 2.2) and Durbin and Koopman (2002)'s book.

Time-invariant parametric models

In particular, if all matrices do not depend deterministically on t the state space system is called *time invariant*. This can be rewritten as

$$\underset{n \times 1}{\mathbf{y}_t} = \underset{n \times r}{\mathbf{H}} \underset{r \times 1}{\boldsymbol{\xi}_t} + \underset{n \times 1}{\mathbf{d}} + \underset{n \times 1}{\mathbf{u}_t} \quad (4)$$

$$\underset{r \times 1}{\boldsymbol{\xi}_t} = \underset{r \times r}{\mathbf{F}} \underset{r \times 1}{\boldsymbol{\xi}_{t-1}} + \underset{r \times 1}{\mathbf{c}} + \underset{r \times 1}{\mathbf{v}_t} \quad (5)$$

$$\text{Cov}(\mathbf{u}_t \mathbf{u}_t) = \underset{n \times n}{\mathbf{R}}, \quad \text{Cov}(\mathbf{v}_t \mathbf{v}_t) = \underset{r \times r}{\mathbf{Q}}, \quad \text{Cov}(\mathbf{u}_t \mathbf{v}_t) = \mathbf{0}.$$

- The state space representation is completed by specifying the behavior of the *initial state* $\boldsymbol{\xi}_0 \sim N(\mathbf{a}_0, \mathbf{P}_0)$.
- Under the weak stationary assumption of $\mathbf{y}_t$, the *unconditional mean* of $\boldsymbol{\xi}_t$, $\mathbf{a}_0$, may be determined using $\mathbf{a}_0 = E[\boldsymbol{\xi}_t] = (\mathbf{I}_r - \mathbf{F})^{-1} \mathbf{c}$.

- Similarly, the *unconditional variance* of ξ_t , $\text{Var}(\xi_0)$ may be determined analytically using

$$\text{vec}(\mathbf{P}_0) = (\mathbf{I}_{r^2} - \mathbf{F} \otimes \mathbf{F})^{-1} \text{vec}(\mathbf{Q}).$$

Here $\text{vec}(\mathbf{P}_0)$ is an $r^2 \times 1$ column vector. It can be easily *reshaped to form* the $r \times r$ matrix $\mathbf{P}_0$.

Examples

(1) AR(p) process

Consider the AR(p) process

$$y_t = \mu + x_t, \quad x_t = \phi_1 x_{t-1} + \cdots + \phi_p x_{t-p} + \varepsilon_t, \quad \varepsilon_t \sim \mathcal{N}(0, \sigma_\varepsilon^2)$$

One way to represent AR(p) in state space form is as follows. Define $\xi_t = (x_t, x_{t-1}, \dots, x_{t-p+1})'$ so that the transition equation for ξ_t becomes

$$\begin{bmatrix} x_t \\ x_{t-1} \\ \vdots \\ x_{t-p+1} \end{bmatrix} = \begin{bmatrix} \phi_1 & \cdot & \cdot & \phi_{p-1} & \phi_p \\ 1 & \cdot & \cdot & 0 & 0 \\ \cdot & \cdot & \vdots & \ddots & \cdot \\ 0 & \cdot & \cdot & 1 & 0 \end{bmatrix} \begin{bmatrix} x_{t-1} \\ x_{t-2} \\ \vdots \\ x_{t-p} \end{bmatrix} + \begin{bmatrix} \varepsilon_t \\ 0 \\ \vdots \\ 0 \end{bmatrix}$$

The transition equation system matrices are

$$\mathbf{F} = \begin{bmatrix} \phi_1 & \cdot & \cdot & \phi_{p-1} & \phi_p \\ 1 & \cdot & \cdot & 0 & 0 \\ \cdot & \cdot & \vdots & \ddots & \cdot \\ 0 & \cdot & \cdot & 1 & 0 \end{bmatrix}, \quad \mathbf{c} = \mathbf{0}_{p \times 1}, \quad \mathbf{Q} = \begin{bmatrix} \sigma_\varepsilon^2 & 0 & \cdots & 0 \\ 0 & 0 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & 0 \end{bmatrix}.$$

The measurement equation is $y_t = \mu + (1, 0, \dots, 0)\boldsymbol{\xi}_t$, which implies that

$$\mathbf{H} = (1, 0, \dots, 0), \quad \mathbf{d} = \mu, \quad \mathbf{R} = 0.$$

(2) ARMA(p, q) model

The general ARMA(p, q) model

$$y_t = \phi_1 y_{t-1} + \cdots + \phi_p y_{t-p} + \varepsilon_t + \theta_1 \varepsilon_{t-1} + \cdots + \theta_q \varepsilon_{t-q}, \quad \varepsilon_t \sim N(0, \sigma_\varepsilon^2)$$

may be put in a state space form in the following way. Let $m = \max(p, q + 1)$ and rewrite the ARMA(p, q) model as

$$y_t = \phi_1 y_{t-1} + \cdots + \phi_m y_{t-m} + \varepsilon_t + \theta_1 \varepsilon_{t-1} + \cdots + \theta_{m-1} \varepsilon_{t-m+1},$$

where some of the AR or MA coefficients will be zero unless $p = q + 1$.

Harvey (1994, Section 4.4) provides a state-space form with an m -dimensional state vector ξ_t , the first element of which is y_t , that is $\xi_{1t} = y_t$. The other elements of ξ_t are obtained recursively.

For the ARMA($m, m - 1$) model, we have

$$\begin{aligned}\xi_{1,t+1} = y_{t+1} &= \phi_1 y_t + \sum_{j=2}^m \phi_j y_{t+1-j} + \sum_{j=1}^{m-1} \theta_j \varepsilon_{t+1-j} + \varepsilon_{t+1} \\ &\equiv \phi_1 \xi_{1t} + \xi_{2t} + \eta_t,\end{aligned}$$

where $\xi_{2t} = \sum_{j=2}^m \phi_j y_{t+1-j} + \sum_{j=1}^{m-1} \theta_j \varepsilon_{t+1-j}$, $\eta_t = \varepsilon_{t+1}$, and as defined earlier $\xi_{1t} = y_t$.

Focusing on $\xi_{2,t+1}$, we have

$$\begin{aligned}\xi_{2,t+1} &= \sum_{j=2}^m \phi_j y_{t+2-j} + \sum_{j=1}^{m-1} \theta_j \varepsilon_{t+2-j} \\ &= \phi_2 y_t + \sum_{j=3}^m \phi_j y_{t+2-j} + \sum_{j=2}^{m-1} \theta_j \varepsilon_{t+2-j} + \theta_1 \varepsilon_{t+1} \\ &\equiv \phi_2 \xi_{1t} + \xi_{3t} + \theta_1 \eta_t\end{aligned}$$

where $\xi_{3t} = \sum_{j=3}^m \phi_j y_{t+2-j} + \sum_{j=2}^{m-1} \theta_j \varepsilon_{t+2-j}$.

Next, considering $\xi_{3,t+1}$, we have

$$\begin{aligned}
 \xi_{3,t+1} &= \sum_{j=3}^m \phi_j y_{t+3-j} + \sum_{j=2}^{m-1} \theta_j \varepsilon_{t+3-j} \\
 &= \phi_3 y_t + \sum_{j=4}^m \phi_j y_{t+3-j} + \sum_{j=3}^{m-1} \theta_j \varepsilon_{t+3-j} + \theta_2 \varepsilon_{t+1} \\
 &= \phi_3 \xi_{1t} + \xi_{4t} + \theta_2 \eta_t,
 \end{aligned}$$

where $\xi_{4t} = \sum_{j=4}^m \phi_j y_{t+2-j} + \sum_{j=3}^{m-1} \theta_j \varepsilon_{t+2-j}$.

Repeating the procedure, we have

$$\xi_{mt} = \sum_{j=m}^m \phi_j y_{t+2-j} + \sum_{j=m-1}^{m-1} \theta_j \varepsilon_{t+2-j} = \phi_m y_{t-1} + \theta_{m-1} \varepsilon_t.$$

Finally,

$$\begin{aligned}\xi_{m,t+1} &= \phi_m y_t + \theta_{m-1} \varepsilon_{t+1} \\ &= \phi_m \xi_{1t} + \theta_{m-1} \eta_t.\end{aligned}$$

Define

$$\boldsymbol{\xi}_t = \begin{bmatrix} y_t \\ \phi_2 y_{t-1} + \cdots + \phi_m y_{t-m+1} + \theta_1 \varepsilon_t + \cdots + \theta_{m-1} \varepsilon_{t-m+2} \\ \vdots \\ \phi_m y_{t-1} + \theta_{m-1} \varepsilon_t \end{bmatrix}$$

and set

$$y_t = \begin{bmatrix} 1 & 0 & \cdots & 0 \end{bmatrix} \boldsymbol{\xi}_t$$

$$\boldsymbol{\xi}_t = \begin{bmatrix} \phi_1 & 1 & 0 & \cdots & 0 \\ \phi_2 & 0 & 1 & \cdots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \phi_{m-1} & 0 & 0 & \cdots & 1 \\ \phi_m & 0 & 0 & \cdots & 0 \end{bmatrix} \boldsymbol{\xi}_{t-1} + \begin{bmatrix} 1 \\ \theta_1 \\ \vdots \\ \theta_{m-2} \\ \theta_{m-1} \end{bmatrix} \varepsilon_t$$

(3) Stochastic volatility

Define a simple *time-vary variance process* as:

$$y_t = \sigma_t \varepsilon_t, \quad \varepsilon_t \sim \mathcal{N}(0, 1) \quad (6)$$

In a stochastic volatility model, $h_t = \log \sigma_t^2$, the logarithm of the variance, behaves exactly as a stochastic process in the mean, such as *random walks* or *autoregression*.

$$h_t = c + \lambda h_{t-1} + \zeta_t, \quad \zeta_t \sim \mathcal{N}(0, \sigma_\zeta^2) \quad (7)$$

Taking the square and the logarithm, Equation (6) can be reexpressed as:

$$\ln(y_t^2) = h_t + \ln(\varepsilon_t^2). \quad (8)$$

A *quasi-maximum likelihood* (QML) estimator has been proposed by Ruiz (1994) and Harvey et al. (1994) who approximate the measurement equation disturbance term by a *normal* random variable with the same mean

$$E[\ln(\varepsilon_t^2)] = -1.2704$$

and the variance

$$\text{Var}[\ln(\varepsilon_t^2)] = \pi^2/2.$$

Note: If x is a standard normal random variable, i.e. $x \sim N(0, 1)$, then the variable y with log-squared transformation of x is defined by

$$y = \ln(x^2).$$

So the density of y is

$$y = \frac{1}{\sqrt{2\pi}} \exp \left\{ \frac{x - \exp(x)}{2} \right\}$$

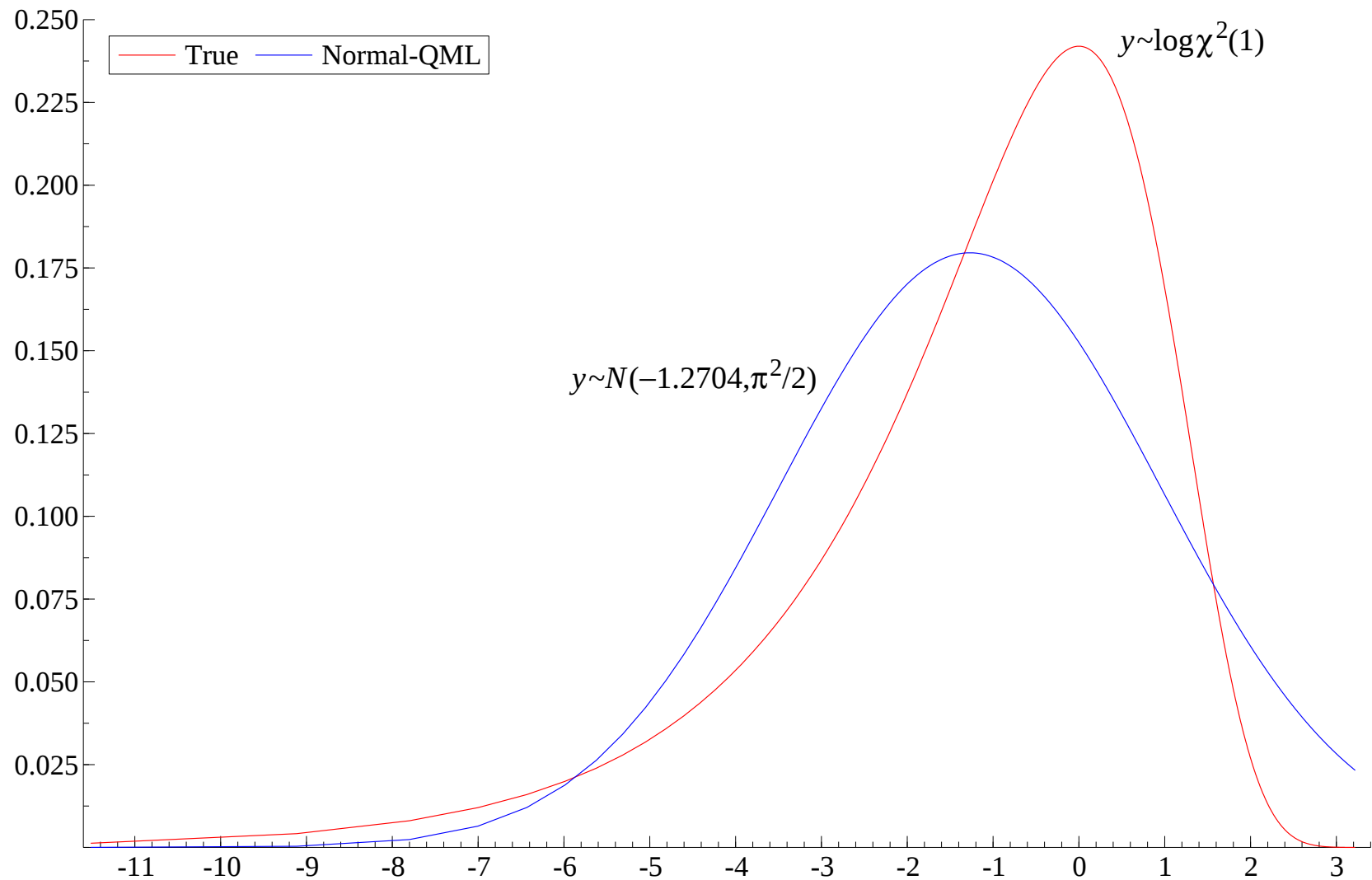


Figure 1 Log-chi-square density and its normal approximation.

(4) Time-varying coefficients

Specify a simple market model modified by using *time-varying coefficients*:

$$\begin{aligned}R_t &= \alpha_t + \beta_t R_{mt} + \varepsilon_t \\ \alpha_t &= \alpha_{t-1} + v_{1t} \\ \beta_t &= \beta_{t-1} + v_{2t}\end{aligned}\tag{9}$$

where R_t return on an individual security, R_{mt} is return on the market, and the coefficients follow a *random walk*.

The *observation equation* is:

$$R_t = [1 \quad R_{mt}] \begin{bmatrix} \alpha_t \\ \beta_t \end{bmatrix} + \varepsilon_t \quad (10)$$

The *state equation* is:

$$\begin{bmatrix} \alpha_t \\ \beta_t \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} \alpha_{t-1} \\ \beta_{t-1} \end{bmatrix} + \begin{bmatrix} v_{1t} \\ v_{2t} \end{bmatrix} \quad (11)$$

(5) Dynamic factor models

Suppose we have two stationary variables, y_{1t} , and y_{2t} with a *common component* c_t :

$$y_{1t} = \gamma_1 c_t + z_{1t}, \quad (12)$$

$$y_{2t} = \gamma_2 c_t + z_{2t}, \quad (13)$$

$$c_t = \phi_1 c_{t-1} + v_t, \quad (14)$$

$$z_{it} = \alpha_i z_{i,t-1} + e_{it}, \quad i = 1, 2, \quad (15)$$

where e_{1t} , e_{2t} , and v_t are independent of one another.

Stock and Watson (1991) employ a general version of the above dynamic factor model to extract an *coincident index* from four *coincident economic variables*.

Then, a state-space representation of the above model is given by

Measurement Equation

$$\begin{bmatrix} y_{1t} \\ y_{2t} \end{bmatrix} = \begin{bmatrix} \gamma_1 & 1 & 0 \\ \gamma_2 & 0 & 1 \end{bmatrix} \begin{bmatrix} c_t \\ z_{1t} \\ z_{2t} \end{bmatrix} \quad (16)$$

Transition Equation:

$$\begin{bmatrix} c_t \\ z_{1t} \\ z_{2t} \end{bmatrix} = \begin{bmatrix} \phi_1 & 0 & 0 \\ 0 & \alpha_1 & 0 \\ 0 & 0 & \alpha_2 \end{bmatrix} \begin{bmatrix} c_{t-1} \\ z_{1,t-1} \\ z_{2,t-1} \end{bmatrix} + \begin{bmatrix} v_t \\ e_{1t} \\ e_{2t} \end{bmatrix} \quad (17)$$

(6) A common stochastic trend model

Hai, Mark, and Wu (1996) consider the following *common stochastic trend* model for the *spot* and *forward* exchange

$$y_{1t} = z_t + x_{1t}, \quad (18)$$

$$y_{2t} = z_t + x_{2t}, \quad (19)$$

$$z_t = z_{t-1} + \varepsilon_t, \quad \varepsilon_t \sim i.i.d.\mathcal{N}(0, \sigma_\varepsilon^2), \quad (20)$$

$$x_{1t} = \mu_1 + \phi_{11}x_{1,t-1} + \phi_{12}x_{2,t-1} + e_{1t} + \theta_{11}e_{1,t-1} + \theta_{12}e_{2,t-1}, \quad (21)$$

$$x_{2t} = \mu_2 + \phi_{21}x_{1,t-1} + \phi_{22}x_{2,t-1} + e_{2t} + \theta_{21}e_{1,t-1} + \theta_{22}e_{2,t-1}, \quad (22)$$

where

$$\begin{pmatrix} e_{1t} \\ e_{2t} \end{pmatrix} \sim i.i.d.\mathcal{N} \left(\begin{bmatrix} 0 \\ 0 \end{bmatrix}, \begin{bmatrix} \sigma_1^2 & \sigma_{12} \\ \sigma_{21} & \sigma_2^2 \end{bmatrix} \right),$$

z_t is a common stochastic trend and the stationary component x_{1t} and x_{2t} are assumed to be generated by a vector ARMA(1,1) process.

A state-space representation of the above model is given by

Measurement Equation

$$\begin{bmatrix} y_{1t} \\ y_{2t} \end{bmatrix} = \begin{bmatrix} 1 & 1 & 0 & 0 & 0 \\ 1 & 0 & 1 & 0 & 0 \end{bmatrix} \begin{bmatrix} z_t \\ x_{1t} \\ x_{2t} \\ e_{1t} \\ e_{2t} \end{bmatrix} \quad (23)$$

Transition Equation:

$$\begin{bmatrix} z_t \\ x_{1t} \\ x_{2t} \\ e_{1t} \\ e_{2t} \end{bmatrix} = \begin{bmatrix} 0 \\ \mu_1 \\ \mu_2 \\ 0 \\ 0 \end{bmatrix} + \begin{bmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & \phi_{11} & \phi_{12} & \theta_{11} & \theta_{12} \\ 0 & \phi_{21} & \phi_{22} & \theta_{21} & \theta_{22} \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \end{bmatrix} \begin{bmatrix} z_{t-1} \\ x_{1,t-1} \\ x_{2,t-1} \\ e_{1,t-1} \\ e_{2,t-1} \end{bmatrix} + \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} \varepsilon_t \\ e_{1t} \\ e_{2t} \end{bmatrix}. \quad (24)$$

(7) GDP decomposition

A classic trend/cycle decomposition model is given as follows:

$$y_t = \tau_t + c_t \quad (25)$$

$$\tau_t = \tau_{t-1} + g_{t-1} + u_t \quad (26)$$

$$g_t = gc + \lambda g_{t-1} + w_t, \quad (27)$$

$$c_t = \phi_1 c_{t-1} + \phi_2 c_{t-2} + v_t \quad (28)$$

where y_t is log GDP; τ_t is its *trend component* follows a *random walk* with a *stochastic drift* or growth rate which is an autoregressive process; c_t is the *cycle component*.

When $\lambda = \sigma_w = 0$, this becomes the Watson (1986) model.

When $gc = 0$ and $\lambda = 1$, this becomes the Clark (1987) model.

Write equations (25) – (28) in the state space form, the observation equation is:

$$y_t = \begin{bmatrix} 1 & 1 & 0 & 0 \end{bmatrix} \begin{bmatrix} \tau_t \\ c_t \\ c_{t-1} \\ g_t \end{bmatrix} \quad (29)$$

The state equation is:

$$\begin{bmatrix} \tau_t \\ c_t \\ c_{t-1} \\ g_t \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 & 1 \\ 0 & \varphi_1 & \varphi_2 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & \lambda \end{bmatrix} \begin{bmatrix} \tau_{t-1} \\ c_{t-1} \\ c_{t-2} \\ g_{t-1} \end{bmatrix} + \begin{bmatrix} 0 \\ 0 \\ 0 \\ gc \end{bmatrix} + \begin{bmatrix} v_t \\ u_t \\ 0 \\ w_t \end{bmatrix} \quad (30)$$

(8) BSM model in STAMP

The basic structural model (BSM) is formulated in terms of trend, seasonal, and irregular components. It can be represented by

$$y_t = \tau_t + \gamma_t + \varepsilon_t, \quad \varepsilon_t \sim iidN(0, \sigma_\varepsilon^2), \quad (31)$$

where τ_t is a stochastic trend, γ_t is a stochastic seasonal component, and ε_t is an irregular component; see Harvey (1989).

The trend is specified in the following way:

$$\tau_t = \tau_{t-1} + \beta_{t-1} + \eta_t, \quad \eta_t \sim iidN(0, \sigma_\eta^2) \quad (32)$$

$$\beta_t = \beta_{t-1} + \zeta_t, \quad \zeta_t \sim iidN(0, \sigma_\zeta^2), \quad (33)$$

where β_t is the slope, η_t and ζ_t are assumed to be mutually independent.

The seasonal component γ_t is often specified in two ways such as dummy-variable seasonality and trigonometric seasonality, see Harvey et al. (1998) and Commandeur and Koopman(2007).

A time-varying dummy seasonal specification is given by

$$\sum_{i=0}^{s-1} \gamma_{t-i} = e_t, \quad (34)$$

where s is the seasonal length, e_t is an i.i.d. normal variable with mean zero and variance σ_e^2 , i.e., $e_t \sim iidN(0, \sigma_e^2)$.

- In the limiting case $\sigma_e^2 = 0$, the seasonal effects are fixed over time.
- Commonly, $s = 4$ for quarterly data and $s = 12$ for monthly data.

(10) Dynamic Nelson and Siegel model

Diebold and Li (2006) and Diebold, Rudebusch and Arouba (2006) extend the model of Nelson and Siegel (1987) to a state space framework and to estimate the factors of term structure of interest rate. We often call this the dynamic Nelson and Siegel model.

The measurement equation of the state-space form is

$$\begin{bmatrix} y_t(\tau_1) \\ y_t(\tau_2) \\ \vdots \\ y_t(\tau_n) \end{bmatrix} = \begin{bmatrix} 1 & \frac{1-e^{-\lambda\tau_1}}{\lambda\tau_1} & \frac{1-e^{-\lambda\tau_1}}{\lambda\tau_1} - e^{-\lambda\tau_1} \\ 1 & \frac{1-e^{-\lambda\tau_2}}{\lambda\tau_2} & \frac{1-e^{-\lambda\tau_2}}{\lambda\tau_2} - e^{-\lambda\tau_2} \\ \vdots & \vdots & \vdots \\ 1 & \frac{1-e^{-\lambda\tau_n}}{\lambda\tau_n} & \frac{1-e^{-\lambda\tau_n}}{\lambda\tau_n} - e^{-\lambda\tau_n} \end{bmatrix} \begin{bmatrix} \beta_{1,t} \\ \beta_{2,t} \\ \beta_{3,t} \end{bmatrix} + \begin{bmatrix} \varepsilon_t(\tau_1) \\ \varepsilon_t(\tau_2) \\ \vdots \\ \varepsilon_t(\tau_n) \end{bmatrix}$$

where $\beta_{1,t}$, $\beta_{2,t}$ and $\beta_{3,t}$ are the level, slope and curvature factors.

The transition equation specifies the dynamics of the state variables

$$\begin{bmatrix} \beta_{1,t} \\ \beta_{2,t} \\ \beta_{3,t} \end{bmatrix} = \begin{bmatrix} \phi_1 & 0 & 0 \\ 0 & \phi_2 & 0 \\ 0 & 0 & \phi_3 \end{bmatrix} \begin{bmatrix} \beta_{1,t-1} \\ \beta_{2,t-1} \\ \beta_{3,t-1} \end{bmatrix} + \begin{bmatrix} \eta_{1,t} \\ \eta_{2,t} \\ \eta_{3,t} \end{bmatrix},$$

$$\begin{bmatrix} \eta_{1,t} \\ \eta_{2,t} \\ \eta_{3,t} \end{bmatrix} \sim N \left(0, \begin{bmatrix} \sigma_1^2 & 0 & 0 \\ 0 & \sigma_2^2 & 0 \\ 0 & 0 & \sigma_3^2 \end{bmatrix} \right),$$

where we have restricted the transition matrix to be diagonal to reflect the theoretical and empirical likelihood that the three factors exhibit little correlation.

2 Kalman Filter and Estimation

Reconsider the model:

$$\begin{aligned}\mathbf{y}_t &= \mathbf{H} \mathbf{x}_t + \mathbf{d} + \mathbf{u}_t \\ \mathbf{x}_t &= \mathbf{F} \mathbf{x}_{t-1} + \mathbf{c} + \mathbf{v}_t\end{aligned}$$

$n \times 1 \quad n \times r \quad r \times 1 \quad n \times 1 \quad n \times 1$
 $r \times 1 \quad r \times r \quad r \times 1 \quad r \times 1 \quad r \times 1$

$$\text{Cov}(\mathbf{u}_t \mathbf{u}_t) = \mathbf{R}, \quad \text{Cov}(\mathbf{v}_t \mathbf{v}_t) = \mathbf{Q}, \quad \text{Cov}(\mathbf{u}_t \mathbf{v}_t) = \mathbf{0}.$$

$n \times n \quad r \times r$

Let the observed data obtained through date $t - 1$ be summarized by

$$\mathcal{F}_{t-1} \equiv (\mathbf{y}'_{t-1}, \mathbf{y}'_{t-2}, \dots, \mathbf{y}'_1)'$$

It is natural to obtain the one-step-ahead prediction density of $\mathbf{y}_t$

conditional on $\mathcal{F}_{t-1}$

$$p(\mathbf{y}_t \mid \mathcal{F}_{t-1}) = \int p(\mathbf{y}_t \mid \mathbf{x}_t) p(\mathbf{x}_t \mid \mathcal{F}_{t-1}) d\mathbf{x}_t$$

where $p(\mathbf{y}_t \mid \mathbf{x}_t)$ denotes the observation density implied by the observation equation.

The prediction density of $\mathbf{x}_t$ given $\mathcal{F}_{t-1}$ is given by

$$p(\mathbf{x}_t \mid \mathcal{F}_{t-1}) = \int p(\mathbf{x}_t \mid \mathbf{x}_{t-1}) p(\mathbf{x}_{t-1} \mid \mathcal{F}_{t-1}) d\mathbf{x}_{t-1}$$

where $p(\mathbf{x}_t \mid \mathbf{x}_{t-1})$ denotes the prediction density implied by the transition equation.

The filtered density of $\mathbf{x}_t$ given by $\mathcal{F}_t$ is given by

$$p(\mathbf{x}_t | \mathcal{F}_t) = \frac{p(\mathbf{y}_t | \mathbf{x}_t)p(\mathbf{x}_t | \mathcal{F}_{t-1})}{p(\mathbf{y}_t | \mathcal{F}_{t-1})} \propto p(\mathbf{y}_t | \mathbf{x}_t)p(\mathbf{x}_t | \mathcal{F}_{t-1})$$

by Bayes rule.

Based on the previous prediction and filtered densities, we can implement a recursion from $t - 1$ to t . If the input density $p(\mathbf{x}_{t-1} | \mathcal{F}_{t-1})$ is Gaussian, then the output density $p(\mathbf{x}_t | \mathcal{F}_t)$ is also Gaussian.

Note that the Gaussian density can be summarized by the corresponding mean and covariance matrix.

This leads to the Kalman filter for the mean and mean-square-error of the state and observation.

The Kalman filter can be better demonstrated in three steps, though at least the first two steps can be easily combined. The three steps are **prediction, updating, and smoothing**.

(1) Prediction

It is noted that, based on information available at $t - 1$, the state vector $\xi_t | \mathcal{F}_{t-1} \sim N(\xi_{t|t-1}, P_{t|t-1})$, where:

$$\xi_{t|t-1} := E(\xi_t | \mathcal{F}_{t-1}) = F \xi_{t-1|t-1} + c, \quad (35)$$

$$P_{t|t-1} := \text{Var}(\xi_t | \mathcal{F}_{t-1}) = F P_{t-1|t-1} F' + Q, \quad (36)$$

where the **input means** $\xi_{t-1|t-1} = E(\xi_{t-1} | \mathcal{F}_{t-1})$ and **variances** $P_{t-1|t-1} = \text{Var}(\xi_{t-1} | \mathcal{F}_{t-1})$ is assumed to be known at time t . These two equations are known as the prediction equations.

The prediction of $\mathbf{y}_t$ conditional on $\mathcal{F}_{t-1}$ can be inferred immediately from (5):

$$\mathbf{y}_{t|t-1} = \mathbb{E}(\mathbf{y}_t | \mathcal{F}_{t-1}) = \mathbf{H}\boldsymbol{\xi}_{t|t-1} + \mathbf{d}. \quad (37)$$

Then the *prediction error* can be written as

$$\boldsymbol{\eta}_t = \mathbf{y}_t - \mathbf{y}_{t|t-1} = \mathbf{H}(\boldsymbol{\xi}_t - \boldsymbol{\xi}_{t|t-1}) + \mathbf{u}_t. \quad (38)$$

Since $\boldsymbol{\xi}_{t|t-1}$ is a function of $\mathcal{F}_{t-1}$, the term $\mathbf{u}_t$ is *independent* of both $\boldsymbol{\xi}_t$ and $\boldsymbol{\xi}_{t|t-1}$.

Thus the *conditional variance* of (38) is

$$\begin{aligned} \boldsymbol{\Sigma}_t &= \text{Var}(\mathbf{y}_t - \mathbf{y}_{t|t-1} | \mathcal{F}_{t-1}) = \mathbf{H}\text{Var}(\boldsymbol{\xi}_t | \mathcal{F}_{t-1})\mathbf{H}' + \mathbf{R} \\ &= \mathbf{H}\mathbf{P}_{t|t-1}\mathbf{H}' + \mathbf{R}. \end{aligned} \quad (39)$$

(2) Updating

First note that the conditional covariance between (38) and the error in forecasting the state vector is

$$\begin{aligned} E\{[\mathbf{y}_t - \mathbf{y}_{t|t-1}][\boldsymbol{\xi}_t - \boldsymbol{\xi}_{t|t-1}]' | \mathcal{F}_{t-1}\} &= \mathbf{H} \cdot E\{[\boldsymbol{\xi}_t - \boldsymbol{\xi}_{t|t-1}][\boldsymbol{\xi}_t - \boldsymbol{\xi}_{t|t-1}]' | \mathcal{F}_{t-1}\} \\ &= \mathbf{H} \mathbf{P}_{t|t-1} \end{aligned}$$

Thus the distribution of the vector $(\mathbf{y}_t', \boldsymbol{\xi}_t')'$ conditional on $\mathcal{F}_{t-1}$ is

$$\begin{bmatrix} \mathbf{y}_t | \mathcal{F}_{t-1} \\ \boldsymbol{\xi}_t | \mathcal{F}_{t-1} \end{bmatrix} \sim N \left(\begin{bmatrix} \mathbf{H} \boldsymbol{\xi}_{t|t-1} + \mathbf{d} \\ \boldsymbol{\xi}_{t|t-1} \end{bmatrix}, \begin{bmatrix} \mathbf{H} \mathbf{P}_{t|t-1} \mathbf{H}' + \mathbf{R} & \mathbf{H} \mathbf{P}_{t|t-1} \\ \mathbf{P}_{t|t-1} \mathbf{H}' & \mathbf{P}_{t|t-1} \end{bmatrix} \right)$$

► *Using this joint conditional distribution, our next purpose is to obtain the distribution of $\boldsymbol{\xi}_t | \mathcal{F}_t$, or $\boldsymbol{\xi}_t | \mathbf{y}_t, \mathcal{F}_{t-1}$. So we apply the following theorem:*

Theorem 1 (Gaussian Conditioning Theorem). *Let z_1 and z_2 denote $(n_1 \times 1)$ and $(n_2 \times 1)$ vectors respectively that have a joint normal distribution:*

$$\begin{bmatrix} z_1 \\ z_2 \end{bmatrix} \sim N \left(\begin{bmatrix} \mu_1 \\ \mu_2 \end{bmatrix}, \begin{bmatrix} \Sigma_{11} & \Sigma_{12} \\ \Sigma_{21} & \Sigma_{22} \end{bmatrix} \right).$$

Then the distribution of z_2 conditional on z_1 is $N(\mu_{2|1}, \Sigma_{22|1})$ where

$$\mu_{2|1} = E(z_2|z_1) = \mu_2 + \Sigma_{21}\Sigma_{11}^{-1}(z_1 - \mu_1) \quad (40)$$

$$\Sigma_{22|1} = E[(z_2 - \mu_{2|1})(z_2 - \mu_{2|1})'|z_1] = \Sigma_{22} - \Sigma_{21}\Sigma_{11}^{-1}\Sigma_{12}. \quad (41)$$

Remark: Thus the *optimal forecast* of z_2 conditional on having observed z_1 is $\mu_{2|1}$, while the *mean squared error* of this forecast is $\Sigma_{22|1}$.

Proof. We have

$$\det(\Sigma) = \det(\Sigma_{11}) \det(\Sigma_{22} - \Sigma_{21}\Sigma_{11}^{-1}\Sigma_{12}) = \det(\Sigma_{11}) \det(\Sigma_{22|1}).$$

$$\text{So } f_{2|1}(\mathbf{z}_2|\mathbf{z}_1) := \frac{f(\mathbf{z})}{f_1(\mathbf{z}_1)} =$$

$$\begin{aligned} &= \frac{(2\pi)^{-n/2} \det(\Sigma)^{-1/2} \exp \left\{ -\frac{1}{2}(\mathbf{z} - \boldsymbol{\mu})' \Sigma^{-1/2}(\mathbf{z} - \boldsymbol{\mu}) \right\}}{(2\pi)^{-n_1/2} \det(\Sigma_{11})^{-1/2} \exp \left\{ -\frac{1}{2}(\mathbf{z}_1 - \boldsymbol{\mu}_1)' \Sigma_{11}^{-1/2}(\mathbf{z}_1 - \boldsymbol{\mu}_1) \right\}} \\ &= (2\pi)^{-n_2/2} \det(\Sigma_{22|1})^{-1/2} \\ &\quad \times \exp \left\{ -\frac{1}{2} \left[(\mathbf{z} - \boldsymbol{\mu})' \Sigma^{-1}(\mathbf{z} - \boldsymbol{\mu}) - (\mathbf{z}_1 - \boldsymbol{\mu}_1)' \Sigma_{11}^{-1}(\mathbf{z}_1 - \boldsymbol{\mu}_1) \right] \right\}. \end{aligned}$$

where the term $(\mathbf{z} - \boldsymbol{\mu})' \Sigma^{-1}(\mathbf{z} - \boldsymbol{\mu}) - (\mathbf{z}_1 - \boldsymbol{\mu}_1)' \Sigma_{11}^{-1}(\mathbf{z}_1 - \boldsymbol{\mu}_1)$

$$= [\mathbf{z}_2 - \boldsymbol{\mu}_2 - \Sigma_{21}\Sigma_{11}^{-1}(\mathbf{z}_1 - \boldsymbol{\mu}_1)]' \Sigma_{22|1}^{-1} [\mathbf{z}_2 - \boldsymbol{\mu}_2 - \Sigma_{21}\Sigma_{11}^{-1}(\mathbf{z}_1 - \boldsymbol{\mu}_1)].$$

We have

$$\Sigma^{-1} = \begin{bmatrix} B_{11} & B_{12} \\ B_{21} & B_{22} \end{bmatrix}$$

where

$$\begin{aligned} B_{11} &:= \Sigma_{11}^{-1} + \Sigma_{11}^{-1} \Sigma_{12} (\Sigma_{22} - \Sigma_{21} \Sigma_{11}^{-1} \Sigma_{12})^{-1} \Sigma_{21} \Sigma_{11}^{-1} \\ &= \Sigma_{11}^{-1} + \Sigma_{11}^{-1} \Sigma_{12} \Sigma_{22|1}^{-1} \Sigma_{21} \Sigma_{11}^{-1}, \\ B_{12} &:= - \Sigma_{11}^{-1} \Sigma_{12} (\Sigma_{22} - \Sigma_{21} \Sigma_{11}^{-1} \Sigma_{12})^{-1} \\ &= - \Sigma_{11}^{-1} \Sigma_{12} \Sigma_{22|1}^{-1} \\ B_{21} &:= - (\Sigma_{22} - \Sigma_{21} \Sigma_{11}^{-1} \Sigma_{12})^{-1} \Sigma_{21} \Sigma_{11}^{-1} \\ &= - \Sigma_{22|1}^{-1} \Sigma_{21} \Sigma_{11}^{-1} \\ B_{22} &:= \Sigma_{22|1}^{-1}. \end{aligned}$$

So

$$\begin{aligned} & (\mathbf{z} - \boldsymbol{\mu})' \boldsymbol{\Sigma}^{-1} (\mathbf{z} - \boldsymbol{\mu}) - (\mathbf{z}_1 - \boldsymbol{\mu}_1)' \boldsymbol{\Sigma}_{11}^{-1} (\mathbf{z}_1 - \boldsymbol{\mu}_1) \\ &= (\mathbf{z}_1 - \boldsymbol{\mu}_1)' \boldsymbol{\Sigma}_{11}^{-1} \boldsymbol{\Sigma}_{12} \boldsymbol{\Sigma}_{22|1}^{-1} \boldsymbol{\Sigma}_{21} \boldsymbol{\Sigma}_{11}^{-1} (\mathbf{z}_1 - \boldsymbol{\mu}_1) - 2(\mathbf{z}_1 - \boldsymbol{\mu}_1)' \boldsymbol{\Sigma}_{11}^{-1} \boldsymbol{\Sigma}_{12} \boldsymbol{\Sigma}_{22|1}^{-1} (\mathbf{z}_2 - \boldsymbol{\mu}_2) \\ & \quad + (\mathbf{z}_2 - \boldsymbol{\mu}_2)' \boldsymbol{\Sigma}_{22|1}^{-1} (\mathbf{z}_2 - \boldsymbol{\mu}_2) \\ &= (\mathbf{z}_2 - \boldsymbol{\mu}_2)' \boldsymbol{\Sigma}_{22|1}^{-1} (\mathbf{z}_2 - \boldsymbol{\mu}_2) - 2(\mathbf{z}_1 - \boldsymbol{\mu}_1)' \boldsymbol{\Sigma}_{22|1}^{-1} \boldsymbol{\Sigma}_{21} \boldsymbol{\Sigma}_{11}^{-1} (\mathbf{z}_2 - \boldsymbol{\mu}_2) \\ & \quad + \left[\boldsymbol{\Sigma}_{21} \boldsymbol{\Sigma}_{11}^{-1} (\mathbf{z}_1 - \boldsymbol{\mu}_1) \right]' \boldsymbol{\Sigma}_{22|1}^{-1} \boldsymbol{\Sigma}_{21} \boldsymbol{\Sigma}_{11}^{-1} (\mathbf{z}_1 - \boldsymbol{\mu}_1) \\ &= \left[(\mathbf{z}_2 - \boldsymbol{\mu}_2) - \boldsymbol{\Sigma}_{21} \boldsymbol{\Sigma}_{11}^{-1} (\mathbf{z}_1 - \boldsymbol{\mu}_1) \right]' \boldsymbol{\Sigma}_{22|1}^{-1} \left[(\mathbf{z}_2 - \boldsymbol{\mu}_2) - \boldsymbol{\Sigma}_{21} \boldsymbol{\Sigma}_{11}^{-1} (\mathbf{z}_1 - \boldsymbol{\mu}_1) \right]. \end{aligned}$$

□

Results from the theorem: It then follows from (40) and (41) that $\xi_t|\mathcal{F}_t = \xi_t|\mathbf{y}_t, \mathcal{F}_{t-1}$ is distributed $N(\xi_{t|t}, \mathbf{P}_{t|t})$, where

$$\xi_{t|t} = \mathbb{E}(\xi_t|\mathcal{F}_t) = \xi_{t|t-1} + \mathbf{K}_t\eta_t, \quad (42)$$

$$\mathbf{P}_{t|t} = \text{Var}(\xi_t|\mathcal{F}_t) = \mathbf{P}_{t|t-1} - \mathbf{K}_t\mathbf{H}\mathbf{P}_{t|t-1}, \quad (43)$$

where the quantity

$$\mathbf{K}_t = \mathbf{P}_{t|t-1}\mathbf{H}'\Sigma_t^{-1} = \mathbf{P}_{t|t-1}\mathbf{H}'(\mathbf{H}\mathbf{P}_{t|t-1}\mathbf{H}' + \mathbf{R})^{-1},$$

which is also called the *Kalman gain*.

Kalman Filtering

Taken together (35), (35), (42) and (43) make up the Kalman filter.

If desired they can be written as a single set of recursions going directly from $\xi_{t-1|t-1}$ to $\xi_{t|t}$ or, alternatively, from $\xi_{t|t-1}$ to $\xi_{t+1|t}$. We might refer to these as, respectively, the *contemporaneous* and *predictive filter*.

In the latter case

$$\xi_{t+1|t} = F\xi_{t|t} + c = F\xi_{t|t-1} + c + FK_t\eta_t \quad (44)$$

The recursion for the covariance matrix (called as a *Riccati equation*)

$$P_{t+1|t} = F(P_{t|t-1} - K_tHP_{t|t-1})F' + Q. \quad (45)$$

Starting values

The starting values for the Kalman filter may be specified in terms of ξ_0 and P_0 or $\xi_{1|0}$ and $P_{1|0}$.

- If ξ_t is stationary, we use the *unconditional mean* and *unconditional covariance matrix* of ξ_t

$$\xi_0 = E[\xi_t] = (\mathbf{I}_r - \mathbf{F})^{-1}\mathbf{c}$$
$$\text{vec}(\mathbf{P}_0) = (\mathbf{I}_{r^2} - \mathbf{F} \otimes \mathbf{F})^{-1}\text{vec}(\mathbf{Q}),$$

where $\text{vec}(\mathbf{P}_0)$ is an $r^2 \times 1$ column vector. It can be easily *reshaped to form* the $r \times r$ matrix $\mathbf{P}_0$.

- If ξ_t is non-stationary, a diffuse prior is often used by setting $\mathbf{P}_0 = \kappa \mathbf{I}$, and letting the scalar κ go to infinity.

(3) Smoothing (for SSM with fixed parameters)

Given initial conditions $\xi_{T|T}$ and $P_{T|T}$ from the Kalman filter.

We can *recursively* obtain $\xi_{t|T}$ and $P_{t|T}$, for $t = T - 1, T - 2, \dots, 1$,

$$\xi_{t|T} := E(\xi_t | \mathcal{F}_T) = \xi_{t|t} + \Upsilon_t(\xi_{t+1|T} - \xi_{t+1|t}) \quad (46)$$

$$P_{t|T} := \text{Var}(\xi_t | \mathcal{F}_T) = P_{t|t} + \Upsilon_t(P_{t+1|T} - P_{t+1|t})\Upsilon_t' \quad (47)$$

where $\Upsilon_t = P_{t|t}F'P_{t+1|t}^{-1}$.

Especially, at time $t = 0$, we can get the estimates of initial state

$$\xi_{0|T} \quad \text{and} \quad P_{0|T},$$

which can be used to set as a_0 and Σ_0 .

Derivation of smoothing algorithm

(a) **Computing** $p(\xi_t, \xi_{t+1} | \mathcal{F}_t)$, which is Gaussian and it suffices to compute its mean vector and its covariance matrix

- The mean vector is $\begin{bmatrix} \xi_{t|t} \\ \xi_{t+1|t} \end{bmatrix}$, where $\xi_{t+1|t} = F\xi_{t|t} + c$.
- We will use this when calculating the cross covariance:

$$\begin{aligned} \text{Cov}(\xi_t, \xi_{t+1} | \mathcal{F}_t) &= E[(\xi_t - \xi_{t|t})(\xi_{t+1} - \xi_{t+1|t})' | \mathcal{F}_t] \\ &= E[(\xi_t - \xi_{t|t})(F\xi_t + c + v_{t+1} - F\xi_{t|t} - c)' | \mathcal{F}_t] \\ &= P_{t|t}F'. \end{aligned}$$

- We have $\begin{bmatrix} \xi_t \\ \xi_{t+1} \end{bmatrix} \Big| \mathcal{F}_t \sim N \left(\begin{bmatrix} \xi_{t|t} \\ \xi_{t+1|t} \end{bmatrix}, \begin{bmatrix} P_{t|t} & P_{t|t}F' \\ FP_{t|t} & P_{t+1|t} \end{bmatrix} \right).$

(b) **Computing** $p(\xi_t|\xi_{t+1}, \mathcal{F}_T) = p(\xi_t|\xi_{t+1}, \mathcal{F}_t)$ by Gaussian Conditioning Theorem.

– The corresponding mean and covariance matrix are

$$E[\xi_t|\xi_{t+1}, \mathcal{F}_T] = E[\xi_t|\xi_{t+1}, \mathcal{F}_t] = \xi_{t|t} + \mathbf{L}_t(\xi_{t+1} - \xi_{t+1|t})$$

$$\text{Cov}(\xi_t|\xi_{t+1}, \mathcal{F}_T) = \text{Cov}(\xi_t|\xi_{t+1}, \mathcal{F}_t) = \mathbf{P}_{t|t} - \mathbf{L}_t \mathbf{P}_{t+1|t} \mathbf{L}_t',$$

where $\mathbf{L}_t = \mathbf{P}_{t|t} \mathbf{F} \mathbf{P}_{t+1|t}^{-1}$.

– Hence we have

$$\xi_t|\xi_{t+1}, \mathcal{F}_T \sim N(\xi_{t|t} + \mathbf{L}_t(\xi_{t+1} - \xi_{t+1|t}), \mathbf{P}_{t|t} - \mathbf{L}_t \mathbf{P}_{t+1|t} \mathbf{L}_t').$$

(c) **Computing** $p(\xi_t|\mathcal{F}_T)$

– Two properties are used:

(c.1.) Law of Iterated Expectation (Tower Property):

$$E[Z|X] = E[E[Z|Y, X]|X].$$

(c.2.) Law of Total Conditional Variance:

$$\text{Cov}[Z|X] = \text{Cov}[E[Z|Y, X]|X] + E[\text{Cov}[Z|Y, X]|X].$$

- Let $Z = \xi_t$, $Y = \xi_{t+1}$, $X = \mathcal{F}_T$, and apply the above two properties.
► First, we compute the smoothed mean

$$\begin{aligned}\xi_{t|T} &= E[\xi_t|\mathcal{F}_T] = E[E[\xi_t|\xi_{t+1}, \mathcal{F}_T]|\mathcal{F}_T] \\ &= E[\xi_{t|t} + \mathbf{L}_t(\xi_{t+1} - \xi_{t+1|t})|\mathcal{F}_T] \\ &= \underbrace{\xi_{t|t}}_{\text{filtered estimate}} + \underbrace{\mathbf{L}_t(\xi_{t+1|T} - \xi_{t+1|t})}_{\text{correction term}}.\end{aligned}$$

- Second, all that remains is to apply the law of total conditional variance to find $P_{t|T}$:

$$\begin{aligned}
 P_{t|T} &= \text{Cov}[\boldsymbol{\xi}_t | \mathcal{F}_T] \\
 &= \text{Cov}[E[\boldsymbol{\xi}_t | \boldsymbol{\xi}_{t+1}, \mathcal{F}_T] | \mathcal{F}_T] + E[\text{Cov}[\boldsymbol{\xi}_t | \boldsymbol{\xi}_{t+1}, \mathcal{F}_T] | \mathcal{F}_T] \\
 &= \text{Cov}[\boldsymbol{\xi}_{t|t} + \mathbf{L}_t(\boldsymbol{\xi}_{t+1} - \boldsymbol{\xi}_{t+1|t}) | \mathcal{F}_T] + E[\mathbf{P}_{t|t} - \mathbf{L}_t \mathbf{P}_{t+1|t} \mathbf{L}_t' | \mathcal{F}_T] \\
 &= \mathbf{L}_t \mathbf{P}_{t+1|T} \mathbf{L}_t' + (\mathbf{P}_{t|t} - \mathbf{L}_t \mathbf{P}_{t+1|t} \mathbf{L}_t') \\
 &= \underbrace{\mathbf{P}_{t|t}}_{\text{filtered estimate}} + \underbrace{\mathbf{L}_t(\mathbf{P}_{t+1|T} - \mathbf{P}_{t+1|t})\mathbf{L}_t'}_{\text{correction term}}.
 \end{aligned}$$

(4) The Lag-One Covariance Smoother *

For the state-space model, with $\mathbf{K}_t$, Υ_t ($t = 1, \dots, T$) obtained from before, and with initial condition

$$\begin{aligned} P_{T,T-1|T} &= \text{Cov}(\xi_T, \xi_{T-1} \mid \mathcal{F}_T) \\ &= (\mathbf{I} - \mathbf{K}_T \mathbf{H}) \mathbf{F} P_{T-1|T-1}, \end{aligned}$$

we have that for $t = T - 1, T - 2, \dots, 1$,

$$\begin{aligned} P_{t,t-1|T} &= \text{Cov}(\xi_t, \xi_{t-1} \mid \mathcal{F}_T) \\ &= P_{t|t} \Upsilon'_{t-1} + \Upsilon_t (P_{t+1,t|T} - \mathbf{F} P_{t|t}) \Upsilon'_{t-1}. \end{aligned}$$

The proof is given in Chapter 6 in Shumway and Stoffer (2006).

1.3 Prediction

As regards multi-step prediction, taking expectations, conditional on the information at time T , of the transition equation at time $T + h$ yields the recursion

$$\boldsymbol{\xi}_{T+h|T} = \mathbf{F}\boldsymbol{\xi}_{T+h-1|T} + \mathbf{c}, \quad h = 1, 2, 3, \dots \quad (48)$$

with initial value $\boldsymbol{\xi}_{T|T}$. Similarly,

$$\mathbf{P}_{T+h|T} = \mathbf{F}\mathbf{P}_{T+h-1|T}\mathbf{F}' + \mathbf{Q}, \quad h = 1, 2, 3, \dots \quad (49)$$

Thus $\boldsymbol{\xi}_{T+h|T}$ and $\mathbf{P}_{T+h|T}$ are evaluated by repeatedly applying the Kalman filter prediction equations.

The MMSE of $\mathbf{y}_{T+h}$ can be obtained directly from $\boldsymbol{\xi}_{T+h|T}$. Taking conditional expectations in the measurement equation for $\mathbf{y}_{T+h}$ gives

$$\hat{\mathbf{y}}_{T+h|T} = E(\mathbf{y}_{T+h}|\mathcal{F}_T) = \mathbf{H}\boldsymbol{\xi}_{T+h|T} + \mathbf{d}, \quad h = 1, 2, 3, \dots \quad (50)$$

with MSE matrix

$$MSE(\hat{\mathbf{y}}_{T+h|T}) = \mathbf{H}\mathbf{P}_{T+h|T}\mathbf{H}' + \mathbf{R}, \quad h = 1, 2, 3, \dots \quad (51)$$

When the normality assumption is relaxed, $\hat{\mathbf{y}}_{T+h|T}$ and $MSE(\hat{\mathbf{y}}_{T+h|T})$ are still minimum mean square linear estimators.

2.4 Maximum likelihood estimation

The joint density function for the T sets of observations, $\mathbf{y}_1, \dots, \mathbf{y}_T$, is

$$p(\mathbf{y}_1, \dots, \mathbf{y}_T; \theta) = \prod_{t=1}^T p(\mathbf{y}_t | \mathcal{F}_{t-1}) \quad (52)$$

where $p(\mathbf{y}_t | \mathcal{F}_{t-1})$ denotes the distribution of $\mathbf{y}_t$ conditional on the information set at time $t - 1$, that is, $\mathcal{F}_{t-1} = \{\mathbf{y}_{t-1}, \dots, \mathbf{y}_1\}$.

In the Gaussian state space model, the conditional distribution of $\mathbf{y}_t$ is normal with mean $\hat{\mathbf{y}}_{t|t-1}$ and covariance matrix Σ_t . Hence the prediction errors or innovations,

$$\boldsymbol{\eta}_t = \mathbf{y}_t - \hat{\mathbf{y}}_{t|t-1} \sim N(0, \Sigma_t), \quad t = 1, \dots, T, \quad (53)$$

is serially independent with mean zero and covariance Σ_t .

Based on the predictive distribution of $\mathbf{y}_t$, the conditional density at time t is:

$$p(\mathbf{y}_t|\mathcal{F}_{t-1}) = (2\pi)^{-n/2}|\boldsymbol{\Sigma}_t|^{-1/2} \exp\left(-\frac{\boldsymbol{\eta}_t'\boldsymbol{\Sigma}_t^{-1}\boldsymbol{\eta}_t}{2}\right) \quad (54)$$

where $\boldsymbol{\eta}_t = \mathbf{y}_t - \mathbf{y}_{t|t-1}$, $\mathcal{F}_{t-1}$ is the information set at time $t - 1$.

The parameters can be estimated by maximizing the log likelihood of the density function:

$$\mathcal{L}(\theta) = -\frac{nT}{2}\log(2\pi) - \frac{1}{2}\sum_{t=1}^T \left\{ \log|\boldsymbol{\Sigma}_t| + (\boldsymbol{\eta}_t'\boldsymbol{\Sigma}_t^{-1}\boldsymbol{\eta}_t) \right\}.$$

Estimated parameters and state variables can be obtained accordingly.

2.5 EM algorithm for dynamic linear models

Shumway and Stoffer (1982) presented a conceptually simpler estimation procedure based on the EM (expectation maximization) algorithm (Dempster et al., 1977).

The basic idea is that if we could observe the states $X_T = \{\xi_0, \xi_1, \dots, \xi_T\}$ and the observations $Y_T = \{\mathbf{y}_1, \dots, \mathbf{y}_T\}$, then we would consider $\{X_T, Y_T\}$ as the complete data.

The joint density for the complete data $\{X_T, Y_T\}$ is

$$f_{\Theta}(X_T, Y_T) = f_{\mathbf{a}_0, \Sigma_0}(\xi_0) \prod_{t=1}^T f_{F, Q}(\xi_t | \xi_{t-1}) \prod_{t=1}^T f_{H, R}(\mathbf{y}_t | \xi_t).$$

Under the Gaussian assumption and ignoring constants, the complete data log-likelihood can be written as

$$\begin{aligned}
 -2 \ln L_{X,Y}(\Theta) = & \ln |\Sigma_0| + (\xi_0 - \mathbf{a}_0)' \Sigma_0^{-1} (\xi_0 - \mathbf{a}_0) \\
 & + T \ln |Q| + \sum_{t=1}^T (\xi_t - F \xi_{t-1})' Q^{-1} (\xi_t - F \xi_{t-1}) \\
 & + T \ln |R| + \sum_{t=1}^T (y_t - H \xi_t)' R^{-1} (y_t - H \xi_t)
 \end{aligned}$$

- If we did have the complete data, we could then use the results from multivariate normal theory to easily obtain the MLEs of Θ .
- However, we do not have the complete data, the EM algorithm gives us an iterative method for finding the MLEs of Θ based on the incomplete data, Y_T , by successively maximizing the conditional expectation of the complete data likelihood.

(1) Expectation step (E-Step):

To implement the EM algorithm, we write, at iteration j , ($j = 1, 2, \dots$),

$$Q(\Theta \mid \Theta^{(j-1)}) = E \left\{ -2 \ln L_{X,Y}(\Theta) \mid Y_T, \Theta^{(j-1)} \right\}$$

Given $\Theta^{(j-1)}$, we can use the Kalman Smoother to obtain the desired conditional expectations as smoothers. This yields

$$\begin{aligned} Q(\Theta \mid \Theta^{(j-1)}) = & \ln |\Sigma_0| + \text{tr} \{ \Sigma_0^{-1} E[(\xi_0 - \xi_{0|T} + \xi_{0|T} - \mathbf{a}_0)(\xi_0 - \xi_{0|T} + \xi_{0|T} - \mathbf{a}_0)' | \mathcal{F}_T] \} \\ & + T \ln |\mathbf{Q}| + \text{tr} \left\{ \mathbf{Q}^{-1} \sum_{t=1}^T E \left[\underbrace{(\xi_t - \mathbf{F}\xi_{t-1})(\xi_t - \mathbf{F}\xi_{t-1})'}_{=\xi_t \xi_t' - \xi_t \xi_{t-1}' \mathbf{F}' - \mathbf{F} \xi_{t-1} \xi_t' + \mathbf{F} \xi_{t-1} \xi_{t-1}' \mathbf{F}'} \mid \mathcal{F}_T \right] \right\} \\ & + T \ln |\mathbf{R}| + \text{tr} \left\{ \mathbf{R}^{-1} \sum_{t=1}^T E \left[\underbrace{(\mathbf{y}_t - \mathbf{H}\xi_t)}_{\text{insert } -\mathbf{H}\xi_{t|T} + \mathbf{H}\xi_{t|T}} \underbrace{(\mathbf{y}_t - \mathbf{H}\xi_t)'}_{\text{insert } -\mathbf{H}\xi_{t|T} + \mathbf{H}\xi_{t|T}} \mid \mathcal{F}_T \right] \right\} \end{aligned}$$

Using the facts that

$$E[\boldsymbol{\xi}_t \boldsymbol{\xi}_t' | \mathcal{F}_T] = E[\boldsymbol{\xi}_t | \mathcal{F}_T] E[\boldsymbol{\xi}_t' | \mathcal{F}_T] + \text{Var}[\boldsymbol{\xi}_t | \mathcal{F}_T] = \boldsymbol{\xi}_{t|T} \boldsymbol{\xi}_{t|T}' + \mathbf{P}_{t|T},$$

$$E[\boldsymbol{\xi}_t \boldsymbol{\xi}_{t-1}' | \mathcal{F}_T] = \boldsymbol{\xi}_{t|T} \boldsymbol{\xi}_{t-1|T}' + \mathbf{P}_{t,t-1|T},$$

$$E[\boldsymbol{\xi}_{t-1} \boldsymbol{\xi}_{t-1}' | \mathcal{F}_T] = \boldsymbol{\xi}_{t-1|T} \boldsymbol{\xi}_{t-1|T}' + \mathbf{P}_{t-1|T},$$

we have

$$\begin{aligned} Q(\boldsymbol{\Theta} \mid \boldsymbol{\Theta}^{(j-1)}) &= \ln |\boldsymbol{\Sigma}_0| + \text{tr}\{\boldsymbol{\Sigma}_0^{-1} [\mathbf{P}_{0|T} + (\boldsymbol{\xi}_{0|T} - \mathbf{a}_0)(\boldsymbol{\xi}_{0|T} - \mathbf{a}_0)']'\} \\ &\quad + T \ln |\mathbf{Q}| + \text{tr}\{\mathbf{Q}^{-1} [\mathbf{S}_{11} - \mathbf{S}_{10} \mathbf{F}' - \mathbf{F} \mathbf{S}_{10}' + \mathbf{F} \mathbf{S}_{00} \mathbf{F}']\} \\ &\quad + T \ln |\mathbf{R}| \\ &\quad + \text{tr} \left\{ \mathbf{R}^{-1} \sum_{t=1}^T [(\mathbf{y}_t - \mathbf{H} \boldsymbol{\xi}_{t|T})(\mathbf{y}_t - \mathbf{H} \boldsymbol{\xi}_{t|T})' + \mathbf{H} \mathbf{P}_{t|T} \mathbf{H}'] \right\}, \end{aligned} \tag{55}$$

where

$$\begin{aligned} S_{11} &= \sum_{t=1}^T (\xi_{t|T} \xi'_{t|T} + P_{t|T}), & S_{10} &= \sum_{t=1}^T (\xi_{t|T} \xi'_{t-1|T} + P_{t,t-1|T}) \\ S_{00} &= \sum_{t=1}^T (\xi_{t-1|T} \xi'_{t-1|T} + P_{t-1|T}) \end{aligned}$$

In the above equations, the smoothers are calculated under the current value of the parameters $\Theta^{(j-1)}$; for simplicity, we have not explicitly displayed this fact.

(2) Maximization step (M-Step):

Minimizing (55) with respect to the unknown parameter set, $\Theta = \{H, F, Q, R\}$, constitutes the maximization step.

At iteration j , for the parameters F and Q , this yields the updated estimates,

$$F^{(j)} = S_{10}S_{00}^{-1}$$
$$Q^{(j)} = T^{-1}(S_{11} - S_{10}S_{00}^{-1}S_{10}').$$

To develop estimators for H and R from the last term of (55), it will reduce to minimizing: $\sum_{t=1}^T E[(y_t - H\xi_t)(y_t - H\xi_t)' | Y_T, \Theta^{(j-1)}]$ over H and then evaluating R .

Performing the first operation, we obtain

$$\mathbf{H}^{(j)} = \sum_{t=1}^T \mathbf{y}_t \boldsymbol{\xi}_{t|T}' \mathbf{S}_{11}^{-1}$$

Then, we can obtain estimator for $\mathbf{R}$ as:

$$\mathbf{R}^{(j)} = T^{-1} \sum_{t=1}^T [(\mathbf{y}_t - \mathbf{H}^{(j)} \boldsymbol{\xi}_{t|T})(\mathbf{y}_t - \mathbf{H}^{(j)} \boldsymbol{\xi}_{t|T})' + \mathbf{H}^{(j)} \mathbf{P}_{t|T} \mathbf{H}^{(j)'}].$$

Finally, the updates for the initial mean and covariance matrix are

$$\mathbf{a}_0^{(j)} = \boldsymbol{\xi}_{0|T} \quad \text{and} \quad \boldsymbol{\Sigma}_0^{(j)} = (\boldsymbol{\xi}_{0|T} - \mathbf{a}_0^{(j)})(\boldsymbol{\xi}_{0|T} - \mathbf{a}_0^{(j)})' + \mathbf{P}_{0|T}.$$

EM algorithm:

1. Initialize the procedure by selecting starting values for the parameters $\Theta^{(0)} = \{\mathbf{a}_0, \Sigma_0, \mathbf{F}, \mathbf{Q}, \mathbf{H}, \mathbf{R}\}$.

On iteration j , ($j = 1, 2, \dots$):

2. Run the Kalman filter and smoother, and compute the incomplete-data likelihood, $-\ln L_{\mathcal{F}}(\Theta^{(j-1)})$.
3. Perform the E-Step.
4. Perform the M-Step.
5. Repeat Steps 2 - 4 to convergence.

3 Approximate Filters for Nonlinear SS Models

3.1 Extended Kalman Filter

Consider next a multidimensional normal state-space model

$$\mathbf{y}_t = \mathbf{h}(\boldsymbol{\xi}_t) + \mathbf{d} + \mathbf{u}_t, \quad (56)$$

$$\boldsymbol{\xi}_t = \boldsymbol{\phi}(\boldsymbol{\xi}_{t-1}) + \mathbf{c} + \mathbf{v}_t, \quad (57)$$

where $\mathbf{v}_t \sim i.i.d.N(\mathbf{0}, \mathbf{Q})$ and $\mathbf{u}_t \sim i.i.d.N(\mathbf{0}, \mathbf{R})$. Suppose $\boldsymbol{\phi}(\cdot)$ in (57) is replaced with a first-order Taylor's approximation around $\boldsymbol{\xi}_{t-1} = \boldsymbol{\xi}_{t-1|t-1}$,

$$\boldsymbol{\xi}_t = \boldsymbol{\phi}_t + \boldsymbol{\Phi}_t \cdot (\boldsymbol{\xi}_{t-1} - \boldsymbol{\xi}_{t-1|t-1}) + \mathbf{c} + \mathbf{v}_t, \quad (58)$$

where

$$\phi_t = \phi(\xi_{t-1|t-1}), \quad \Phi_t = \left. \frac{\partial \phi(\xi_{t-1})}{\partial \xi'_{t-1}} \right|_{\xi_{t-1} = \xi_{t-1|t-1}}.$$

If the form of ϕ and any parameters it depends on are known, then the inference $\xi_{t|t}$, can be constructed as a function of variables observed at date t through a recursion to be described in a moment.

Similarly the function $h(\cdot)$ in (56) can be linearized around $\xi_{t|t-1}$ to produce

$$y_t = h_t + H_t \cdot (\xi_t - \xi_{t|t-1}) + d + u_t. \quad (59)$$

where

$$h_t = h(\xi_{t|t-1}), \quad H_t = \left. \frac{\partial h(\xi_t)}{\partial \xi'_t} \right|_{\xi_t = \xi_{t|t-1}}$$

Treat (58) and (59) as if they were the true model, then the *extended Kalman filter* is implemented as follows:

$$\xi_{t|t-1} = \phi(\xi_{t-1|t-1}) + c \quad (60)$$

$$P_{t|t-1} = \Phi_t P_{t-1|t-1} \Phi_t' + Q \quad (61)$$

$$y_{t|t-1} = h_t + d = h(\xi_{t|t-1}) + d \quad (62)$$

$$\eta_t = y_t - y_{t|t-1} = H_t(\xi_t - \xi_{t|t-1}) + u_t \quad (63)$$

$$\Sigma_t = H_t P_{t|t-1} H_t' + R \quad (64)$$

$$\xi_{t|t} = \xi_{t|t-1} + P_{t|t-1} H_t (H_t' P_{t|t-1} H_t + R)^{-1} \eta_t \quad (65)$$

$$P_{t|t} = P_{t|t-1} - P_{t|t-1} H_t (H_t' P_{t|t-1} H_t + R)^{-1} H_t' P_{t|t-1} \quad (66)$$

3.2 State Space Models with conditional heteroskedasticity

We consider the following state-space model with ARCH disturbances proposed by Harvey, Ruiz and Sentana (1992):

$$\mathbf{y}_t = \mathbf{H}\boldsymbol{\xi}_t + \mathbf{d} + \mathbf{u}_t + \Lambda\epsilon_t \quad (67)$$

$$\boldsymbol{\xi}_t = \mathbf{F}\boldsymbol{\xi}_{t-1} + \mathbf{c} + \mathbf{v}_t + \lambda\omega_t \quad (68)$$

$$\mathbf{u}_t \sim N(0, \mathbf{R}), \quad \mathbf{v}_t \sim N(0, \mathbf{Q}), \quad (69)$$

where Λ is $n \times 1$ and λ is $k \times 1$, and ϵ_t and ω_t are 1×1 . The ARCH effects are introduced via the following scalar disturbances:

$$\epsilon_t | I_{t-1} \sim N(0, h_{1t}), \quad h_{1t} = \alpha_0 + \alpha_1 \epsilon_{t-1}^2,$$

$$\omega_t | I_{t-1} \sim N(0, h_{2t}), \quad h_{2t} = \gamma_0 + \gamma_1 \omega_{t-1}^2.$$

- If $\Lambda = 0$ and $\lambda = 0$, then there is no ARCH effect and the model reduces to the usual time-invariant model.
- If the first element of Λ is nonzero with the others being zero, only the first shock in the measurement equation is subject to an ARCH effect.
- If all the elements of Λ are nonzero, then all the shocks in the measurement equation are subject to a common ARCH effect.
- The same argument applies for the transition equation.
- If we define $\mathbf{u}_t^* = \mathbf{u}_t + \Lambda\epsilon_t$ and $\mathbf{v}_t^* = \mathbf{v}_t + \lambda\omega_t$, then we have $\mathbf{u}_t^*|I_{t-1} \sim N(0, \mathbf{R}_t)$ and $\mathbf{v}_t^*|I_{t-1} \sim N(0, \mathbf{Q}_t)$, where

$$\mathbf{R}_t = \mathbf{R} + \Lambda h_{1t} \Lambda' \quad \text{and} \quad \mathbf{Q}_t = \mathbf{Q} + \lambda h_{2t} \lambda'.$$

Difficulty in Kalman filter

$$\xi_{t|t-1} = F\xi_{t-1|t-1} + c$$

$$P_{t|t-1} = FP_{t-1|t-1}F' + Q + \lambda h_{2t}\lambda'$$

$$\eta_t = y_t - H\xi_{t|t-1} - d$$

$$\Sigma_t = HP_{t|t-1}H' + R + \Lambda h_{1t}\Lambda'$$

$$\xi_{t|t} = \xi_{t|t-1} + P_{t|t-1}H\Sigma_t^{-1}\eta_t$$

$$P_{t|t} = P_{t|t-1} - P_{t|t-1}H\Sigma_t^{-1}H'P_{t|t-1}$$

However, we need to calculate $h_{1t} = \alpha_0 + \alpha_1\epsilon_{t-1}^2$ and $h_{2t} = \gamma_0 + \gamma_1\omega_{t-1}^2$, which are functions of the squares of past unobserved shocks ϵ_{t-1}^2 and ω_{t-1}^2 . Thus, the above Kalman filter is not operable.

Approximate Kalman filter

Harvey, Ruiz and Sentana (1992) solve the problem of unobserved ϵ_{t-1}^2 and ω_{t-1}^2 by replacing them with their conditional expectations:

$$h_{1t} = \alpha_0 + \alpha_1 E[\epsilon_{t-1}^2 | I_{t-1}] \quad \text{and} \quad h_{2t} = \gamma_0 + \gamma_1 E[\omega_{t-1}^2 | I_{t-1}].$$

To get the conditional expectations of ϵ_{t-1}^2 and ω_{t-1}^2 , Harvey, Ruiz and Sentana augment the heteroskedastic shocks into the original state vector in the transition equation as follows:

$$\begin{bmatrix} \xi_t \\ \epsilon_t \\ \omega_t \end{bmatrix} = \begin{bmatrix} c \\ 0 \\ 0 \end{bmatrix} + \begin{bmatrix} F & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{bmatrix} \begin{bmatrix} \xi_{t-1} \\ \epsilon_{t-1} \\ \omega_{t-1} \end{bmatrix} + \begin{bmatrix} I & 0 & \lambda \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} v_t \\ \epsilon_t \\ \omega_t \end{bmatrix}$$
$$(\xi_t^* = c^* + F^* \xi_{t-1}^* + G^* v_t^*)$$

where

$$E[\mathbf{v}_t^* \mathbf{v}_t^{*'} | I_{t-1}] = \begin{bmatrix} \mathbf{Q} & 0 & 0 \\ 0 & h_{1t} & 0 \\ 0 & 0 & h_{2t} \end{bmatrix} = \mathbf{Q}_t^*$$

The measurement equation should be replaced by

$$\mathbf{y}_t = \begin{bmatrix} \mathbf{H} & \Lambda & 0 \end{bmatrix} \begin{bmatrix} \boldsymbol{\xi}_t \\ \epsilon_t \\ \omega_t \end{bmatrix} + \mathbf{d} + \mathbf{u}_t, \quad \mathbf{u}_t \sim N(0, \mathbf{R})$$

$$(\mathbf{y}_t = \mathbf{H}^* \boldsymbol{\xi}_t^* + \mathbf{d} + \mathbf{u}_t)$$

Applying the Kalman filter recursions to the model, we have

$$\begin{aligned}
 \xi_{t|t-1}^* &= F^* \xi_{t-1|t-1}^* + c^* \\
 P_{t|t-1}^* &= F^* P_{t-1|t-1}^* F^{*'} + G^* Q_t^* G^{*'} \\
 \eta_t^* &= y_t - H^* \xi_{t|t-1}^* - d^* \\
 \Sigma_t^* &= H^* P_{t|t-1}^* H^{*'} + R^* \\
 \xi_{t|t}^* &= \xi_{t|t-1}^* + P_{t|t-1}^* H^* \Sigma_t^{*-1} \eta_t^* \\
 P_{t|t}^* &= P_{t|t-1}^* - P_{t|t-1}^* H^* \Sigma_t^{*-1} H^{*'} P_{t|t-1}^*
 \end{aligned}$$

In the Q_t^* matrix,

$$h_{1t} = \alpha_0 + \alpha_1 E[\epsilon_{t-1}^2 | I_{t-1}] \quad \text{and} \quad h_{2t} = \gamma_0 + \gamma_1 E[\omega_{t-1}^2 | I_{t-1}].$$

It is easy to show that

$$E[\epsilon_{t-1}^2 | I_{t-1}] = E[\epsilon_{t-1} | I_{t-1}]^2 + E[(\epsilon_{t-1} - E[\epsilon_{t-1} | I_{t-1}])^2],$$

$$E[\omega_{t-1}^2 | I_{t-1}] = E[\omega_{t-1} | I_{t-1}]^2 + E[(\omega_{t-1} - E[\omega_{t-1} | I_{t-1}])^2],$$

where

- $E[\epsilon_{t-1} | I_{t-1}]$ and $E[\omega_{t-1} | I_{t-1}]$ are obtained from the last two elements of $\xi_{t-1|t-1}^*$;
- $E[(\epsilon_{t-1} - E[\epsilon_{t-1} | I_{t-1}])^2]$ and $E[(\omega_{t-1} - E[\omega_{t-1} | I_{t-1}])^2]$ are obtained from the last two diagonal elements of $P_{t-1|t-1}^*$.

Example 1 (TVP model with GARCH disturbances).

$$y_t = X_{t-1}\beta_t + \varepsilon_t^*, \quad \varepsilon_t^* | I_{t-1} \sim N(0, h_t)$$

$$\beta_t = \beta_{t-1} + v_t, \quad v_t \sim N(0, Q)$$

$$h_t = \alpha_0 + \alpha_1 \varepsilon_{t-1}^{*2} + \alpha_2 h_{t-1},$$

3.3 State Space Models with Markov Switching

Consider the following state-space representation of a dynamic linear model with Markov switching in both measurement and transition equations:

$$\mathbf{y}_t = \mathbf{H}_{S_t} \boldsymbol{\xi}_t + \mathbf{d}_{S_t} + \mathbf{u}_t, \quad (70)$$

$$\boldsymbol{\xi}_t = \mathbf{F}_{S_t} \boldsymbol{\xi}_{t-1} + \mathbf{c}_{S_t} + \mathbf{v}_t, \quad (71)$$

$$\begin{pmatrix} \mathbf{u}_t \\ \mathbf{v}_t \end{pmatrix} \sim N \left(\mathbf{0}, \begin{pmatrix} \mathbf{R}_{S_t} & \mathbf{0} \\ \mathbf{0} & \mathbf{Q}_{S_t} \end{pmatrix} \right), \quad (72)$$

$$S_t | S_{t-1} \sim \text{Markov}(\mathbf{P}) \quad (73)$$

where the transition probability matrix $\mathbf{P} = (p_{ij})$, $p_{ij} = \Pr(S_t = j | S_{t-1} = i)$ satisfying that $\sum_{j=1}^M p_{ij} = 1$ for all i .

(1) Multiprocess Kalman filter

In derivation of the filtering and estimation of the model, we need to solve two filtering problems.

We begin with the recursion between $\xi_{t-1|t-1}^i$ (and $P_{t-1|t-1}^i$) and $\xi_{t|t}^j$ (and $P_{t|t}^j$). Conditional on $S_{t-1} = i$ and $S_t = j$, $i, j = 1, \dots, M$, we have

$$\xi_{t|t-1}^{ij} = F_j \xi_{t-1|t-1}^i + c_j, \quad (74)$$

$$P_{t|t-1}^{ij} = F_j P_{t-1|t-1}^i F_j' + Q_j, \quad (75)$$

$$\eta_{t|t-1}^{ij} = y_t - H_j \xi_{t|t-1}^{ij} - d_j, \quad (76)$$

$$\Sigma_{t|t-1}^{ij} = H_j P_{t|t-1}^{ij} H_j' + R_j, \quad (77)$$

$$\xi_{t|t}^{ij} = \xi_{t|t-1}^{ij} + P_{t|t-1}^{ij} H_j' [\Sigma_{t|t-1}^{ij}]^{-1} \eta_{t|t-1}^{ij}, \quad (78)$$

$$P_{t|t}^{ij} = (I - P_{t|t-1}^{ij} H_j' [\Sigma_{t|t-1}^{ij}]^{-1} H_j) P_{t|t-1}^{ij}. \quad (79)$$

To reduce the dimension of the estimation problem, a *re-collapsing* procedure is suggested by Harrison and Stevens (1976) which effectively provides *re-approximation* at each time t .

These approximations are given by

$$\begin{aligned}
 \boldsymbol{\xi}_{t|t}^j &= E[\boldsymbol{\xi}_t | S_t = j, \mathcal{F}_t] \\
 &= \sum_{i=1}^M \Pr(S_{t-1} = i | S_t = j, \mathcal{F}_t) E[\boldsymbol{\xi}_t | S_t = j, S_{t-1} = i, \mathcal{F}_t] \\
 &= \sum_{i=1}^M w_{t|t}^{ij} \boldsymbol{\xi}_{t|t}^{ij},
 \end{aligned} \tag{80}$$

$$\begin{aligned}
 \mathbf{P}_{t|t}^j &= \sum_{i=1}^M w_{t|t}^{ij} \{ \mathbf{P}_{t|t}^{ij} + (\boldsymbol{\xi}_{t|t}^j - \boldsymbol{\xi}_{t|t}^{ij})(\boldsymbol{\xi}_{t|t}^j - \boldsymbol{\xi}_{t|t}^{ij})' \},
 \end{aligned} \tag{81}$$

where

$$\begin{aligned} w_{t|t}^{ij} &= \Pr(S_{t-1} = i | S_t = j, \mathcal{F}_t) \\ &= \frac{\Pr(S_{t-1} = i, S_t = j | \mathcal{F}_t)}{\Pr(S_t = j | \mathcal{F}_t)}, \end{aligned} \tag{82}$$

$\xi_{t|t}^j$ is an inference about ξ_t based on information up to time t given $S_t = i$, and $P_{t|t}^j$ is the corresponding mean squared error.

(2) Hamilton filter

In equation (82), the probability terms $\Pr(S_{t-1} = i, S_t = j | \mathcal{F}_t)$ and $\Pr(S_t = j | \mathcal{F}_t)$ remain unknown. This leads to solving *the second filtering problem* as suggested by Hamilton (1989).

Given the term $\Pr(S_{t-1} = i | \mathcal{F}_{t-1})$ (for $i = 1, 2, \dots, M$) at the beginning of the t -th iteration, we can compute the joint density of $S_{t-1} = i$ and $S_t = j$ conditional on $\mathcal{F}_{t-1}$, that is

$$\Pr(S_{t-1} = i, S_t = j | \mathcal{F}_{t-1}) = p_{ij} \Pr(S_{t-1} = i | \mathcal{F}_{t-1}) \quad (83)$$

for $i, j = 1, \dots, M$, where $p_{ij} = \Pr(S_t = j | S_{t-1} = i)$ is the *transition probability*.

By *Bayes' theorem*, we have

$$\begin{aligned}\Pr(S_{t-1} = i, S_t = j | \mathcal{F}_t) &= \frac{f(\mathbf{y}_t, S_{t-1} = i, S_t = j | \mathcal{F}_{t-1})}{f(\mathbf{y}_t | \mathcal{F}_{t-1})} \\ &= \frac{f(\mathbf{y}_t | S_{t-1} = i, S_t = j, \mathcal{F}_{t-1}) \Pr(S_{t-1} = i, S_t = j | \mathcal{F}_{t-1})}{f(\mathbf{y}_t | \mathcal{F}_{t-1})}.\end{aligned}\tag{84}$$

where

$$f(\mathbf{y}_t | S_{t-1} = i, S_t = j, \mathcal{F}_{t-1}) = \frac{1}{(2\pi)^{\frac{n}{2}} |\boldsymbol{\Sigma}_{t|t-1}^{ij}|^{\frac{1}{2}}} \exp \left\{ -\frac{\boldsymbol{\eta}_{t|t-1}^{ij}{}' (\boldsymbol{\Sigma}_{t|t-1}^{ij})^{-1} \boldsymbol{\eta}_{t|t-1}^{ij}}{2} \right\},$$

and so the *conditional likelihood function* at time t is given by

$$f(\mathbf{y}_t | \mathcal{F}_{t-1}) = \sum_{j=1}^M \sum_{i=1}^M \Pr(S_{t-1} = i, S_t = j | \mathcal{F}_{t-1}) f(\mathbf{y}_t | S_{t-1} = i, S_t = j, \mathcal{F}_{t-1}).\tag{85}$$

For given parameters of the model, the Kalman filter provides us with the forecast error ($\eta_{t|t-1}^{ij}$) and its variance ($\Sigma_{t|t-1}^{ij}$) as by-products.

Also, the filtered density of S_t based on the information up to time t is given by

$$\Pr(S_t = j|\mathcal{F}_t) = \sum_{i=1}^M \Pr(S_{t-1} = i, S_t = j|\mathcal{F}_t). \quad (86)$$

Then we finish the second recursion between $\Pr(S_{t-1} = i|\mathcal{F}_{t-1})$ and $\Pr(S_t = j|\mathcal{F}_t)$.

(3) Kim's (1994) filter

Kim's (1994) filter for the state-space model with Markov switching may be viewed as a *combination* of the *Kalman filter* and the filter provided by Hamilton (1989) along with appropriate approximations.

The Kim's (1994) filter consists of the following steps:

Algorithm 1 (Kim filter).

1. *Run the Kalman filter given in the Eqs. (74) to (79) for $i, j = 1, \dots, M$.*
2. *Calculate $\Pr(S_t = j, S_{t-1} = i | \mathcal{F}_t)$ and $\Pr(S_t = j | \mathcal{F}_t)$ for $i, j = 1, \dots, M$.*
3. *Using these probability terms, collapse the $(M \times M)$ posteriors in the Eqs. (78) and (79) into $(M \times 1)$ using the Eqs. (80) and (81).*

(4) Log-likelihood

As a by-product of running the above Kalman filter, the conditional log-likelihood function can be obtained from step 2 using the equation (85).

The sample conditional log-likelihood is

$$\begin{aligned}\mathcal{L}(\theta) &= \sum_{t=1}^T \log f(\mathbf{y}_t | \mathcal{F}_{t-1}) \\ &= \sum_{t=1}^T \log \left\{ \sum_{i,j} \Pr(S_{t-1} = i, S_t = j | \mathcal{F}_{t-1}) f(\mathbf{y}_t | S_t = j, S_{t-1} = i, \mathcal{F}_{t-1}) \right\} \\ &= \sum_{t=1}^T \log \left\{ \sum_{i,j} \frac{\Pr(S_{t-1} = i, S_t = j | \mathcal{F}_{t-1})}{(2\pi)^{\frac{n}{2}} |\Sigma_{t|t-1}^{ij}|^{\frac{1}{2}}} \exp \left(-\frac{\boldsymbol{\eta}_{t|t-1}^{ij'} (\Sigma_{t|t-1}^{ij})^{-1} \boldsymbol{\eta}_{t|t-1}^{ij}}{2} \right) \right\},\end{aligned}$$

where $i, j \in \{1, \dots, M\}$.

(5) Smoothing algorithm

- Once parameters are estimated, we can make inferences about S_t and ξ_t by using the entire information in the sample $\mathcal{F}_T$, that is we can determine $\Pr(S_t = j|\mathcal{F}_T)$ and $\xi_{t|T}$ for $t = 1, 2, \dots, T$.
- For this, consider the *joint probability* of the event that $S_t = j$ and $S_{t+1} = k$ based on all information from the full sample:

$$\begin{aligned}\Pr(S_t = j, S_{t+1} = k|\mathcal{F}_T) &= \Pr(S_{t+1} = k|\mathcal{F}_T) \times \Pr(S_t = j|S_{t+1} = k, \mathcal{F}_T) \\ &= \Pr(S_{t+1} = k|\mathcal{F}_T) \times \Pr(S_t = j|S_{t+1} = k, \mathcal{F}_t) \\ &= \Pr(S_{t+1} = k|\mathcal{F}_T) \times \frac{\Pr(S_t = j, S_{t+1} = k|\mathcal{F}_t)}{\Pr(S_{t+1} = k|\mathcal{F}_t)} \\ &= \Pr(S_{t+1} = k|\mathcal{F}_T) \times \frac{p_{jk} \Pr(S_t = j|\mathcal{F}_t)}{\Pr(S_{t+1} = k|\mathcal{F}_t)}\end{aligned}\tag{87}$$

and

$$\Pr(S_t = j | \mathcal{F}_T) = \sum_{k=1}^M \Pr(S_t = j, S_{t+1} = k | \mathcal{F}_T). \quad (88)$$

- Keeping Eq.(87) in mind, we now turn to the derivation of the smoothing algorithm for the vector ξ_t .
- Given that $S_t = j$ and $S_{t+1} = k$, the smoothing algorithm can be written as follows:

$$\xi_{t|T}^{(j,k)} = \xi_{t|t}^j + \Upsilon_t^{(j,k)} (\xi_{t+1|T}^k - \xi_{t+1|t}^{(j,k)}), \quad (89)$$

$$P_{t|T}^{(j,k)} = P_{t|t}^j + \Upsilon_t^{(j,k)} (P_{t+1|T}^k - P_{t+1|t}^{(j,k)}) \Upsilon_t^{(j,k)'} \quad (90)$$

where $\Upsilon_t^{(j,k)} = P_{t|t}^j F_k' [P_{t+1|t}^{(j,k)}]^{-1}$.

Since $\Pr(S_t = j|\mathcal{F}_T)$ does not depend on $\xi_{t|T}$, we can first calculate smoothed probabilities which we can use then to obtain smoothed values of ξ_t , i.e. $\xi_{t|T}$.

So the above smoothing algorithm can be performed by applying the following approximation approach:

Algorithm 2 (Smoothing algorithm).

1. *We run through the basic filter and store the resulting sequences $\xi_{t|t}^{(i,j)}$, $P_{t|t}^{(i,j)}$, $\xi_{t|t}^j$, $P_{t|t}^j$, $\Pr(S_t = j|\mathcal{F}_{t-1})$ and $\Pr(S_t = j|\mathcal{F}_t)$.*
2. *For $t = T - 1, T - 2, \dots, 1$, we obtain the smoothed probabilities $\Pr(S_t = j, S_{t+1} = k|\mathcal{F}_T)$ and $\Pr(S_t = j|\mathcal{F}_T)$ according to (87) and (88) and save these. Here, $\Pr(S_T = j|\mathcal{F}_T)$, the initial value for smoothing, is given by the final iteration of the basic filter.*

3. At each iteration of (89) and (90) for $t = T-1, T-2, \dots, 1$, we collapse the $M \times M$ elements of the matrices $\boldsymbol{\xi}_{t|T}^{(j,k)}$ and $\boldsymbol{P}_{t|T}^{(j,k)}$ into M weighted averages, by taking weighted averages over the state variable S_{t+1} :

$$\begin{aligned}\boldsymbol{\xi}_{t|T}^j &= \sum_{k=1}^M w_{t|T}^{jk} \boldsymbol{\xi}_{t|T}^{(j,k)}, \\ \boldsymbol{P}_{t|T}^j &= \sum_{k=1}^M w_{t|T}^{jk} \left\{ \boldsymbol{P}_{t|T}^{(j,k)} + (\boldsymbol{\xi}_{t|T}^j - \boldsymbol{\xi}_{t|T}^{(j,k)})(\boldsymbol{\xi}_{t|T}^j - \boldsymbol{\xi}_{t|T}^{(j,k)})' \right\},\end{aligned}$$

where

$$w_{t|T}^{jk} = \frac{\Pr(S_t = j, S_{t+1} = k | \mathcal{F}_T)}{\Pr(S_t = j | \mathcal{F}_T)}.$$

4. By taking a weighted average over the M Markov-regime at time t , the smoothed value is computed as $\boldsymbol{\xi}_{t|T} = \sum_{j=1}^M \Pr(S_t = j | \mathcal{F}_T) \boldsymbol{\xi}_{t|T}^j$.

(6) Markov switching SV model

Consider the following Markov switching stochastic volatility model:

$$y_t = \sigma_t \varepsilon_t, \quad \varepsilon_t \sim \mathcal{N}(0, 1), \quad (91)$$

$$h_t = c_{s_t} + \lambda h_{t-1} + \zeta_t, \quad \zeta_t \sim \mathcal{N}(0, \sigma_\zeta^2). \quad (92)$$

The discrete state variable s_t is assumed to be a first-order Markov process with the *transition probabilities*

$$\Pr(s_t = 0 | s_{t-1} = 0) = p_{00}, \quad \Pr(s_t = 0 | s_{t-1} = 1) = 1 - p_{11},$$

$$\Pr(s_t = 1 | s_{t-1} = 0) = 1 - p_{00}, \quad \Pr(s_t = 1 | s_{t-1} = 1) = p_{11}.$$

Taking the square and the logarithm, Equation (6) can be rewritten as:

$$\ln(y_t^2) = h_t + \ln(\varepsilon_t^2).$$

We then represent the model as follows:

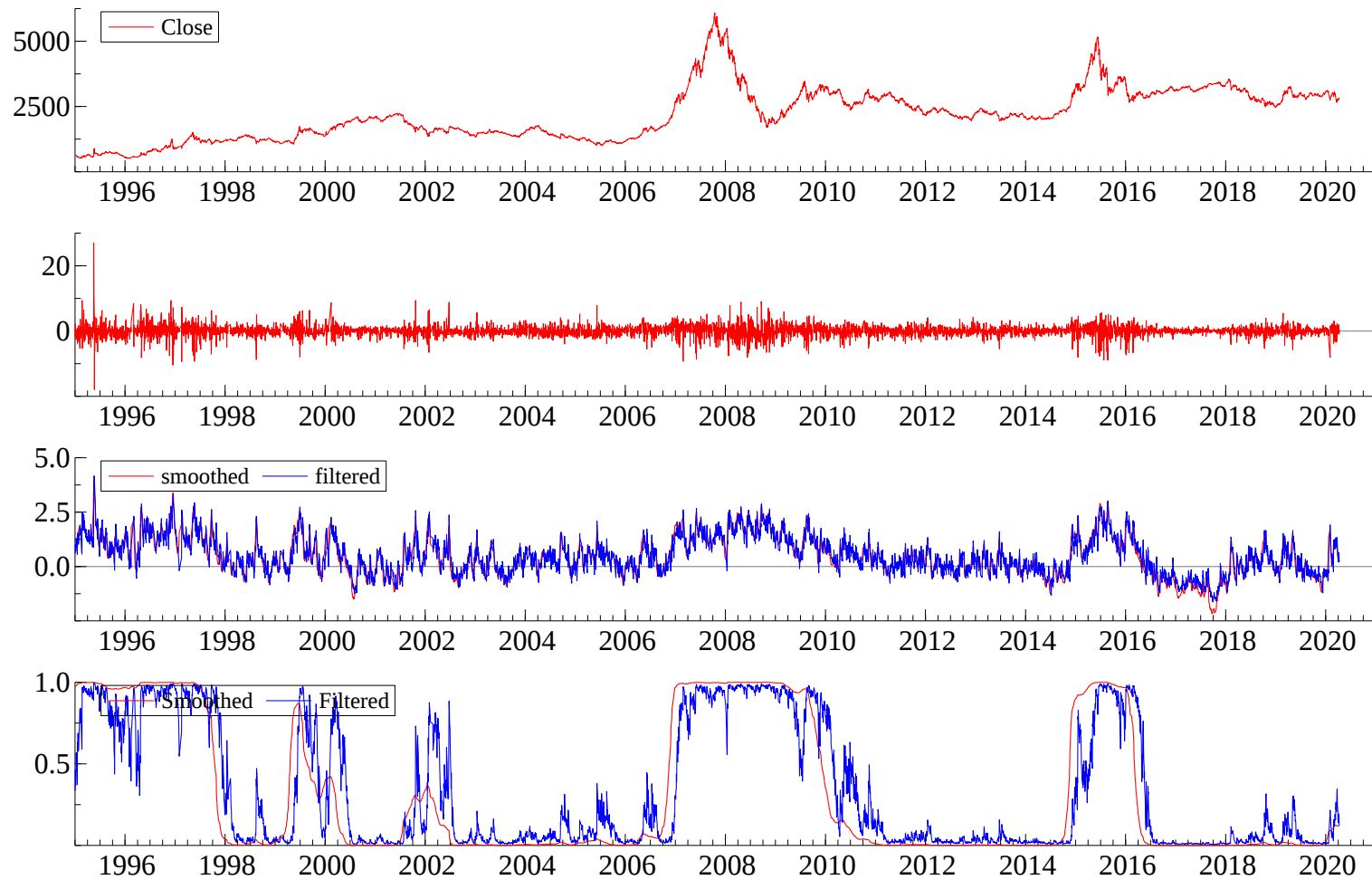
$$\ln(y_t^2) = h_t - 1.2704 + u_t, \quad u_t \sim N(0, \pi^2/2),$$

$$h_t = c_{s_t} + \lambda h_{t-1} + \zeta_t, \quad \zeta_t \sim \mathcal{N}(0, \sigma_\zeta^2),$$

$$s_t | s_{t-1}, \mathcal{F}_{t-1} \sim \text{Markov}(P),$$

$$P = \begin{bmatrix} p_{00} & 1 - p_{11} \\ 1 - p_{00} & p_{11} \end{bmatrix}.$$

Example 2. Fitting Shanghai composite index with the MSSV model. The sample period is from 1995/01/01 to 2020/4/10.



The parameter estimates are reported as follows:

	Coef.Est	Std.Err	t-value	t-prob
p11	0.998794	0.000697	1433.840978	0.000000
p22	0.997625	0.001342	743.540344	0.000000
nu1	0.005172	0.006971	0.741902	0.458175
nu2	0.135570	0.025345	5.348979	0.000000
phi	0.906501	0.016639	54.480587	0.000000
sigma	0.341337	0.030876	11.055149	0.000000
=====				
Log likelihood	= -13571.1			
Akaika Information Criterion	= 4.42468			
Schwartz Bayesian Criterion	= 4.43125			

4 Mixed-Frequency Models and Applications

- Mariano and Murasawa (2003): A new coincident index of business cycles based on monthly and quarterly series. *Journal of Applied Econometrics*, 18, 427 – 443.
- 郑挺国、王霞 (2013): 中国经济周期的混频数据测度及实时分析, 《经济研究》第6期, 58-70.
- Mariano and Murasawa (2010): A Coincident Index, Common Factors, and Monthly Real GDP, *Oxford Bulletin of Economics and Statistics*, 72(1), 27-46.

4.1 Mixed-frequency dynamic factor model

- $\{Y_{1,t}\}_{t=-\infty}^{\infty}$: an N_1 -variate random sequence of quarterly indicators observable every third period; $\{Y_{2,t}\}_{t=-\infty}^{\infty}$: an N_2 -variate random sequence of monthly indicators;
- $N = N_1 + N_2$;
- Assume that logs of the indicators are integrated of order 1;
- Let $Y_{1,t}$ be the geometric mean of $Y_{1,t}^*$, $Y_{1,t-1}^*$, and $Y_{1,t-2}^*$, i.e.,

$$\ln Y_{1,t} = \frac{1}{3}(\ln Y_{1,t}^* + \ln Y_{1,t-1}^* + \ln Y_{1,t-2}^*)$$

where $\{Y_{1,t}^*\}_{t=-\infty}^{\infty}$: an N_1 -variate latent random sequence such that for all t ,

- Taking the three-period differences, for all t ,

$$\begin{aligned}
y_{1,t} &= \ln Y_{1,t} - \ln Y_{1,t-3} \\
&= \frac{1}{3}(\ln Y_{1,t}^* - \ln Y_{1,t-3}^*) + \frac{1}{3}(\ln Y_{1,t-1}^* - \ln Y_{1,t-4}^*) \\
&\quad + \frac{1}{3}(\ln Y_{1,t-2}^* - \ln Y_{1,t-5}^*) \\
&= \frac{1}{3}(y_{1,t}^* + y_{1,t-1}^* + y_{1,t-2}^*) + \frac{1}{3}(y_{1,t-1}^* + y_{1,t-2}^* + y_{1,t-3}^*) \\
&\quad + \frac{1}{3}(y_{1,t-2}^* + y_{1,t-3}^* + y_{1,t-4}^*) \\
&= \frac{1}{3}y_{1,t}^* + \frac{2}{3}y_{1,t-1}^* + y_{1,t-2}^* + \frac{2}{3}y_{1,t-3}^* + \frac{1}{3}y_{1,t-4}^*, \tag{93}
\end{aligned}$$

where $y_{1,t} = \Delta_3 \ln Y_{1,t}$ and $y_{1,t}^* = \Delta \ln Y_{1,t}^*$.

- $y_{1,t}$ is observed every third period, but $y_{1,t}^*$ is unobserved.

Now build the dynamic factor model for $y_{1,t}^*$ and $y_{2,t}$, which are in the same frequency. The model is given by

$$\begin{bmatrix} y_{1,t}^* \\ y_{2,t} \end{bmatrix} = \begin{bmatrix} \mu_1^* \\ \mu_2 \end{bmatrix} + \beta f_t + u_t, \quad (94)$$

$$\phi_f(L)f_t = v_{1,t}, \quad (95)$$

$$\Phi_u(L)u_t = v_{2,t}, \quad (96)$$

$$\begin{bmatrix} v_{1,t} \\ v_{2,t} \end{bmatrix} = NID \left(0, \begin{bmatrix} \sigma_1^2 & 0 \\ 0 & \Sigma_{22} \end{bmatrix} \right). \quad (97)$$

For identification, assume that: (i) the first component of β is 1; (ii) $\Phi_u(\cdot)$ and Σ_{22} is diagonal.

Since we never observe $y_{1,t}^*$, we cannot estimate the above model directly.

Hence we consider the corresponding dynamic one-factor model for $\{y_t\}_{t=-\infty}^{\infty}$ such that for all t ,

$$\begin{bmatrix} y_{1,t} \\ y_{2,t} \end{bmatrix} = \begin{bmatrix} \mu_1 \\ \mu_2 \end{bmatrix} + \begin{bmatrix} \beta_1 \left(\frac{1}{3}f_t + \frac{2}{3}f_{t-1} + f_{t-2} + \frac{2}{3}f_{t-3} + \frac{1}{3}f_{t-4} \right) \\ \beta_2 f_t \end{bmatrix} \\ + \begin{bmatrix} \frac{1}{3}u_{1,t} + \frac{2}{3}u_{1,t-1} + u_{1,t-2} + \frac{2}{3}u_{1,t-3} + \frac{1}{3}u_{1,t-4} \\ u_{2,t} \end{bmatrix},$$

where $\mu_1 = 3\mu_1^*$, $(\beta'_1, \beta'_2)' = \beta$, and $(u'_{1,t}, u'_{2,t})' = u_t$.

4.2 Mixed-frequency Markov switching dynamic factor model

By including regime switching dynamics in the intercept, we have

$$\begin{bmatrix} y_{1,t} \\ y_{2,t} \end{bmatrix} = \begin{bmatrix} \mu_1 \\ \mu_2 \end{bmatrix} + \begin{bmatrix} \beta_1 \left(\frac{1}{3}f_t + \frac{2}{3}f_{t-1} + f_{t-2} + \frac{2}{3}f_{t-3} + \frac{1}{3}f_{t-4} \right) \\ \beta_2 f_t \end{bmatrix} \\ + \begin{bmatrix} \frac{1}{3}u_{1,t} + \frac{2}{3}u_{1,t-1} + u_{1,t-2} + \frac{2}{3}u_{1,t-3} + \frac{1}{3}u_{1,t-4} \\ u_{2,t} \end{bmatrix},$$

$$\phi_f(L)f_t = \nu(S_t) + v_{1,t},$$

$$\Phi_u(L)u_t = v_{2,t},$$

$$S_t|S_{t-1} \sim \text{Markov}(P),$$

$$\begin{bmatrix} v_{1,t} \\ v_{2,t} \end{bmatrix} = NID \left(0, \begin{bmatrix} \sigma_1^2 & 0 \\ 0 & \Sigma_{22} \end{bmatrix} \right).$$

Example 3. Using the data in Mariano and Murasawa (2003).

```

Strong convergence using analytical derivatives
Log-likelihood value = -1273.59
=====
Results of Parameter Estimations for Models
-----

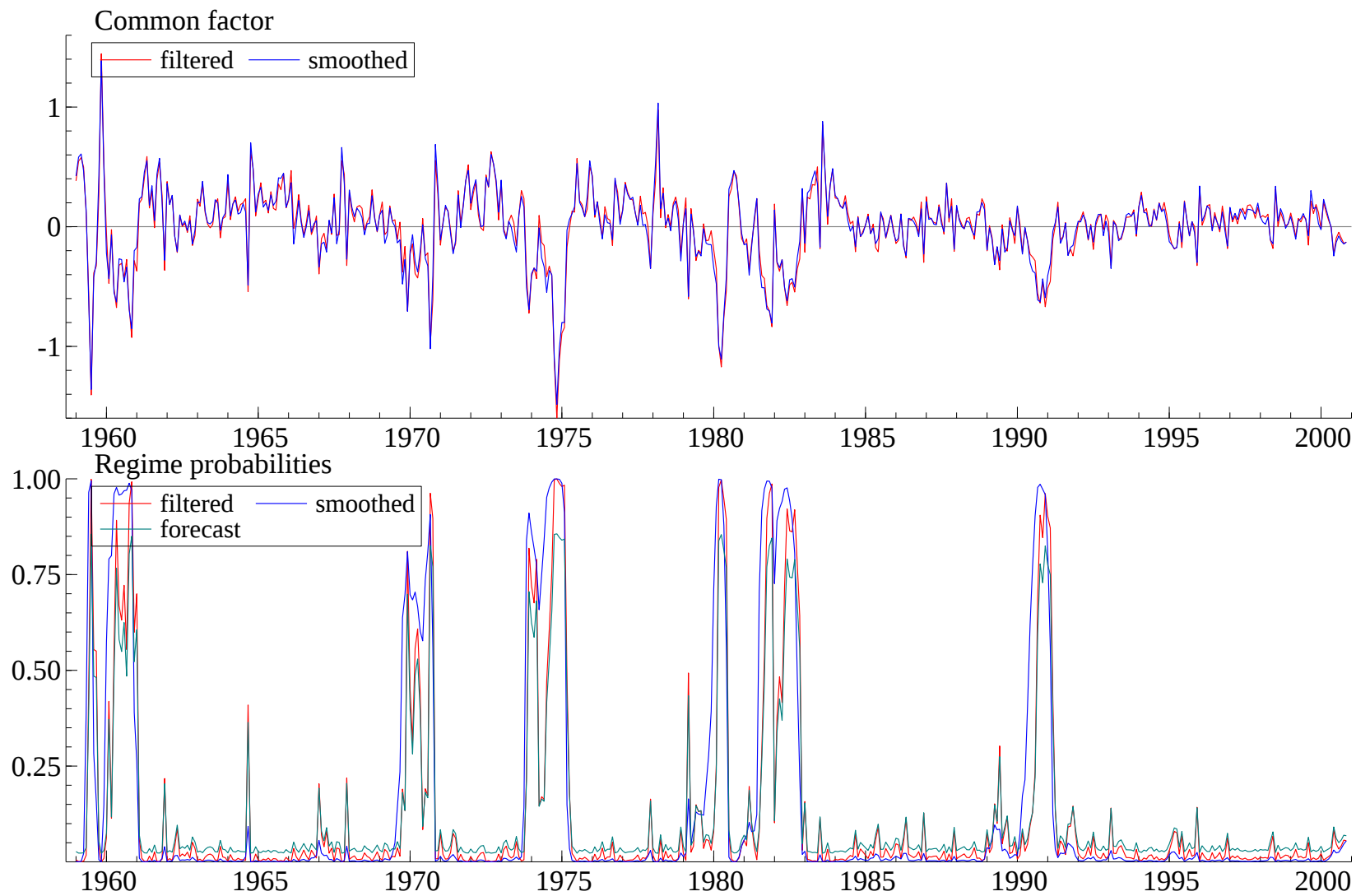
```

	Coef.Est	Std.Err	t-value	t-prob
p11	0.856730	0.068624	12.484343	0.000000
p22	0.977691	0.009061	107.904145	0.000000
beta2	0.508209	0.039172	12.973795	0.000000
beta3	0.817620	0.062672	13.046035	0.000000
beta4	2.154900	0.136030	15.841397	0.000000
beta5	1.755300	0.120708	14.541699	0.000000
nu1	-0.361690	0.066312	-5.454362	0.000000
nu2	0.057276	0.015802	3.624598	0.000320
phi_f1	0.334640	0.060548	5.526812	0.000000
sig_f	0.244907	0.018122	13.514373	0.000000
phi_u_11	-0.821450	0.199619	-4.115092	0.000046
phi_u_12	0.034668	0.253728	0.136634	0.891378
phi_u_21	0.094243	0.045042	2.092351	0.036934
phi_u_22	0.448270	0.050138	8.940794	0.000000
phi_u_31	-0.055209	0.051342	-1.075326	0.282771
phi_u_32	0.032623	0.051215	0.636987	0.524439
phi_u_41	-0.025256	0.064229	-0.393216	0.694335
phi_u_42	-0.058011	0.060438	-0.959844	0.337619
phi_u_51	-0.406940	0.049169	-8.276412	0.000000
phi_u_52	-0.192700	0.047797	-4.031632	0.000064
sig_u_1	0.494360	0.091277	5.416042	0.000000
sig_u_2	0.130292	0.006423	20.285157	0.000000
sig_u_3	0.296111	0.010835	27.328795	0.000000
sig_u_4	0.505410	0.023909	21.138977	0.000000
sig_u_5	0.783600	0.027914	28.071702	0.000000

```

=====

```



Example 4. Large dynamic factor models, forecasting, and nowcasting

See <https://gist.github.com/ChadFulton/a9172cd6f41947333f47>

5 Nonlinear and non-Gaussian state space models

- **Reference:**

- Cappé, O., Godsill, S. J., and Moulines, E., (2007), An overview of existing methods and recent advances in sequential Monte Carlo, IEEE Proc., 95, 899 – 924.

- **Outline:**

- Nonlinear and non-Gaussian state space models
- Importance sampling and resampling
- Sequential Monte Carlo filters

5.1 The Nonlinear and Non-Gaussian SS Models

A statistical model generates the observed data, a vector y_t , and the conditional distribution of y_t depends on a latent state variable, x_t . See the Figure.

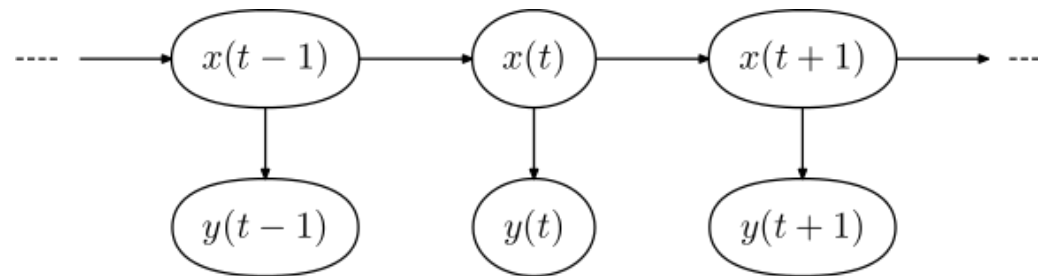


Figure 2 The State-Space (The Hidden states given by the vector x , and the observation states given by the vector y).

Formally, the data is generated by the following state space model

$$\text{Observation equation:} \quad y_t = f(x_t, \varepsilon_t^y) \quad \Rightarrow \quad p(y_t|x_t) \quad (98)$$

$$\text{State equation:} \quad x_t = g(x_{t-1}, \varepsilon_t^x) \quad \Rightarrow \quad p(x_t|x_{t-1}), \quad (99)$$

where ε_t^y is the observation error or “noise”, and ε_t^x is the state shock.

The observation equation is often written as a conditional likelihood, $p(y_t|x_t)$, and the state evolution as $p(x_t|x_{t-1})$.

For simplicity, we ignore the static parameters θ in these distributions.

Our interest is mainly in $p(x_t|y_{0:t})$, where $y_{0:t} = (y_0, y_1, \dots, y_t)$.

Beginning with the Kalman filter, computing $p(x_t|y_{0:t})$ uses a two-step procedure: prediction and Bayesian updating.

- Prediction

$$p(x_t|y_{0:t-1}) = \int p(x_t|x_{t-1})p(x_{t-1}|y_{0:t-1})dx_{t-1},$$

- Updating by Bayes rule

$$\underbrace{p(x_t|y_{0:t})}_{\text{Posterior}} = \frac{p(x_t, y_t|y_{0:t-1})}{p(y_t|y_{0:t-1})} = \frac{p(y_t|x_t)p(x_t|y_{0:t-1})}{\int p(y_t|x_t)p(x_t|y_{0:t-1})dx_t}$$
$$\propto \underbrace{p(y_t|x_t)}_{\text{Likelihood}} \underbrace{p(x_t|y_{0:t-1})}_{\text{Prior}}.$$

We now need the computation of the two integrals.

- Only for a limited number of settings such as linear Gaussian model, $p(x_t|y_{0:t-1})$ and $p(x_t|y_{0:t})$ are Gaussian. In this case, the Kalman filter recursion is often used to get the mean and variance of $p(x_t|y_{0:t})$.
- However, in nonlinear and non-Gaussian models, it is not possible to analytically compute $p(x_t|y_{0:t})$ since it is a complicated function of $y_{0:t}$. Simulation methods are typically required to characterize $p(x_t|y_{0:t})$.

5.2 Importance Sampling and Resampling

In the Monte Carlo method, we are concerned with estimating the properties of some highly complex probability distribution p , for example computing expectations of the form:

$$\bar{h} = \int h(x)p(x)dx,$$

where $h(\cdot)$ is some useful function for estimation, for example $h(x) = x$.

Sample from p

When the integral cannot be achieved analytically, we can generate random samples (called as *particles*) from p , denote these $\{x^{(i)}\}_{1 \leq i \leq N}$, and then approximate the distribution p by point masses so that

$$\bar{h} \approx \frac{1}{N} \sum_{i=1}^N h(x^{(i)}).$$

Clearly N needs to be large in order to give a good coverage of all regions of interest.

Sample from q

When we cannot sample directly from the distribution p , we can sample from another distribution q (the importance distribution, or instrumental distribution). This is also called as the importance sampling.

Make N random draws (*particles*) $x^{(i)}$, $i = 1, \dots, N$ from q instead of p . Consequently, $\bar{h}$ can be estimated using a weighted average:

$$\begin{aligned}\bar{h} &= \int h(x) \frac{p(x)}{q(x)} q(x) dx = \int h(x) r(x) q(x) dx \\ &= \sum_{j=1}^N \frac{\tilde{w}^{(j)}}{\sum_{j=1}^N \tilde{w}^{(j)}} h(x^{(j)}),\end{aligned}$$

where $\tilde{w}^{(j)} = r(x^{(j)}) = p(x^{(j)})/q(x^{(j)})$ is termed the unnormalised importance weight.

Definition: A sample $(x^{(j)}, \tilde{w}^{(j)})$, $j = 1, \dots, N$ is said to be properly weighted with respect to distribution p if for all integrable function h , we have

$$\sum_{j=1}^N \frac{\tilde{w}^{(j)}}{\sum_{i=1}^N \tilde{w}^{(i)}} h(x^{(j)}) \rightarrow E_p[h(X)] \quad \text{as } N \rightarrow \infty.$$

Sampling importance resampling (SIR)

In the *first stage of importance sampling*, an i.i.d. sample $(\tilde{x}^{(1)}, \dots, \tilde{x}^{(M)})$ is drawn from the importance distribution q , and one computes the normalised version of the importance weights,

$$w^{(i)} = \frac{\tilde{w}^{(i)}}{\sum_{i=1}^M \tilde{w}^{(i)}}, \quad i = 1, \dots, M. \quad (100)$$

In the *resampling stage*, a sample of size N denoted by $x^{(1)}, \dots, x^{(N)}$ is drawn from the intermediate set of points $\tilde{x}^{(1)}, \dots, \tilde{x}^{(M)}$, taking proper account of the weights computed in (100).

This principle is illustrated in Figure 3.

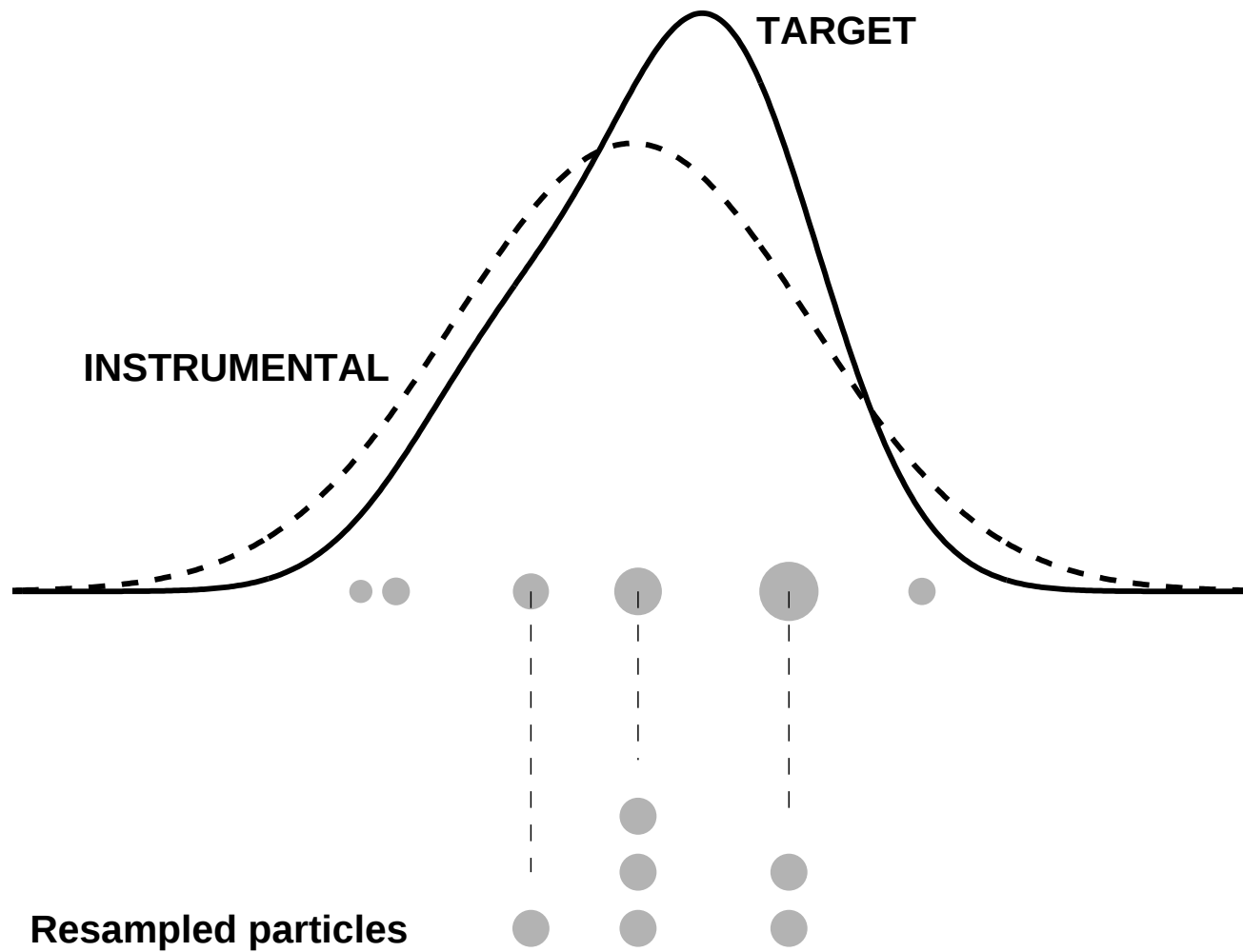


Figure 3 Principle of resampling. ($M = N = 6$)

- The particle approximation can be transformed into an equally weighted random sample from $p(x)$ by resampling from the discrete distribution, $\{\tilde{x}^{(j)}, \tilde{w}^{(j)}\}_{j=1}^N$.
- This procedure produces a new sample with uniformly distributed weights, $w_t^{(j)} = 1/N$.
- Resampling can be done in many ways, but the simplest is multinomial sampling. Other methods include stratified sampling, residual resampling, and systematic resampling.

5.3 Sequential Monte Carlo methods

Starting with the initial, or ‘prior’, distribution $p(x_0)$, *the (on-line) posterior density* $p(x_{0:t}|y_{0:t})$ can be obtained *recursively*:

- Prediction

$$p(x_{0:t}|y_{0:t-1}) = p(x_{0:t-1}|y_{0:t-1})p(x_t|x_{t-1}) \quad (101)$$

- Correction (or updating)

$$\begin{aligned} p(x_{0:t}|y_{0:t}) &= \frac{p(y_t|x_t)p(x_{0:t}|y_{0:t-1})}{\ell(y_t|y_{0:t-1})} \\ &= \frac{p(y_t|x_t)p(x_t|x_{t-1})}{\ell(y_t|y_{0:t-1})}p(x_{0:t-1}|y_{0:t-1}). \end{aligned} \quad (102)$$

We would like to sample from $p(x_{0:t}|y_{0:t})$; since it is generally impossible to sample directly from this distribution, we resort to a sequential version of the importance sampling and resampling procedure.

Conceptually, we $\tilde{x}_{0:t}^{(i)}$, $i = 1, \dots, N$, from a convenient importance distribution $q(x_{0:t}|y_{0:t})$, and compute the unnormalised importance weights

$$\tilde{w}_t^{(i)} = \frac{p(\tilde{x}_{0:t}^{(i)}|y_{0:t})}{q(\tilde{x}_{0:t}^{(i)}|y_{0:t})}, \quad i = 1, \dots, N.$$

We can get the normalised weights with $w_t^{(i)} = \frac{\tilde{w}_t^{(i)}}{\sum_{i=1}^N \tilde{w}_t^{(i)}}$. Consequently,

$$\hat{p}(x_{0:t}|y_{0:t}) = \sum_{i=1}^N w_t^{(i)} \delta_{\tilde{x}_{0:t}^{(i)}}(x_{0:t}), \quad (103)$$

where $\delta(\cdot)$ is a dirac function.

To achieve the importance distribution we construct the proposal such that it factorises in a form similar to that of the target posterior distribution:

$$q(x_{0:t}|y_{0:t}) = \overbrace{q(x_{0:t-1}|y_{0:t-1})}^{\text{Keep existing path}} \times \overbrace{q(x_t|x_{t-1}, y_t)}^{\text{extend path}}.$$

The unnormalised importance weights then take the following form

$$\tilde{w}_t^{(i)} = \frac{p(\tilde{x}_{0:t}^{(i)}|y_{0:t})}{q(\tilde{x}_{0:t}^{(i)}|y_{0:t})} \propto w_{t-1}^{(i)} \times \frac{p(y_t|\tilde{x}_t^{(i)})p(\tilde{x}_t^{(i)}|\tilde{x}_{t-1}^{(i)})}{q(\tilde{x}_t^{(i)}|\tilde{x}_{t-1}^{(i)}, y_t)}. \quad (104)$$

In the sequential Monte Carlo literature, the multiplicative update factor on the right-hand side of (104) is often called the *incremental weight*.

(1) SIS/SIR algorithm

The sequential importance sampling (SIS) algorithm is summarized as follows:

Algorithm 3 (*Sequential Importance Sampling (SIS)*).

1. At initial time $t = 0$,

(a) Sample $\tilde{x}_0^{(i)} \sim q(x_0)$, for $i = 1, \dots, N$.

(b) Compute the importance weights $w_0^{(i)} = \tilde{w}_0^{(i)} / \sum_{i=1}^N \tilde{w}_0^{(i)}$, where

$$\tilde{w}_0^{(i)} = \frac{p(y_0 | \tilde{x}_0^{(i)}) p(\tilde{x}_0^{(i)})}{q(\tilde{x}_0^{(i)})}, \quad \text{for } i = 1, \dots, N.$$

2. For $t = 1, \dots, T$,

(a) For $i = 1, \dots, N$, propagate particles: $\tilde{x}_t^{(i)} \sim q(x_t | \tilde{x}_{t-1}^{(i)}, y_t)$,

compute weights:

$$\tilde{w}_t^{(i)} = w_{t-1}^{(i)} \frac{p(y_t | \tilde{x}_t^{(i)}) p(\tilde{x}_t^{(i)} | \tilde{x}_{t-1}^{(i)})}{q(\tilde{x}_t^{(i)} | \tilde{x}_{t-1}^{(i)}, y_t)}.$$

(b) Normalise weights: $w_t^{(i)} = \tilde{w}_t^{(i)} / \sum_{i=1}^N \tilde{w}_t^{(i)}, \quad i = 1, \dots, N,$
compute filtering estimate (on-line filtered estimate): $\bar{h}_t = \sum_{i=1}^N w_t^{(i)} h_t(\tilde{x}_t^{(i)}).$

The problem in SIS is that the algorithm is degenerate. It is known that the variance of the weights increases at every step. This means that we will always converge to a single non-zero weight $w^{(i)} = 1$ for some i and the rest being zero.

Resampling strategy is commonly used to avoid the problem of degeneracy of the algorithm.

The sampling importance resampling (SIR) algorithm can be summarized as follows:

Algorithm 4 (*Sampling Importance Resampling (SIR)*).

1. At initial time $t = 0$.

(a) Sample $\tilde{x}_0^{(i)} \sim q(x_0)$, for $i = 1, \dots, N$.

(b) Compute importance weights $w_0^{(i)} = \tilde{w}_0^{(i)} / \sum_{i=1}^N \tilde{w}_0^{(i)}$, where

$$\tilde{w}_0^{(i)} = \frac{p(y_0 | \tilde{x}_0^{(i)}) p(\tilde{x}_0^{(i)})}{q(\tilde{x}_0^{(i)})}, \quad \text{for } i = 1, \dots, N.$$

(c) Resample N particles $\{\tilde{x}_0^{(i)}\}_{i=1}^N$ with weights $\{\tilde{w}_0^{(i)}\}_{i=1}^N$ to obtain the new particles $\{x_0^{(i)}, 1/N\}_{i=1}^N$.

2. For $t = 1, \dots, T$,

(a) For $i = 1, \dots, N$, propagate particles: $\tilde{x}_t^{(i)} \sim q(x_t | x_{t-1}^{(i)}, y_t)$, compute weights:

$$\tilde{w}_t^{(i)} = w_{t-1}^{(i)} \frac{p(y_t | \tilde{x}_t^{(i)}) p(\tilde{x}_t^{(i)} | x_{t-1}^{(i)})}{q(\tilde{x}_t^{(i)} | x_{t-1}^{(i)}, y_t)}.$$

(b) Normalise weights: $w_t^{(i)} = \tilde{w}_t^{(i)} / \sum_{i=1}^N \tilde{w}_t^{(i)}$, $i = 1, \dots, N$, compute filtering estimate (on-line filtered estimate): $\bar{h}_t = \sum_{i=1}^N w_t^{(i)} h_t(\tilde{x}_t^{(i)})$.

(c) Resample N particles $\{\tilde{x}_t^{(i)}\}_{i=1}^N$ with weights $\{\tilde{w}_t^{(i)}\}_{i=1}^N$ to obtain the new particles $\{x_t^{(i)}, 1/N\}_{i=1}^N$.

Effective sample size (ESS)

An important technique for resampling and rejuvenating a sequential importance sampler is the *effective sample size* (ESS), which is a useful metric to determine whether resampling is necessary.

Define the coefficients of variation of the importance weights as $cv_t^2 = \frac{\text{Var}(w_t^{(i)})}{E^2(w_t^{(i)})}$, Liu and Chen (1995) and Liu (1996) propose tracking the ESS

$$N_{\text{eff},t} = \frac{N}{1 + cv_t^2},$$

which varies between 1 (N copies of a single particle) and N (equal weights).

If the ESS is large, say, $N_{\text{eff},t} > \alpha_0 N$ ($0 < \alpha_0 < 1$), resampling is undesirable since the pdf over the trajectories is well represented.

Replace Step 2(c) in SIR with the following step

(c) If resampling when $N_{eff,t} \leq \alpha_0 N$, select N particle indices $j_i \in \{1, \dots, N\}$ according to weights $\{w_t^{(i)}\}_{1 \leq i \leq N}$, and set $x_t^{(i)} = \tilde{x}_t^{(j_i)}$ and $w_t^{(i)} = 1/N$, $i = 1, \dots, N$; otherwise, set $x_t^{(i)} = \tilde{x}_t^{(i)}$, $i = 1, \dots, N$.

(2) The bootstrap filter

The bootstrap filter is proposed by Gordon, Salmond, and Smith (1993)

- It uses the state transition density or “prior kernel” as importance distribution, i.e., $q(x_t|x_{t-1}, y_t) = p(x_t|x_{t-1})$.
- The importance weight then simplifies to

$$w_t^{(i)} \propto w_{t-1}^{(i)} p(y_t|x_t^{(i)}).$$

- In the original version of the algorithm, resampling is carried out at each and every time step, in which case the term $w_{t-1}^{(i)} = 1/N$ is a constant, which may thus be ignored.

The bootstrap filter can be summarized as follows:

Algorithm 5 (*The bootstrap filter*).

1. At initial time $t = 0$.

(a) Sample $x_0^{(i)} = \tilde{x}_0^{(i)} \sim p(x_0)$ with associated weights $w_0^{(i)} = 1/N$.

2. For $t = 1, \dots, T$,

(a) For $i = 1, \dots, N$, propagate particles: $\tilde{x}_t^{(i)} \sim p(x_t | x_{t-1}^{(i)})$, compute weights:

$$\tilde{w}_t^{(i)} = p(y_t | \tilde{x}_t^{(i)}).$$

(b) Normalise weights: $w_t^{(i)} = \tilde{w}_t^{(i)} / \sum_{i=1}^N \tilde{w}_t^{(i)}$, $i = 1, \dots, N$, compute filtering estimate (on-line filtered estimate): $\bar{h}_t = \sum_{i=1}^N w_t^{(i)} h_t(\tilde{x}_t^{(i)})$.

(c) Resample according to weights $\{w_t^{(i)}\}_{1 \leq i \leq N}$, and set $x_t^{(i)} = \tilde{x}_t^{(j_i)}$ and $w_t^{(i)} = 1/N$, $i = 1, \dots, N$.

(3) Auxiliary sampling

The auxiliary particle filter (APF) of Pitt and Shephard (1999, 2001) is a popular algorithm that is simple to implement and works well in many cases. They approximate the incremental target distribution in (102) with the importance distribution

$$\begin{aligned} p(y_t|x_t)p(x_t|x_{t-1}) &\approx g_1(y_t|x_t)g_2(x_t|x_{t-1}) \\ &= g_1(y_t|x_{t-1})g_2(x_t|x_{t-1}, y_t) \end{aligned} \quad (105)$$

The proposal distribution in (105) is decomposed into two parts implying that the sampling of new values $\{x_t^{(i)}\}_{i=1}^N$ from this distribution can potentially be performed in two steps.

The APF is given as follows:

Algorithm 6 (*Auxiliary Particle Filter (APF)*).

1. At initial time $t = 0$.

(a) Sample $\tilde{x}_0^{(i)} \sim g_2(x_0)$, for $i = 1, \dots, N$.

(b) Compute importance weights $w_0^{(i)} = \tilde{w}_0^{(i)} / \sum_{i=1}^N \tilde{w}_0^{(i)}$, where

$$\tilde{w}_0^{(i)} = \frac{p(y_0|\tilde{x}_0^{(i)})p(\tilde{x}_0^{(i)})}{g_2(\tilde{x}_0^{(i)})}, \quad \text{for } i = 1, \dots, N.$$

(c) Resample if necessary.

2. For $t = 1, \dots, T$,

(a) For $i = 1, \dots, N$, compute $\tau_{t-1}^{(i)} = g_1(y_t|\tilde{x}_{t-1}^{(i)})$ and normalize $v_{t-1}^{(i)} =$

$\frac{\tau_{t-1}^{(i)}}{\sum_{j=1}^N \tau_{t-1}^{(j)}} \cdot$
(b) Select N particle indices $j_i \in \{1, \dots, N\}$ according to weights $\{v_{t-1}^{(i)}\}_{i=1}^N$. For $i = 1, \dots, N$, set $x_{t-1}^{(i)} = \tilde{x}_{t-1}^{(j_i)}$ and set first stage weights

$$u_{t-1}^{(i)} = \frac{w_{t-1}^{(j_i)}}{v_{t-1}^{(j_i)}}.$$

(c) For $i = 1, \dots, N$, propagate particles: $\tilde{x}_t^{(i)} \sim g_2(x_t | x_{t-1}^{(i)}, y_t)$, compute weights:

$$\tilde{w}_t^{(i)} = u_{t-1}^{(i)} \frac{p(y_t | \tilde{x}_t^{(i)}) p(\tilde{x}_t^{(i)} | x_{t-1}^{(i)})}{g_2(\tilde{x}_t^{(i)} | x_{t-1}^{(i)}, y_t)} = \frac{p(y_t | \tilde{x}_t^{(i)}) p(\tilde{x}_t^{(i)} | x_{t-1}^{(i)})}{g_1(y_t | \tilde{x}_{t-1}^{(j_i)}) g_2(\tilde{x}_t^{(i)} | x_{t-1}^{(i)}, y_t)}.$$

(d) Normalise weights: $w_t^{(i)} = \tilde{w}_t^{(i)} / \sum_{i=1}^N \tilde{w}_t^{(i)}$, $i = 1, \dots, N$, compute filtering estimate (on-line filtered estimate): $\bar{h}_t = \sum_{i=1}^N w_t^{(i)} h_t(\tilde{x}_t^{(i)})$.

6 Score-Driven Models

6.1 Basic GAS model

Let $\{y_t\}$ be a time series and $\mathcal{F}_t = \{y_t, y_{t-1}, \dots\}$ be available information up to t . Creal et al. (2013) proposed the framework of Generalized Autoregressive Score Models (GAS) under a general conditional distribution $p(y_t \mid \mathcal{F}_{t-1})$, which can be parameterized as

$$y_t \mid \mathcal{F}_{t-1} \sim p(y_t \mid f_t, \theta), \quad (106)$$

- f_t is the time-varying parameter of interest which fully characterizes $p(\cdot)$ and only depends on $\mathcal{F}_{t-1}$;
- θ is a collection of static additional parameters.

Furthermore, the time-varying parameter f_t is assumed to follow the familiar autoregressive updating equation

$$f_{t+1} = \omega + \sum_{j=1}^p A_j f_{t-j+1} + \sum_{j=1}^q B_j s_{t-j+1}, \quad (107)$$

where ω is a vector of constants, A_i and B_j are coefficient matrices. In (107), s_t is a vector which is proportional to the score of (106):

$$s_t = S_t \nabla_t \quad (108)$$

where S_t is a positive definite scaling matrix known at time t and

$$\nabla_t = \frac{\partial \log p(y_t \mid f_t, \theta)}{\partial f_t} \quad (109)$$

is the score of (107) evaluated at f_t .

Creal et al. (2013) suggest to set the scaling matrix S_t to a power $\gamma > 0$ of the inverse of the information matrix of f_t to account for the variance of ∇_t . More precisely:

$$S_t = I_t^{-\gamma} \quad \text{and} \quad I_t = E_{t-1} [\nabla_t \nabla_t'] \quad (110)$$

where the expectation is taken with respect to the conditional distribution of y_t given $\mathcal{F}_{t-1}$.

- The additional parameter γ is fixed by the user and usually takes values 0, 1/2 or 1.
- When $\gamma = 0$, $S_t = I$ and there is no scaling.
- If $\gamma = 1$ ($\gamma = 1/2$), the conditional score ∇_t is premultiplied by the inverse of (the square root of) its covariance matrix I_t .

6.2 Special cases of GAS models

(1) Regression model:

$$y_t = x_t' \beta_{t-1} + \varepsilon_t, \quad \varepsilon_t \sim N(0, \sigma^2),$$

where x_t is a $k \times 1$ vector x_t of exogenous variables, β_{t-1} is a $k \times 1$ vector of time-varying regression coefficients.

Let $f_t = \beta_t$. The scaled score function based on $S_t = I_t^{-1}$ is given by

$$s_t = (x_t' x_t)^{-1} x_t' (y_t - x_t' f_{t-1}) = (x_t' x_t)^{-1} x_t' \varepsilon_t.$$

The GAS(1,1) specification for the time-varying regression coefficient is

$$f_t = \omega + A_0 s_t + B_1 f_{t-1} \tag{111}$$

In case $x_t \equiv 1$, the above equation for the time-varying intercept reduces to the EWMA recursion by setting $\omega = 0$ and $B_1 = 1$, i.e.

$$f_t = f_{t-1} + A_0(y_t - f_{t-1}).$$

In this case, we obtain the observation driven analogue of the local level (parameter driven) model,

$$y_t = \mu_{t-1} + \varepsilon_t, \quad \mu_t = \mu_{t-1} + \eta_t,$$

see Durbin and Koopman (2001, Chapter 2).

(2) Unobserved component models

Consider a two component decomposition of a time series $\{y_t\}$. Let $f_t = (f_{1t}, f_{2t})$ be the vector of decomposed components.

For this decomposition we obtain the GAS(1,2) model with observation model

$$y_t = f_{1,t-1} + f_{2,t-1} + \varepsilon_t, \quad \varepsilon_t \sim N(0, \sigma^2),$$

The updating equation is specified as

$$f_t = \begin{bmatrix} \omega \\ 0 \end{bmatrix} + \begin{bmatrix} \alpha_1 \\ \alpha_2 \end{bmatrix} s_t + \begin{bmatrix} 1 & 0 \\ 0 & \phi_1 \end{bmatrix} f_{t-1} + \begin{bmatrix} 0 & 0 \\ 0 & \phi_2 \end{bmatrix} f_{t-2}.$$

The scaled score function is given by

$$s_t = y_t - f_{1,t-1} - f_{2,t-1} = \varepsilon_t$$

This GAS decomposition can be regarded as the observation driven equivalent of the UC models of Watson (1986) and Clark (1989), who also aim to decompose macroeconomic time series into trend and cycle factors.

The UC trend-cycle decomposition model is then given by

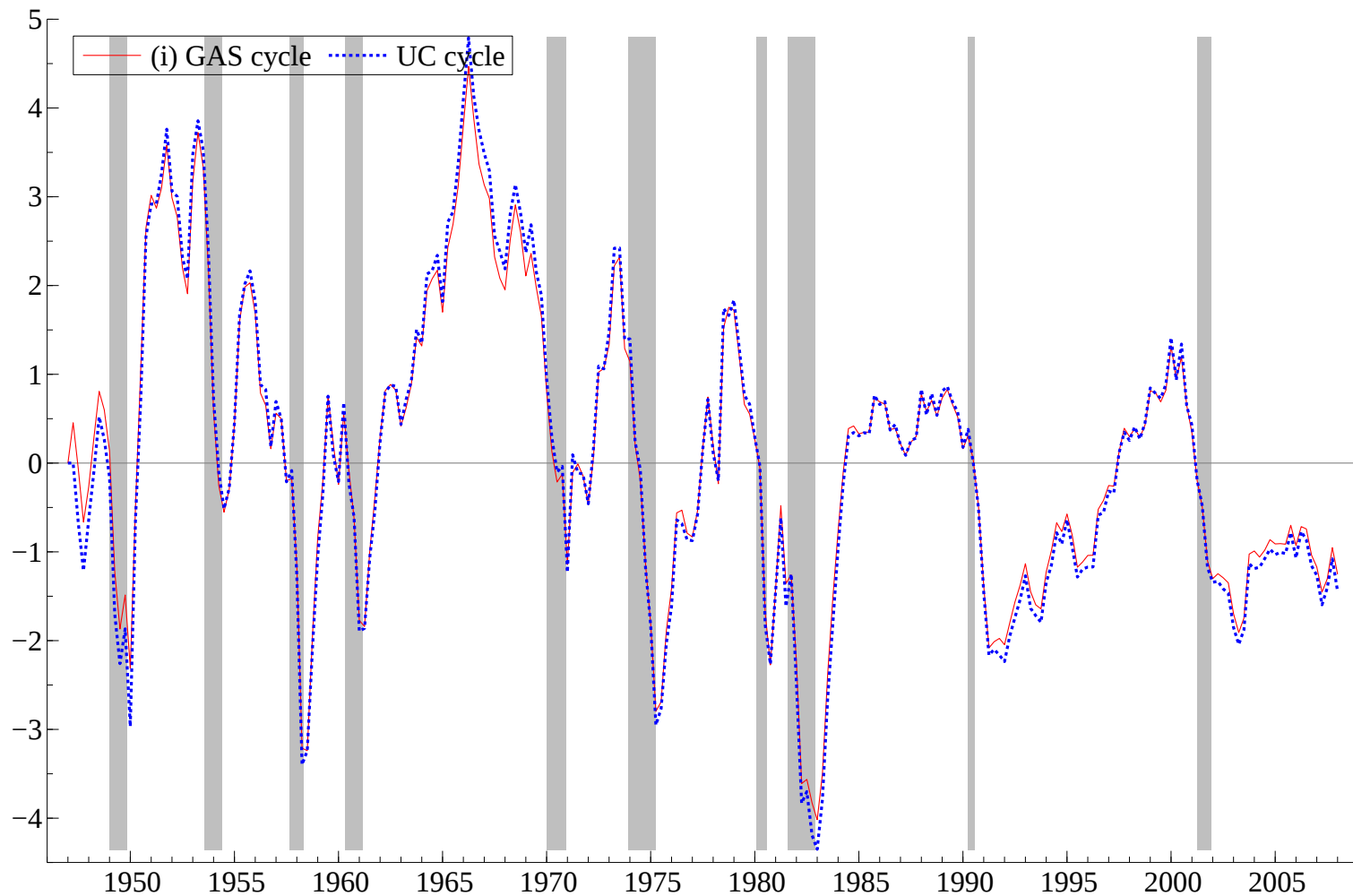
$$y_t = f_{1,t} + f_{2,t}$$

$$f_{1t} = \omega + f_{1,t-1} + a_1\xi_{1,t}, \quad \xi_{1,t} \sim N(0, 1),$$

$$f_{2t} = \phi_1 f_{2,t-1} + \phi_2 f_{2,t-2} + a_2\xi_{2,t}, \quad \xi_{2,t} \sim N(0, 1),$$

where the disturbances $\xi_{1,t}$ and $\xi_{2,t}$ are mutually and serially independent.

Example 5. Trend-cycle illustration: estimated cycles from the GAS and UC trend-cycle models based on quarterly log of U.S. real gdp from 1947(1) through 2008(1).



More GAS models with Python examples are released in

<https://pyflux.readthedocs.io/en/latest/gas.html>.

 PyFlux

latest

Search docs

Getting Started with Time Series

ARIMA models

ARIMAX models

DAR models

Dynamic Linear regression models

Beta-t-EGARCH models

Beta-t-EGARCH in-mean models

Beta-t-EGARCH in-mean regression models

Beta-t-EGARCH long memory models

Beta Skew-t GARCH models

Beta Skew-t in-mean GARCH models

GARCH models

⊖ GAS models

Introduction

Example

Class Description

References

GAS local level models

GAS local linear trend models

Docs » GAS models

 Edit on GitHub

GAS models

Introduction

Generalized Autoregressive Score (GAS) models are a recent class of observation-driven time-series model for non-normal data, introduced by Creal et al. (2013) and Harvey (2013). For a conditional observation density $p(y_t | x_t)$ with an observation y_t and a latent time-varying parameter x_t , we assume the parameter x_t follows the recursion:

$$x_t = \mu + \sum_{i=1}^p \phi_i x_{t-i} + \sum_{j=1}^q \alpha_j S(x_{t-j}) \frac{\partial \log p(y_{t-j} | x_{t-j})}{\partial x_{t-j}}$$

For example, for a Poisson distribution density, where the default scaling is $\exp(x_j)$, the time-varying parameter follows:

$$x_t = \mu + \sum_{i=1}^p \phi_i x_{t-i} + \sum_{j=1}^q \alpha_j \left(\frac{y_{t-j}}{\exp(x_{t-j})} - 1 \right)$$

These types of model can be viewed as approximations to parameter-driven state space models, and are often competitive in predictive performance. See **GAS State Space models** for a more general class of models that extend beyond the simple autoregressive form. The simple GAS models considered here in this notebook can be viewed as an approximation to non-linear ARIMA processes.

Chapter Six

State-Space Models © 2026 Tingguo Zheng

135