

Schlicht, Ekkehart; Ludsteck, Johannes

Working Paper

Variance Estimation in a Random Coefficients Model

Munich Discussion Paper, No. 2006-12

Provided in Cooperation with:

University of Munich, Department of Economics

Suggested Citation: Schlicht, Ekkehart; Ludsteck, Johannes (2006) : Variance Estimation in a Random Coefficients Model, Munich Discussion Paper, No. 2006-12, Ludwig-Maximilians-Universität München, Volkswirtschaftliche Fakultät, München, <https://doi.org/10.5282/ubm/epub.904>

This Version is available at:

<https://hdl.handle.net/10419/104208>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.



Ekkehart Schlicht und Johannes Ludsteck: Variance Estimation in a Random Coefficients Model

Munich Discussion Paper No. 2006-12

Department of Economics
University of Munich

Volkswirtschaftliche Fakultät
Ludwig-Maximilians-Universität München

Online at <http://epub.ub.uni-muenchen.de/904/>

Variance Estimation in a Random Coefficients Model

by

EKKEHART SCHLICHT and JOHANNES LUDSTECK¹

Abstract

This paper describes an estimator for a standard state-space model with coefficients generated by a random walk that is statistically superior to the Kalman filter as applied to this particular class of models. Two closely related estimators for the variances are introduced: A maximum likelihood estimator and a moments estimator that builds on the idea that some moments are equalized to their expectations. These estimators perform quite similar in many cases. In some cases, however, the moments estimator is preferable both to the proposed likelihood estimator and the Kalman filter, as implemented in the program package Eviews.

¹ Ekkehart Schlicht (schlicht@lmu.de) is at the Department of Economics, University of Munich, Schackstr. 4, 80539 Munich, Germany; Johannes Ludsteck (Johannes.Ludsteck@iab.de) is at the Institut für Arbeitsmarkt- und Berufsforschung, Regensburger Straße 104, 90478 Nürnberg, Germany. Ekkehart Schlicht takes main responsibility for the theoretical part and Johannes Ludsteck for the simulations. An earlier version of the paper has been presented at the Econometric Society European Meeting in Munich 1989.

The fact that the general conditions of life are not stationary is the source of many of the difficulties that are met with in applying economic doctrines to practical problems.

ALFRED MARSHALL (1949, v.iii.24)

1 Introduction

Economists, from KEYNES (1939) to LUCAS (1976), have have been uneasy about using regression analysis for estimating behavioral coefficients from economic time-series. The assumption that the coefficients describing economic behavior remain invariant over decades, let alone centuries, seems preposterous. KEYNES (1973, 294) objected against regression analysis (“Tinbergen’s method”) for this, as well as other, reasons: “One of the chief dilemmas facing you is, of course, ... that the method requires not too short a series whereas it is only in a short series, in most cases, that there is a reasonable expectation that the coefficients will be fairly constant.”

Further, the assumption that alternative courses of economic policy will leave the coefficients describing economic behavior unaffected points in a similar direction: “What happens if the phenomenon under investigation itself reacts on the factors by which we are explaining it?” (KEYNES 1939, 561). This point – known as the “Lucas Critique” – led to a rejection of allegedly Keynesian models in modern macroeconomics (LUCAS and SARGENT, 1979).

Some economists tried to account for these problems by allowing coefficients to change over time in a diffuse way, formalized by a random walk (COOLEY and PRESCOTT 1973, SCHLICHT 1973, ATHANS 1974). Here the standard linear model

$$y_t = a' x_t + u_t, \quad a, x_t \in \mathbb{R}^n, \quad u_t \sim \mathcal{N}(0, \sigma^2), \quad t = 1, 2, \dots, T$$

is replaced by

$$y_t = a_t' x_t + u_t, \quad u_t \sim \mathcal{N}(0, \sigma^2) \tag{1}$$

$$a_{t+1} = a_t + v_t, \quad v_t \sim \mathcal{N}(0, \Sigma) \tag{2}$$

with $y_t \in \mathbb{R}$, $x_t \in \mathbb{R}^n$ observations, $a_t \in \mathbb{R}^n$ coefficients to be estimated, $u_t \in \mathbb{R}$ normal disturbances with variance σ^2 , and $v_t \in \mathbb{R}^n$ normal perturbations in the random walk

(2) with variances $\sigma_1^2, \sigma_2^2, \dots, \sigma_n^2$, viz. a covariance matrix

$$\Sigma = \begin{pmatrix} \sigma_1^2 & & & 0 \\ & \sigma_2^2 & & \\ & & \ddots & \\ 0 & & & \sigma_n^2 \end{pmatrix}.$$

As a straightforward estimation procedure for the time-paths of the coefficients $a' = (a'_1, a'_2, \dots, a'_T)$, the Kalman-filter has been proposed.¹ It calculates $a_t^{Kalman} = E\{a_t | (y_1, y_2, \dots, y_T), (x_1, x_2, \dots, x_T)\}$ for all $t = 1, 2, \dots, T$.

This paper describes another method for estimating the time-paths of the coefficients. This method – termed VC (for “varying coefficients”) – offers the following *theoretical* advantages:

- The state of the coefficients at time t , viz. a_t , is estimated by using *all* information, rather than merely observations up to t : $a_t^{VC} = E\{a_t | (y_1, y_2, \dots, y_T), (x_1, x_2, \dots, x_T)\}$. This is preferable from a theoretical point of view.
- The estimator uses an orthogonal parametrization, instead of the usual parametrization by initial values. This effaces the problems associated with estimating, or otherwise providing, initial values.
- For the variances, a maximum likelihood estimator and a related moments estimator are derived. The moments estimator has a straightforward interpretation in small samples and coincides with the likelihood estimator in large samples. This lends intuitive appeal to the maximum likelihood estimator and statistical appeal to the moments estimator.

Further, VC provides the following *practical* advantages:

- It is implemented in software packages that are freely available and easy to use (LUDSTECK 2004, SCHLICHT 2005a,b). The more elaborate implementation (LUDSTECK 2004) relies on the commercial Mathematica program and uses a global maximization of the criterion function, rather than some gradient method as typically found in implementations of the Kalman filter. It permits, in addition, the analysis of panel data, a feature that is not available in current software packages, whether free or commercial.

¹ATHANS (1974), SCHLICHT (1977). For the treatment of time-varying coefficients with the Kalman filter see e.g. HAMILTON (1994, Sect. 13.8).

- Although the VC moments estimator, the VC likelihood estimator, and the Kalman filter yield almost identical results in many cases, the VC moments estimator is, from a computational point of view, sometimes superior to the other estimators in poorly conditioned (yet, we submit, practically important) cases.

From a *descriptive* point of view, the VC method has a clear-cut interpretation as a straightforward generalization of the method of least squares:

- While the method of ordinary least squares selects estimates that minimize the sum of squares $\sum_{t=1}^T u_t^2$, VC selects estimates that minimize the weighted sum of squares $\sum_{t=1}^T u_t^2 + \gamma_1 \sum_{t=2}^T v_1^2 + \gamma_2 \sum_{t=2}^T v_2^2 + \dots + \gamma_n \sum_{t=2}^T v_n^2$, where the weights for the changes in the coefficients $\gamma_1, \gamma_2, \dots, \gamma_n$ are determined by the inverse variance ratios, *i.e.* $\gamma_i = \sigma^2 / \sigma_i^2$. In other words, it balances the desiderata of a good fit and parameter stability.
- The time-averages of the regression coefficients are GLS estimates of the corresponding regression with fixed coefficients, *i.e.* $\frac{1}{T} \sum_t a_t = a_{GLS}$.

Both features permit an easier interpretation of the results in relation to the original data than seems possible from the Kalman filter perspective.

The paper is organized as follows: Sections 2 to 4 gives some notation, introduce the filter and present some preliminary results. Sections 5 to 8 derive the likelihood estimator and the moments estimator and compare them. Sections 10 to 11 illustrate the estimators by means of some simulations and compare the performance with the Kalman filter, as implemented in the popular program package EViews (2005). The final section 12 summarizes the argument and findings.

2 Notation

Define

$$\begin{array}{cccc}
 y := \begin{pmatrix} y_2 \\ \cdot \\ \cdot \\ y_T \end{pmatrix}, & u := \begin{pmatrix} u_1 \\ u_2 \\ \cdot \\ \cdot \\ u_T \end{pmatrix}, & a := \begin{pmatrix} a_1 \\ a_2 \\ \cdot \\ \cdot \\ a_T \end{pmatrix}, & v := \begin{pmatrix} v_2 \\ v_3 \\ \cdot \\ \cdot \\ v_T \end{pmatrix} \\
 \text{order} & T \times 1 & T \times 1 & Tn \times 1 & (T-1)n \times 1
 \end{array}$$

$$X := \begin{pmatrix} x'_1 & & 0 \\ & x'_2 & \\ & & \ddots \\ 0 & & & x'_T \end{pmatrix}, \quad P := \begin{pmatrix} -I_n & I_n & & 0 \\ & -I_n & I_n & \\ & & \ddots & \ddots \\ 0 & & & -I_n & I_n \end{pmatrix} \quad (3)$$

$$\begin{array}{ccc} \text{order} & T \times Tn & (T-1)n \times Tn \end{array}$$

and write (1), (2) as

$$y = Xa + u, \quad u \sim \mathcal{N}(0, \sigma^2 I_T) \quad (4)$$

$$Pa = v, \quad v \sim \mathcal{N}(0, V), \quad V := I_{T-1} \otimes \Sigma. \quad (5)$$

Denote by $e_i \in \mathbb{R}^n$ the n -th column of an $n \times n$ identity matrix and define the $(T-1) \times (T-1)n$ -matrix

$$F_i := I \otimes e'_i \quad (6)$$

that picks the time-path of the i -th disturbance $v_i = (v_{i,2}, v_{i,3}, \dots, v_{i,T})'$ from the disturbance vector v :

$$v_i := F_i v.$$

Define further the $Tn \times n$ matrix

$$Z := \frac{1}{\sqrt{T}} \begin{pmatrix} I_n \\ I_n \\ \vdots \\ I_n \end{pmatrix} \quad (7)$$

and note that the square matrix (P', Z) is of full rank. Note further that

$$PZ = 0, \quad Z'Z = I_n, \quad P'(PP')^{-1}P + ZZ' = I_{Tn} \quad (8)$$

where the last equality is implied by the identity

$$\begin{pmatrix} P' & Z \end{pmatrix} \left(\begin{pmatrix} P \\ Z' \end{pmatrix} \begin{pmatrix} P' & Z \end{pmatrix} \right)^{-1} \begin{pmatrix} P \\ Z' \end{pmatrix} = I_{Tn}.$$

3 Orthogonal Parametrization

For purposes of estimation we need a model that explains the observation x as a function of the random variables u and v . This would permit calculating the probability distribution of the observations x contingent on the parameters of the distributions of u and v , viz. σ^2 and Σ . The true model does not permit such an inference, though, because the matrix P in (3) is of rank $(T-1)n$ rather than of rank Tn and cannot be inverted. Hence v does not determine a unique y but rather the set of solutions

$$A := \left\{ a = P'(PP')^{-1}v + Z\lambda \mid \lambda \in \mathbb{R}^n \right\}. \quad (9)$$

For any v we have $a \in A \Leftrightarrow Pa = v$. Hence equation (4) and the set (9) give equivalent descriptions of the relationship between a and v in this sense.

In view of (9), any solution a to $Pa = v$ can be written as

$$a = P'(PP')^{-1}v + Z\lambda \quad (10)$$

for some $\lambda \in \mathbb{R}^n$. As $y = Xa + u$, equation (4) can be re-written as

$$y = u + XP'(PP')^{-1}v + XZ\lambda. \quad (11)$$

The model (10), (11) will be referred to as the *equivalent orthogonally parametrized model*. It implies the *true model* (4), (5). It implies, in particular, that a_t is a random walk even though a_t depends, according to (10), on past *and* future realizations of v_t .

Equation (11) permits calculating the density of y dependent upon the parameters of the distributions of u and v and the formal parameters λ . In a second step, all these parameters – σ_u^2 , Σ , and λ – can be determined by the maximum likelihood principle. This will give our likelihood estimates. Our moments estimates – to be introduced later – will build on the equivalent orthogonally parametrized model as well.

The orthogonal parametrization entails some advantages with respect to symmetry and mathematical transparency, as compared to more usual procedures, such as parametrization by initial values. By assuming some initial values $a_1 = \bar{a}_1$, the system (2) can be solved recursively, giving a as a function of v and $\bar{a}_1$, and the analysis would then proceed in a similar way as indicated above. Theoretically speaking, and with regard to likelihood estimation, all parameterizations are equivalent, but initial values are more cumbersome to implement than the formal parameters λ . It is for this reason that, in the context of Kalman filtering, initial values are estimated as posterior

means, *viz.* by running the filter back and forth in order to determine the necessary initial values $\bar{a}_1$ iteratively (AKAIKE, 1989, 61-2). The orthogonal parametrization proposed by SCHLICHT (1985, 55-8) will permit us to write down an explicit likelihood function and estimate all relevant parameters in a unified one-shot procedure.

Although the orthogonal parametrization may appear not very intuitive at first sight, it has a straightforward interpretation: The formal parameter vector $\lambda \in \mathbb{R}^n$ expresses additive shifts in the level of the parameters a that leave the disturbance vector v unaffected.

The formal parameter vector λ relates directly to the coefficient estimates of a standard GLS regression. Equation (11) can be interpreted as a standard regression for this parameter vector:

$$y = XZ\lambda + w \quad (12)$$

with $w := XP'(PP')^{-1}v + u$ distributed normally:

$$w \sim \mathcal{N}(0, W), \quad W := XBX' + \sigma^2 I_T \quad (13)$$

with

$$B := P'(PP')^{-1}V(PP')^{-1}P \quad (14)$$

The maximum likelihood (Aitken, GLS) estimate $\hat{\lambda}$ satisfies

$$Z'X'W^{-1}(y - XZ\hat{\lambda}) = 0 \quad (15)$$

or

$$\hat{\lambda} = (Z'X'W^{-1}XZ)^{-1}Z'X'W^{-1}y$$

This can be related to the corresponding standard linear regression with constant coefficients. Define the matrix of observations in conventional format as

$$X^* := \sqrt{T}XZ = X \begin{pmatrix} I_n \\ I_n \\ \cdot \\ I_n \end{pmatrix} = \begin{pmatrix} x'_1 \\ x'_2 \\ \cdot \\ x'_T \end{pmatrix}. \quad (16)$$

This matrix is assumed to be of full rank:

$$r(X^*) = r(XZ) = n \quad (17)$$

Inserting (16) into (12) gives rise to

$$y = \frac{1}{\sqrt{T}} X^* \lambda + w$$

and $\hat{\beta} = \frac{1}{\sqrt{T}} \hat{\lambda}$ turns out to be the standard GLS estimator for the regression $y = X^* \beta + w$.

4 The Filter

This section derives the filter for given variances σ^2 and Σ .

For given λ and X , the vectors y and a can be viewed as realizations of random variables determined jointly by the system (10), (12) as brought about by the disturbances u and v :

$$\begin{pmatrix} y \\ a \end{pmatrix} = \begin{pmatrix} XZ \\ Z \end{pmatrix} \lambda + \begin{pmatrix} I_T & XP'(PP')^{-1} \\ 0 & P'(PP')^{-1} \end{pmatrix} \begin{pmatrix} u \\ v \end{pmatrix}$$

The marginal distribution of y is as given by (12) and (13). The conditional distribution of a for given y is

$$a \sim \mathcal{N}(Z\lambda + BX'W^{-1}(y - XZ\lambda), A) \quad (18)$$

with

$$A := B - BW^{-1}B'$$

Hence a likelihood estimator for a can be derived by plugging the Aitken estimator $\hat{\lambda}$ from (15) into (18) and calculating the mean:

$$\hat{a} := Z\hat{\lambda} + BX'W^{-1}(y - XZ\hat{\lambda}). \quad (19)$$

Note that $Z'B = 0$ implies

$$Z'\hat{a} = \hat{\lambda}$$

and hence

$$\frac{1}{T} \sum_{t=1}^T a_{i,t} = \beta_i, \quad i = 1, 2, \dots, n \quad (20)$$

with $a_{i,t}$ as the i -th element of a_t and β_i as the i -th element of the Aitken estimate in the GLS regression $y = X^* \beta + w$. Note further that the variance-covariance matrix of w , as given in equation (13), tends to $\sigma^2 I_T$ if the variances σ_i^2 go to zero. In this

sense, the standard regression model is covered as a special limiting case of the model discussed here: The time-averages of the estimated coefficients $a_{i,t}$ are equal to the corresponding Aitken estimates β_i . For zero variance in the coefficients, equation (11) turns into a standard unweighted linear regression.

The estimator (19) can be characterized in an easy way. For given observations X and y and any given $\hat{a}$, the estimated disturbances are

$$\hat{u} = y - X\hat{a} \quad (21)$$

$$\hat{v} = P\hat{a} \quad (22)$$

and the weighted sum of squares of these disturbances is

$$Q := \frac{1}{\sigma^2} \hat{u}' \hat{u} + \hat{v}' V^{-1} \hat{v}. \quad (23)$$

The following theorem states that the likelihood estimator $\hat{a}$ can be obtained by minimizing this expression.

Theorem 1. *Minimizing the sum of squares (23) with respect to a yields a unique solution that is numerically identical to estimator $\hat{a}$ given in (19).*

Proof. Consider first the necessary conditions for a minimum of (23). Inserting (21) and (22) into (23) and putting the derivative with respect to a to zero gives the necessary condition

$$(X'X + \sigma^2 P' V^{-1} P) a = X' y \quad (24)$$

1. It will be shown first that the system matrix

$$M := (X'X + \sigma^2 P' V^{-1} P) \quad (25)$$

is of full rank. This will establish uniqueness. For $S = \frac{1}{\sigma^2} V$ we have

$$\begin{aligned} r(M) &= r \left(\begin{pmatrix} X' & P' S^{-\frac{1}{2}} \end{pmatrix} \begin{pmatrix} X \\ S^{-\frac{1}{2}} P \end{pmatrix} \right) \\ &= r \left(\begin{pmatrix} X' & P' S^{-\frac{1}{2}} \end{pmatrix} \right). \end{aligned}$$

If $\begin{pmatrix} X' & P' \end{pmatrix}$ were not of full rank, there would exist vectors $c_t \in \mathbb{R}^n$, $T = 1, 2, \dots, T$, not

all of them zero, such that

$$X'c_1 = P'S^{-\frac{1}{2}} \begin{pmatrix} c_2 \\ c_3 \\ \cdot \\ c_T \end{pmatrix} \quad (26)$$

Pre-multiplication by Z' gives

$$Z'X'c_1 = 0.$$

Together with (17) this implies $c_1 = 0$. As $P'S^{-\frac{1}{2}}$ has full rank, (26) would imply all c_t to be zero – a contradiction. Hence the system matrix is of full rank.

2. As the system matrix is positive definite, this establishes that the second order condition is fulfilled and that the solution to (24) gives indeed a global minimum.

3. We show now that (19) implies (24). Pre-multiplication of (19) by $(X'X + \sigma^2 P'V^{-1}P)$ gives

$$\begin{aligned} (X'X + \sigma^2 P'V^{-1}P) \hat{a} &= (X'X + \sigma^2 P'V^{-1}P) Z \hat{\lambda} \\ &\quad + X'XBX'W^{-1}(y - XZ\hat{\lambda}) \\ &\quad + \sigma^2 P'V^{-1}PBX'W^{-1}(y - XZ\hat{\lambda}). \end{aligned}$$

Since $PZ = 0$, the first term reduces to $X'XZ\hat{\lambda}$. Since $W = XBX' + \sigma^2 I_T$ we have $XBX'W^{-1} = I_T - \sigma^2 W^{-1}$ and the second term reduces to

$$X'(I_T - \sigma^2 W^{-1})(y - XZ\hat{\lambda}) = X'(y - XZ\hat{\lambda}) - \sigma^2 X'W^{-1}(y - XZ\hat{\lambda}).$$

Because $P'S^{-1}PB = P'(PP')^{-1}P$ and (8), the third term reduces to

$$\sigma^2 (I_{Tn} - ZZ')X'W^{-1}(y - XZ\hat{\lambda}).$$

By using the definition of $\hat{\lambda}$ in (15), this simplifies further to

$$\sigma^2 X'W^{-1}(y - XZ\hat{\lambda}).$$

Collecting terms gives

$$(X'X + \sigma^2 P'V^{-1}P) \hat{a} = X'y \quad (27)$$

which is identical to the definition of a given in (24).

The normal equation (27) can be equivalently expressed in terms of the variance

ratios

$$r_i := \frac{\sigma_i^2}{\sigma^2}, \quad r = \begin{pmatrix} r_1 \\ r_2 \\ \vdots \\ r_n \end{pmatrix}. \quad (28)$$

Define the matrix of variance ratios

$$R := \text{diag} r = \frac{1}{\sigma^2} \Sigma$$

and the covariance matrix

$$S := I_{T-1} \otimes R = \frac{1}{\sigma^2} V. \quad (29)$$

Using these expressions, the normal equation (27) can be written as

$$(X'X + P'S^{-1}P) \hat{a} = X'y. \quad (30)$$

or, for short

$$M \hat{a} = X'y \quad (31)$$

with M of order $nT \times nT$ as the system matrix, as in (25).

The estimates for $\hat{a}$ and for the corresponding disturbances $\hat{u}$ and $\hat{v}$ are

$$\begin{aligned} \hat{a} &= M^{-1} X' y \\ \hat{u} &= (I - X M^{-1} X') y \\ \hat{v} &= P M^{-1} X' y. \end{aligned}$$

All these are functions of the variance ratios r and the observations (y, X) .

Next turn to the covariance of the estimate $\hat{a}$. As the estimate satisfies the normal equation (31) it follows together with (1) and (2) that

$$\begin{aligned} \hat{a} &= M^{-1} X' (Xa + u) \\ &= M^{-1} (X' X a + X' u + P' S^{-1} P a - P' S^{-1} P a) \\ &= a + M^{-1} (X' u - P' S^{-1} v). \end{aligned} \quad (32)$$

Given a realization of the time-path of the coefficients a , the estimator $\hat{a}$ is normally distributed with mean a and covariance

$$E\{(\hat{a} - a)^2\} = \sigma^2 M^{-1}. \quad (33)$$

The system matrix (25) is determined by the observations X and the variance ratios

r . If we know in addition the variance of the error term σ^2 , equation (33) gives the precision of our estimate. The next step is to determine the variance σ^2 and the variances Σ .

5 Maximum Likelihood Estimation of the Variances

This section derives a maximum-likelihood estimator for the variances.

The likelihood function for (12), (13) is

$$L\left(\sigma^2, \Sigma\right) = -\log \det W - \hat{w}' W^{-1} \hat{w} \quad (34)$$

with

$$\hat{w} = (y - XZ\hat{\lambda})$$

The construction of the matrix W according to (13) and (14) is involved and requires the inversion of large full matrices. A considerable simplification can be achieved by expressing the likelihood in terms of the system matrix (25), the sum of squares (23), and the variances σ^2, Σ .

Theorem 2: *Disregarding constants, the likelihood function (34) can be written equivalently in terms of the variances σ^2, Σ as*

$$\begin{aligned} \mathcal{L}\left(\sigma^2, \Sigma\right) &= -(\log \det M + (T-1) \log \det \Sigma - (Tn - T - n) \log \sigma^2) \\ &\quad - \left(\frac{1}{\sigma^2} \hat{u}' \hat{u} + \hat{v}' V^{-1} \hat{v}\right) \end{aligned} \quad (35)$$

Proof. Consider the first term first. Define

$$N := XP'(PP')^{-1} S^{\frac{1}{2}}.$$

With (13), (14), and (29) we have

$$W = \sigma^2 (NN' + I_T). \quad (36)$$

Consider now

$$\begin{aligned} \frac{\det M}{\det PP'} &= \det M \det (PP')^{-1} \\ &= \det M \det \begin{pmatrix} (PP')^{-1} & 0 \\ 0 & I_n \end{pmatrix} \end{aligned}$$

$$\begin{aligned}
 &= \det M \det \begin{pmatrix} (PP')^{-1} & 0 \\ 0 & I_n \end{pmatrix} \begin{pmatrix} (PP')^{-1} & 0 \\ 0 & I_n \end{pmatrix} \begin{pmatrix} P \\ Z' \end{pmatrix} \begin{pmatrix} P' & Z \end{pmatrix} \\
 &= \det M \det \begin{pmatrix} P' & Z \end{pmatrix} \begin{pmatrix} (PP')^{-1} & 0 \\ 0 & I_n \end{pmatrix} \begin{pmatrix} (PP')^{-1} & 0 \\ 0 & I_n \end{pmatrix} \begin{pmatrix} P \\ Z' \end{pmatrix} \\
 &= \det \begin{pmatrix} (PP')^{-1} & 0 \\ 0 & I_n \end{pmatrix} \begin{pmatrix} P \\ Z' \end{pmatrix} M \begin{pmatrix} P' & Z \end{pmatrix} \begin{pmatrix} (PP')^{-1} & 0 \\ 0 & I_n \end{pmatrix} \\
 &= \det \begin{pmatrix} (PP')^{-1} P M P' (PP')^{-1}, & (PP')^{-1} P M Z \\ Z' M P' (PP')^{-1}, & Z' M Z \end{pmatrix} \\
 &= \det \begin{pmatrix} (PP')^{-1} P X' X P' (PP')^{-1} + S^{-1}, & (PP')^{-1} P X' X Z \\ Z' X' X P' (PP')^{-1}, & Z' X' X Z \end{pmatrix} \\
 &= \det \begin{pmatrix} (PP')^{-1} P X' X P' (PP')^{-1} + S^{-1}, & (PP')^{-1} P X' X Z \\ Z' X' X P' (PP')^{-1}, & Z' X' X Z \end{pmatrix} \\
 &= \det \begin{pmatrix} S & 0 \\ 0 & I_n \end{pmatrix}^{-1} \det \begin{pmatrix} N' N + I_{(T-1)n}, & N' X Z \\ Z' X' N, & Z' X' X Z \end{pmatrix} \\
 &= \det S^{-1} \det (N' N + I_{(T-1)n}) \cdot \\
 &\quad \det (Z' X' X Z - N' X Z (N' N + I_{(T-1)n})^{-1} Z' X' N).
 \end{aligned}$$

Because

$$Z' X' X P' = 0$$

we have

$$Z' X' N = 0$$

and we obtain the result

$$\frac{\det M}{\det P P'} = \det S^{-1} \det (N' N + I_{(T-1)n}) \det (Z' X' X Z). \quad (37)$$

As $N' N$ and $N N'$ have all non-zero eigenvalues in common, and adding the identity matrix shifts the eigenvalues by unity, the products of the eigenvalues of $(N' N + I_{(T-1)n})$ and of $(N N' + I_T)$ and the corresponding determinants are identical. In view of (36) we obtain from (37)

$$\frac{\det M}{\det (P P')} = \frac{\det W \det (Z' X' X Z)}{(\det \Sigma)^{T-1}} (\sigma^2)^{Tn-T-n}$$

and hence

$$\begin{aligned} \log \det W &= \log \det M + (T-1) \log \det \Sigma - (Tn - T - n) \log \sigma^2 \\ &\quad - \log \det (PP') - \log \det (Z'X'XZ). \end{aligned}$$

For any given problem, the last two terms are constant and can be dropped. Hence the first term in the likelihood function (35) is established.

Consider the second term next. The sum of squares at $a = \hat{a}$ is

$$\begin{aligned} Q &= \frac{1}{\sigma^2} \hat{u}' \hat{u} + \hat{v}' V^{-1} \hat{v} \\ &= \frac{1}{\sigma^2} (y' - \hat{a}' X') (y - X \hat{a}) + \hat{a}' P' V^{-1} P \hat{a} \\ &= \frac{1}{\sigma^2} y' y - \frac{2}{\sigma^2} y' X \hat{a} + \frac{1}{\sigma^2} \hat{a}' (X' X + \sigma^2 P' V^{-1} P) \hat{a} \\ &= \frac{1}{\sigma^2} y' y - \frac{1}{\sigma^2} y' X \hat{a} \end{aligned}$$

From (19) and (30) we find $y' X' \hat{a} = y' y - y' \hat{w} + y' X B X' W^{-1} \hat{w}$. Inserting this into (35) yields

$$Q = \frac{1}{\sigma^2} ((\hat{w}' + \hat{\lambda}' Z' X') \hat{w} - (\hat{w}' + \hat{\lambda}' Z' X') X B X' W^{-1} \hat{w}). \quad (38)$$

From the definition (13) we find

$$X B X' W^{-1} = I_T - \sigma^2 W^{-1}.$$

and (38) reduces to

$$Q = \hat{w}' W^{-1} \hat{w} - \hat{\lambda}' Z' X' W^{-1} \hat{w}. \quad (39)$$

As $w = y - X Z \hat{\lambda}$, (15) implies that the last term in (39) vanishes and we obtain

$$Q = \frac{1}{\sigma^2} \hat{u}' \hat{u} + \hat{v}' V^{-1} \hat{v} = \hat{w}' W^{-1} \hat{w}$$

which completes the proof.

6 Moments Estimation of the Variances

The maximum likelihood estimates of the variances lack intuitive appeal. The moments estimator that will be developed in this section has, for any sample size, a straightforward interpretation: It is defined by the property that estimated variances are equalized to their expectations. It will turn out that the estimators coincide for

large samples, and do not differ much even in small samples, the intuitive appeal of the moments estimator carries over to the likelihood estimator, and the attractive large-sample properties of the likelihood estimator carry over to the moments estimator.

The estimated coefficients $\hat{a}$ along with the estimated disturbances $\hat{u}$ and $\hat{v}$ are random variables brought about by realizations of the random variables u and v . Consider $\hat{u} = y - X\hat{a} = X(a - \hat{a}) + u$ first. With (32) we obtain

$$\hat{u} = (I_T - XM^{-1}X')u + XM^{-1}P'S^{-1}v.$$

Regarding $\hat{v}$, pre-multiplying (32) by F_i from (6) yields

$$\hat{v}_i = F_i(I_{(T-1)n} - PM^{-1}P'S^{-1})v + F_iPM^{-1}X'u$$

Thus $\hat{u}$ and $\hat{v}_i$ are linear functions of the normal random variables u and v , and their expected squared errors can be calculated.

$$\begin{aligned} E\{\hat{u}'\hat{u}\} &= E\{(u'(I_T + XM^{-1}X') + v'S^{-1}PM^{-1}X') \cdot \\ &\quad ((I_T - XM^{-1}X')u + XM^{-1}P'S^{-1}v)\} \\ &= E\{u'(I_T - XM^{-1}X')(I_T - XM^{-1}X')u\} + \\ &\quad + E\{v'S^{-1}PM^{-1}X'XM^{-1}P'S^{-1}v\} \\ &= \text{tr}E\{u'(I_T - XM^{-1}X')(I_T - XM^{-1}X')u\} + \\ &\quad + \text{tr}E\{v'S^{-1}PM^{-1}X'XM^{-1}P'S^{-1}v\} \\ &= \text{tr}E\{(I_T - XM^{-1}X')uu'(I_T - XM^{-1}X')\} + \\ &\quad + \text{tr}E\{XM^{-1}P'S^{-1}vv'S^{-1}PM^{-1}X'\} \\ &= \text{tr}\sigma^2(I_T - XM^{-1}X')(I_T - XM^{-1}X') + \text{tr}\sigma^2XM^{-1}P'S^{-1}PM^{-1}X' \\ &= \sigma^2\text{tr}((I_T - XM^{-1}X')(I_T - XM^{-1}X') + XM^{-1}P'S^{-1}PM^{-1}X') \\ &= \sigma^2\text{tr}(I - 2XM^{-1}X' + XM^{-1}X'XM^{-1}X' + XM^{-1}P'S^{-1}PM^{-1}X') \\ &= \sigma^2\text{tr}(I_T - 2XM^{-1}X' + XM^{-1}(X'X + P'S^{-1}P)M^{-1}X') \\ &= \sigma^2\text{tr}(I_T - XM^{-1}X') \\ &= \sigma^2(T - \text{tr}XM^{-1}X') \end{aligned}$$

Next, consider the variance of $\hat{v}_i$.

$$\begin{aligned}
E\{\hat{v}_i' \hat{v}_i\} &= E\{(u' X M^{-1} P' F_i' + v' (I_{(T-1)n} - S^{-1} P M^{-1} P') F_i') \cdot \\
&\quad (F_i P M^{-1} X' u + F_i (I_{(T-1)n} - P M^{-1} P' S^{-1}) v)\} \\
&= E\{u' X M^{-1} P' F_i' F_i P M^{-1} X' u\} + \\
&\quad + E\{v' (I_{(T-1)n} - S^{-1} P M^{-1} P') F_i' F_i (I_{(T-1)n} - P M^{-1} P' S^{-1}) v\} \\
&= E\{\text{tr} u' X M^{-1} P' F_i' F_i P M^{-1} X' u\} + \\
&\quad + E\{\text{tr} v' (I_{(T-1)n} - S^{-1} P M^{-1} P') F_i' F_i (I_{(T-1)n} - P M^{-1} P' S^{-1}) v\} \\
&= \text{tr} E\{F_i P M^{-1} X' u u' X M^{-1} P' F_i'\} + \\
&\quad + \text{tr} E\{F_i (I_{(T-1)n} - P M^{-1} P' S^{-1}) v v' (I_{(T-1)n} - S^{-1} P M^{-1} P') F_i'\} \\
&= \sigma^2 \text{tr} F_i P M^{-1} X' X M^{-1} P' F_i' + \\
&\quad + \sigma^2 \text{tr} F_i (I_{(T-1)n} - P M^{-1} P' S^{-1}) S (I_{(T-1)n} - S^{-1} P M^{-1} P') F_i' \\
&= \sigma^2 \text{tr} F_i P M^{-1} X' X M^{-1} P' F_i' + \\
&\quad + \sigma^2 \text{tr} F_i (I_{(T-1)n} - P M^{-1} P' S^{-1}) S (I_{(T-1)n} - S^{-1} P M^{-1} P') F_i' \\
&= \sigma^2 \text{tr} F_i S F_i' + \sigma^2 \text{tr} F_i P M^{-1} X' X M^{-1} P' F_i' + \\
&\quad + \sigma^2 \text{tr} F_i (-2 P M^{-1} P' + P M^{-1} P' S^{-1} P M^{-1} P') F_i' \\
&= (T-1) \sigma_i^2 + \\
&\quad + \sigma^2 \text{tr} F_i (P M^{-1} (X' X + P' S^{-1} P) M^{-1} P' - 2 P M^{-1} P') F_i' \\
&= (T-1) \sigma_i^2 - \sigma^2 \text{tr} F_i P M^{-1} P' F_i'
\end{aligned}$$

The moments estimators are defined by selecting variances σ^2 and $\sigma_i^2, i = 1, 2, \dots, n$ such that the expected moments $E\{\hat{u}' \hat{u}\}$ and $E\{\hat{v}_i' \hat{v}_i\}, i = 1, 2, \dots, n$ are equalized to the estimated moments $\hat{u}' \hat{u}$ and $\hat{v}_i' \hat{v}_i, i = 1, 2, \dots, n$. As both the expected moments and the estimated moments are functions of the variances, the moments estimators, denoted by $\check{\sigma}^2$ and $\check{\sigma}_i^2, i = 1, 2, \dots, n$, respectively, are defined as a fix point of the system

$$\frac{\hat{u}' \hat{u}}{\check{\sigma}^2} = T - \text{tr} X M^{-1} X' \quad (40)$$

$$\frac{\hat{v}_i' \hat{v}_i}{\check{\sigma}_i^2} = (T-1) - \frac{\check{\sigma}^2}{\check{\sigma}_i^2} \text{tr} F_i P M^{-1} P' F_i', \quad i = 1, 2, \dots, n \quad (41)$$

Summing (40) and (41) gives the corresponding calculated sum of squares (23):

$$Q = (n+1)T - n - \text{tr} M^{-1} X'X - \text{tr} \left(\sum_{i=1}^n \frac{\check{\sigma}^2}{\check{\sigma}_i^2} M^{-1} P' F_i' F_i P \right) \quad (42)$$

By definitions (3), (6), and (29) of P , F_i , and S we have

$$\sum_{i=1}^n \frac{\check{\sigma}^2}{\check{\sigma}_i^2} P' F_i' F_i P = P' S^{-1} P \quad (43)$$

and (42) reduces to

$$Q = (n+1)T - n - \text{tr} (M^{-1} (X'X + P' S^{-1} P)) = T - n$$

which implies

$$\check{\sigma}^2 = \frac{1}{T-n} (\hat{u}' \hat{u} + \hat{v}' S^{-1} \hat{v}).$$

7 Another Representation of the Moments Estimator

The relationship between the likelihood estimator and the moments estimator can be elucidated with the aid of the criterion function

$$\mathcal{K} \left(\sigma^2, \Sigma \right) = -(\log \det M + (T-1) \log \det \Sigma - T(n-1) \log \sigma^2) - \frac{1}{\sigma^2} \hat{u}' \hat{u} - \hat{v}' V^{-1} \hat{v} \quad (44)$$

defined in analogy to the concentrated likelihood function (35).

Theorem 3. *Minimization of the criterion function (44) with respect to the variances σ^2 and Σ yields the moments estimators as defined in (40), (41).*

Proof. Note that the envelope theorem together with (43) implies

$$\frac{\partial}{\partial \sigma^2} \left(\frac{1}{\sigma^2} \hat{u}' \hat{u} + \hat{v}' V^{-1} \hat{v} \right) = -\frac{1}{\sigma^4} \hat{u}' \hat{u} \quad (45)$$

$$\frac{\partial}{\partial \sigma_i^2} \left(\frac{1}{\sigma^2} \hat{u}' \hat{u} + \hat{v}' V^{-1} \hat{v} \right) = -\frac{1}{\sigma_i^4} \hat{v}_i' \hat{v}_i \quad (46)$$

In view of (43) we obtain further

$$\begin{aligned}\frac{\partial \log \det M}{\partial \sigma^2} &= \operatorname{tr} P' S^{-1} P M^{-1} \\ \frac{\partial \log \det M}{\partial \sigma_i^2} &= -\frac{\sigma^2}{\sigma_i^4} \operatorname{tr} (M^{-1} P' F_i' F_i P).\end{aligned}$$

With these results we find

$$\frac{\partial \mathcal{K}}{\partial \sigma^2} = -\operatorname{tr} P' S^{-1} P M^{-1} + \frac{T(n-1)}{\sigma^2} + \frac{1}{\sigma^4} \hat{u}' \hat{u} = 0 \quad (47)$$

$$\frac{\partial \mathcal{K}}{\partial \sigma_i^2} = -\frac{\sigma^2}{\sigma_i^4} \operatorname{tr} P' F_i' F_i P M^{-1} - (T-1) \frac{1}{\sigma_i^2} + \frac{1}{\sigma_i^4} \hat{v}_i' \hat{v}_i = 0 \quad (48)$$

By definition (25) we have

$$(X'X + \sigma^2 P' S^{-1} P) M^{-1} = I$$

and therefore

$$\begin{aligned}\sigma^2 \operatorname{tr} (P' S^{-1} P M^{-1}) &= \operatorname{tr} (I_{Tn} - X' X M^{-1}) \\ &= Tn - \operatorname{tr} (X M^{-1} X')\end{aligned}$$

and (47) can be written as

$$T - \operatorname{tr} (X M^{-1} X') = \frac{\hat{u}' \hat{u}}{\sigma^2}$$

which is identical to (40), while (48) reduces directly to (41).

8 The Relationship Between the Likelihood and the Moments Estimator

The likelihood estimates $\hat{\sigma}^2$ and $\hat{\Sigma}$ maximize $\mathcal{L}()$ and the moments estimates $\check{\sigma}^2$ and $\check{\Sigma}$ maximize $\mathcal{K}()$. Hence we have

$$\begin{aligned}\mathcal{L}(\hat{\sigma}^2, \hat{\Sigma}) &\geq \mathcal{L}(\check{\sigma}^2, \check{\Sigma}) \\ \mathcal{K}(\hat{\sigma}^2, \hat{\Sigma}) &\leq \mathcal{K}(\check{\sigma}^2, \check{\Sigma})\end{aligned}$$

and

$$\mathcal{K}(\check{\sigma}^2, \check{\Sigma}) - \mathcal{L}(\check{\sigma}^2, \check{\Sigma}) \geq \mathcal{K}(\hat{\sigma}^2, \hat{\Sigma}) - \mathcal{L}(\hat{\sigma}^2, \hat{\Sigma}) \quad (49)$$

with strict inequality if the estimators differ.

The difference between the moments criterion (44) and the likelihood criterion (35) is

$$\mathcal{K}(\sigma^2, \Sigma) - \mathcal{L}(\sigma^2, \Sigma) = n \log \sigma^2$$

and (49) implies

$$\check{\sigma}^2 \geq \hat{\sigma}^2 \quad (50)$$

with strict inequality if the estimators differ.

For purposes of comparison and computation it is useful to parametrize the likelihood function and the criterion function by the variance σ^2 and the variance ratios $r = (r_1, r_2, \dots, r_n)$ instead of the variances σ^2 and Σ . As $\Sigma = \sigma^2 \text{diag } r$ we have

$$\begin{aligned} \mathcal{L}(\sigma^2, \sigma^2 \text{diag } r) &= -\log \det M - (T-1) \sum_{i=1}^n \log r_i + (T-1) \log \sigma^2 - \frac{1}{\sigma^2} Q^*(r) \end{aligned} \quad (51)$$

$$\begin{aligned} \mathcal{K}(\sigma^2, \sigma^2 \text{diag } r) &= -\log \det M - (T-1) \sum_{i=1}^n \log r_i + (T-n) \log \sigma^2 - \frac{1}{\sigma^2} Q^*(r) \end{aligned} \quad (52)$$

with

$$Q^*(r) := \hat{u}' \hat{u} + \hat{v}' S^{-1} \hat{v} = \sigma^2 Q.$$

Maximizing these functions with respect to σ^2 yields the necessary conditions

$$\hat{\sigma}^2 = \frac{1}{T-1} Q^*(\hat{r}) \quad (53)$$

$$\check{\sigma}^2 = \frac{1}{T-n} Q^*(\check{r}) \quad (54)$$

for maxima of the functions $\mathcal{L}$ and $\mathcal{K}$, where $\hat{r}$ and $\check{r}$ denote the variance ratios of the likelihood and the moments estimator, respectively.

Inserting (53) into (51) and (54) into (52), and disregarding constants, yields the concentrated likelihood function $\mathcal{L}^*$ and the concentrated criterion function $\mathcal{K}^*$:

$$\mathcal{L}^*(r) = -\log \det M - (T-1) \log Q^*(r) - (T-1) \sum_{i=1}^n \log r_i \quad (55)$$

$$\mathcal{K}^*(r) = -\log \det M - (T-n) \log Q^*(r) - (T-1) \sum_{i=1}^n \log r_i \quad (56)$$

We note that

$$\begin{aligned} \mathcal{L}^*(\hat{r}) &\geq \mathcal{L}^*(\check{r}) \\ \mathcal{K}^*(\hat{r}) &\leq \mathcal{K}^*(\check{r}) \end{aligned}$$

and hence

$$\mathcal{K}^*(\check{r}) - \mathcal{L}^*(\check{r}) \geq \mathcal{K}^*(\hat{r}) - \mathcal{L}^*(\hat{r})$$

or equivalently

$$Q^*(\check{r}) \geq Q^*(\hat{r}).$$

As Q^* is decreasing in the r_i 's (compare (46), we will expect by and large smaller variance ratios for the moments estimator as compared to the likelihood estimator, and hence slightly greater stability in the estimated coefficients over time, and, in view of (50) a larger variance in the disturbance u .

For long time series, the likelihood criterion (55) and the moments criterion (56) are practically identical and the estimates will coincide.

9 Computation

Note that it is only necessary to estimate the variance ratios in a first step. Once these ratios are determined, the variances follow directly through (53) or (54) and (28).

Regarding computation of the variance ratios, there are basically two strategies available. One is to start from (55) or (56) and determine the variance ratios r by a maximization routine. This strategy is implemented in LUDSTECK (2004). The other is to use a gradient method. The VC-package by SCHLICHT (2005a) uses this approach for the moments estimator. It proceeds from (40), (41) and (54). These equations imply

$$\frac{\hat{v}_i' \hat{v}_i}{Q^*} = \frac{\check{r}_i (T-1) + \text{tr} F_i P M^{-1} P' F_i'}{(T-n)}, \quad i = 1, 2, \dots, n. \quad (57)$$

and therefore

$$\check{r}_i = \frac{1}{T-1} \left(\frac{(T-n) \hat{v}_i' \hat{v}_i}{Q^*} - \text{tr} F_i P M^{-1} P' F_i' \right)$$

or

$$\hat{r}_i = \frac{1}{T-1} \left(\frac{(T-1) \hat{v}_i' \hat{v}_i}{Q^*} - \text{tr} F_i P M^{-1} P' F_i' \right)$$

The iteration starts with some initial values for the variance ratios r^0 and computes new ratios of estimated variances on the left-hand side of (57) and the corresponding ratio of expected variances on the right-hand side. If the ratio of estimated variances is larger than the ratio of expected variances, the corresponding variance ratio r_i is increased, or, in the opposite case, reduced. This adjustment continues until (57) is satisfied for all variance ratios.

Both approaches – the direct and the gradient approach – do not require matrix inversions, which may pose problems for large systems as the the inverse of the sparse $Tn \times Tn$ band matrix M is full. The solution to the normal equation can be computed in both cases conveniently by a Cholesky decomposition of the system matrix M . The determinant of system matrix M – which is needed in case of direct maximization – or the traces in (57) – which are needed in the case of the gradient method – can be determined in the course of the Cholesky decomposition as by-products.

10 Practical Performance

In order to convey an impression about how the estimation procedure works in practice, we discuss in the following a small example and some simulations. Estimation involves in all cases minimizing (23), which reduces to minimizing the weighted sum of squares

$$u' u + \sum_i \gamma_i v_i' v_i$$

with weights $(\gamma_1, \gamma_2, \dots, \gamma_n) = \left(\frac{\sigma^2}{\sigma_1^2}, \frac{\sigma^2}{\sigma_2^2}, \dots, \frac{\sigma^2}{\sigma_n^2} \right)$. These weights, rather than the variances, determine the estimated time-paths of the coefficients for any given set of observations. Hence we concentrate in the following on the estimation of these weights, rather than the variances.

Consider first the simplest case of two explanatory variables – an intercept term

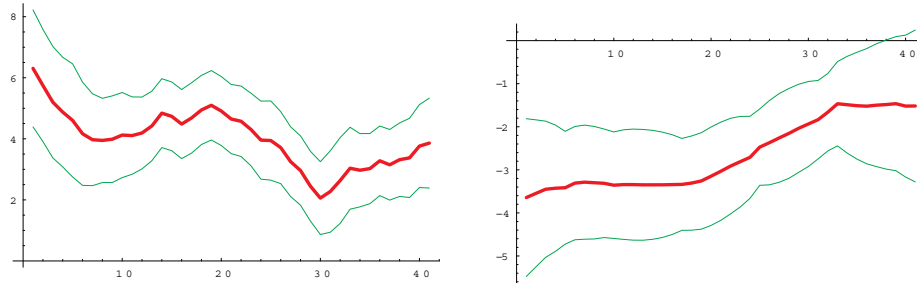


Figure 1: Time-varying intercept and slope for equation (58) with 95 percent confidence bands with data for Germany 1951-91.

and a single exogenous variable. The following model

$$dY_t = a_t + b_t du_t + \varepsilon_t \quad (58)$$

with dY_t as the percentage change in the GDP, du_t the change in the unemployment rate, ε_t as the error term, and a_t, b_t the time-varying coefficients of intercept and slope has been estimated. Taking data from Germany (provided in SCHLICHT 2005a), the moments estimator yields $\check{\gamma}_a = 3.48$ and $\check{\gamma}_b = 11.8$, $\check{\sigma}_\varepsilon^2 = 1.22$, $\check{\sigma}_a^2 = 0.35$, and $\check{\sigma}_b^2 = 0.10$. The corresponding likelihood estimates are $\hat{\gamma}_a = 2.83$, $\hat{\gamma}_b = 10.3$, $\hat{\sigma}_\varepsilon^2 = 1.15$, $\hat{\sigma}_a^2 = 0.4$, and $\hat{\sigma}_b^2 = 0.11$. The estimates obtained by the Kalman-Filter, as implemented in the program EViews (2005), are $\gamma_a^{EV} = 2.25$, $\gamma_b^{EV} = 8.88$, $\sigma_\varepsilon^{2EV} = 1.01$, $\sigma_a^{2EV} = 0.45$, and $\sigma_b^{2EV} = 0.11$. The estimates for the time-paths of the coefficients obtained from these three sets of estimators are practically indistinguishable. The estimated time paths for intercept a_t and slope b_t together with their 95 percent confidence bands (here for the moments estimator) are given in Figure 1. For comparison we note the OLS estimates $\bar{a}^{OLS} = 4.22$ and $\bar{b}^{OLS} = -2.98$ with standard errors $SE(a)^{OLS} = 0.27$ and $SE(b)^{OLS} = 0.30$ whereas the corresponding averages of the time-varying specification are $\bar{a}^{VC} = 4.0$ and $\bar{b}^{VC} = -2.65$ with average standard errors $SE(a^{VC}) = 0.66$ and $SE(b^{VC}) = 0.62$. Hence the averages of the results of the VC method are similar to the results obtained by ordinary least squares – see the discussion around equation (20) above.

All this does not tell very much, however, about how well the method recovers the variances and time-paths of the coefficients. In order to obtain an impression about this, some simulations were conducted.

The first exploration was to assume a model with an intercept term a_t and a single explanatory variable x_t with coefficient b_t . A time series for the explanatory

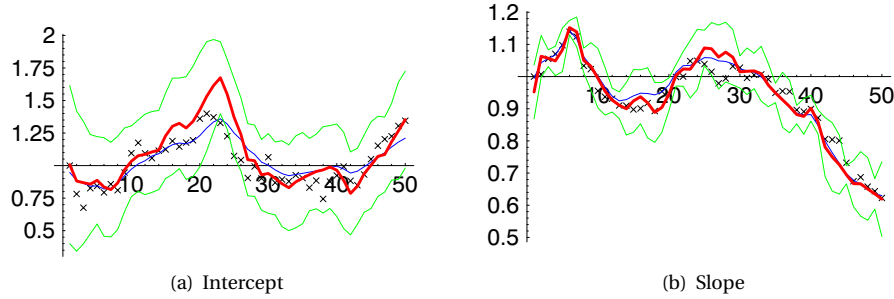


Figure 2: Optimally calculated expectations (thin lines) and VC estimates (thick lines) for intercept (left) and slope (right), together with the realizations of the coefficients (x) and the VC confidence bands. The example has been selected to visually exhibit differences between the true expectations and the VC estimates; usually the weights are estimated better and the curves lie quite close together. As the estimated smoothing weights are considerably smaller than the true weights, the time-paths of the VC estimates are less smooth than the true expectations (True weights are $\gamma_a = 10$ and $\gamma_b = 100$, while the estimated weights are $\hat{\gamma}_a = 1.60$ and $\hat{\gamma}_b = 14.76$ here. The true variances are $\sigma_u^2 = 0.1$, $\sigma_a^2 = 0.01$, and $\sigma_b^2 = 0.001$, the estimated variances are $\hat{\sigma}_u^2 = 0.040$, $\hat{\sigma}_a^2 = 0.025$, and $\hat{\sigma}_b^2 = 0.0029$.)

variable was generated with $x_t \sim \mathcal{N}(0, 100)$, $t = 1, 2, \dots, 50$. Further it was assumed that $u_t \sim \mathcal{N}(0, 0.1)$, $(a_t - a_{t-1}) \sim \mathcal{N}(0, 0.01)$, and $(b_t - b_{t-1}) \sim \mathcal{N}(0, 0.001)$. Typically the optimally computed expectations of the time paths (calculated by using the true variances) and the VC estimates lie very close together. Figure 2 illustrates a somewhat atypical run with estimated smoothing weights that deviate from the true smoothing weights by the order of five. The optimally estimated time-paths of the coefficients (based on the true variances) and the estimated time-paths (based on the estimated coefficients) move together. This illustrates the general impression that the filtering results, especially the qualitative time-patterns, are not extremely sensitive with regard to the weights used for filtering.

It is, obviously, never possible to extract the movement of the true coefficients from the data, irrespective how long the time series is. (Only the estimation of the weights will improve with the length of the time series.) The best that can be done is to estimate the expectations of the coefficients. Given the variances, the VC estimate (which is the mean of a normally distributed vector) is optimal and cannot be improved upon, and the standard of comparison must be the estimates obtained with optimal weights, as in Figure 2.

The distribution of the weights in the above setting is illustrated in Figure 3. The time series for x , u , and v have been generated as described above and the VC

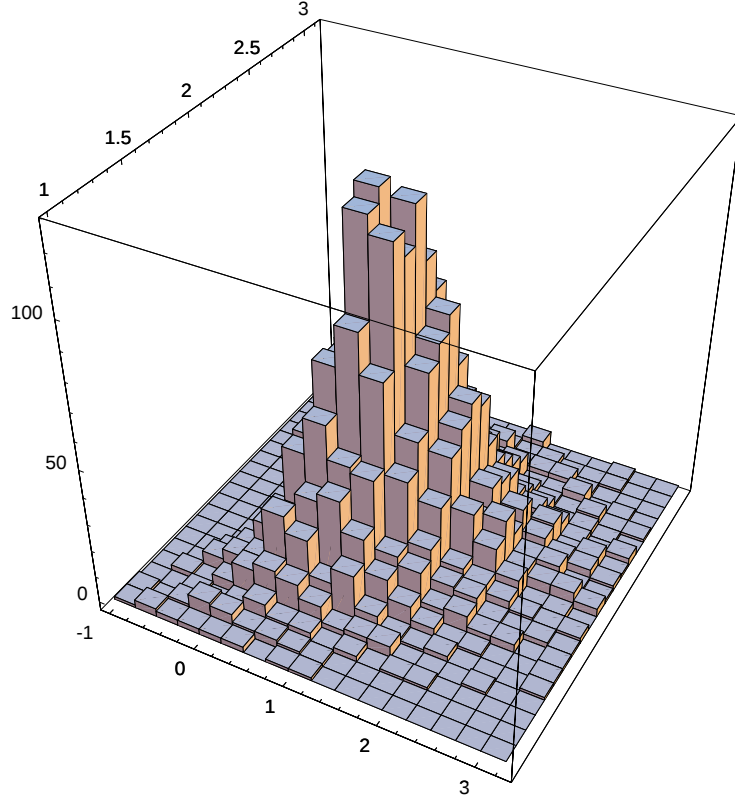


Figure 3: Histogram of estimates for the $\log_{10}$ weights. The theoretical values are $\log_{10} \gamma_a = 1$ and $\log_{10} \gamma_b = 2$. The distribution of estimates clusters around this peak. ($T = 50$, 5000 trials.)

moments estimation applied 5000 times. The histogram illustrates that the estimates cluster around their theoretical values.

Consider now the case of three exogenous variables: An intercept term $a_{1,t}$, and two explanatory variables $x_{2,t}$ and $x_{3,t}$ with associated coefficients $a_{2,t}$ and $a_{3,t}$. (The first explanatory variable is always taken as unity, *i.e.* $x_{1,t} = 1$.) Thus the model is summarized by

$$\begin{aligned} y_t &= a_{1,t} + a_{2,t}x_{2,t} + a_{3,t}x_{3,t} + u_t, \quad t = 1, 2, \dots, T \\ a_{i,t} &= a_{i,t-1} + v_i, \quad i = 1, 2, 3, \dots, n, \quad t = 2, 3, \dots, T \end{aligned} \quad (59)$$

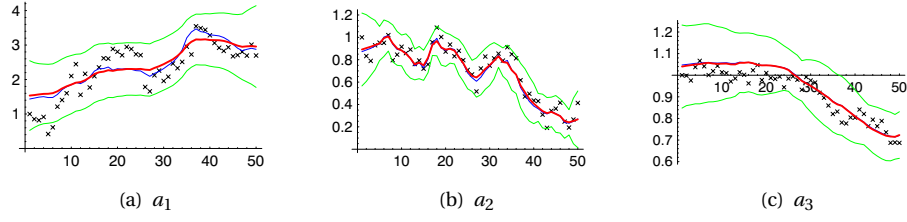


Figure 4: The time paths of the coefficients a_1 , a_2 , and a_3 : Theoretical expectations (thin lines), VC estimates (thick lines) and values (x), with estimated confidence bands. The theoretical weights are $\gamma_1 = 10$, $\gamma_2 = 100$, and $\gamma_3 = 1000$, the estimated weights are $\check{\gamma}_1 = 28.2$, $\check{\gamma}_2 = 177.4$, and $\check{\gamma}_3 = 994.9$. The theoretical variances are $\sigma_u^2 = 1$, $\sigma_1^2 = 0.1$, $\sigma_2^2 = 0.01$, and $\sigma_3^2 = 0.001$, the estimated variances are $\check{\sigma}_u^2 = 1.30$, $\check{\sigma}_1^2 = 0.046$, $\check{\sigma}_2^2 = 0.0073$, and $\check{\sigma}_3^2 = 0.0013$.

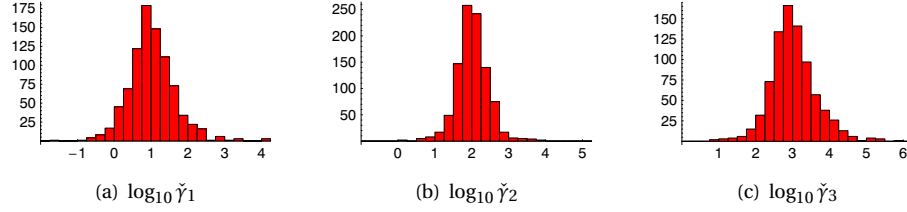


Figure 5: Histograms for the $\log_{10}$ weight estimates. The theoretical weights are $\log_{10} \gamma_1 = 1$, $\log_{10} \gamma_2 = 2$, and $\log_{10} \gamma_3 = 3$. The estimates cluster around these values. ($T = 50$, 1000 trials.)

with

$$u_t \sim \mathcal{N}(0, \sigma_u^2), \quad v_{i,t} \sim \mathcal{N}(0, \sigma_i^2).$$

We follow, again, the example given in the Mathematica notebook (LUDSTECK, 2004) to permit easy replication. The explanatory variables $x_{2,t}$ and $x_{3,t}$ are generated as normally distributed random variables, both with variance 100 and contemporaneous covariance 50. The number of periods is taken as $T = 50$. For our example we take $\sigma_u^2 = 1$, $\sigma_1^2 = 0.1$, $\sigma_2^2 = 0.01$, and $\sigma_3^2 = 0.001$. A typical run is illustrated Figure 4. Taking 1000 replications yields the distribution of estimates for the variance ratios given in Figure 5.

11 The VC Estimator and the Kalman Filter: Some Simulations

The VC estimator has been devised specifically for the model with time-varying coefficients (1), (2) that can also be estimated by means of the Kalman filter (HAMILTON 1994, Sect. 13.8). It is of interest, therefore, to compare the practical performance of the estimators by means of a simulation study. We concentrate on the estimation of weights, again, because the weights are determining the estimation of the time path of the coefficients in both cases. Given the weights, the estimation of the time-paths of the coefficients in VC is theoretically optimal, anyway, and the Kalman filter, being one-sided, can not do better. (We note however that, for given weights, both methods seem to produce nearly identical estimates for the time-paths of the coefficients.)

For purposes of comparison, we have conducted a small Monte-Carlo simulation study for some plausible parameter combinations. The Kalman-filter computations are performed using the popular statistical program package EViews (2005). With regard to the Eviews estimates it is to be noted that the package uses a gradient method that requires initial values for the variances. We have supplied the *true* values (which are, of course, in practical applications, unavailable) as initial values for the Eviews estimates throughout. Hence the practical performance of the program may be overstated in our examples.² For the VC estimator, we have used the moments rather than the likelihood version because some preliminary checks suggested to us that the moments estimator is to preferred in short time series and equivalent to the likelihood estimator for longer time series.

As in the previous illustrations, the simulation is based on a simple model with an intercept term a_t and a slope coefficient b_t that describes the effect of a single explanatory variable $x_t \sim \mathcal{N}(0, 5)$. The disturbance term is taken as $u_t \sim \mathcal{N}(0, 1)$ in all simulations. The starting values of the coefficients are taken as $a_1 = 1$ and $b_1 = 1$. Thus the variance of the disturbance u_t is relatively large, as compared to the variance in the explanatory variable x_t , and given the average levels of the coefficients a_t and b_t around unity. This comparatively ill-conditioned setting has been selected to bring out differences between the methods that would not be visible otherwise.

We have simulated time series of different length and for different combinations of variances. The cases studied are given by all combinations of

- the length of the time series from $T \in \{40, 60, 80, 100\}$
- variances from $\sigma_a^2 \in \{0.001, 0.01\}$ and $\sigma_b^2 \in \{0.001, 0.01\}$

²The fact that Eviews performed much worse if good starting values were not supplied indicates a possibly crucial problem for real applications.

For each combination of these parameters we have generated $K = 1000$ realizations of the time series by obtaining pseudorandom draws for the $x_t, u_t, v_{a,t}$, and $v_{b,t}$. It is, however, problematic, to directly compare the estimates of σ_u, σ_a , and σ_b with their theoretical values because the pseudorandom variates may deviate considerably from their theoretical values in small samples. In order to avoid this possible source of error, we have standardized the pseudorandom series by multiplying them with appropriate scale factors.

The relative performance of the VC moments estimator in comparison to that of the Kalman filter is then evaluated by computing the mean squared error of the logarithms of the estimated weights

$$MSE_{10}(\gamma_i) := \frac{1}{K} \sum_{k=1}^K \left(\log_{10} \gamma_{i,k}^{estimated} - \log_{10} \gamma_{i,k}^{true} \right)^2$$

for $i = a, b$ and all combinations of the parameters.

The reason for concentrating on the deviations of the logarithms is that the weights seem to act multiplicatively in the decomposition. Yet the use of logarithms gives rise to the problem that the logarithm of a weight approaches infinity if the variance of the corresponding coefficient is estimated as being close to zero, *viz.* that the coefficient is estimated as remaining constant over time. This renders the computation of mean squared errors infeasible. In order to avoid artifacts produced in this way, we have computed VC estimates under the restriction $\sigma_u^2 / \sigma_i^2 \leq 10^{10}$ and forced the Eviews estimates to conform to this restriction as well. (A weight of 10^{10} enforces, in the setting studied here, constancy of the corresponding parameter anyway. By imposing the restriction, we avoid the problem that differences in weights above this threshold may affect our results.)

Table 1 summarizes our findings. It gives the mean squared errors of the estimated logarithms of the weights for the various parameter constellations, as obtained by the Eviews Kalman filter, and the VC moments estimator, respectively. The performance of the VC estimator is slightly better.

A closer look at the results reconfirms the impression that both estimators perform very similar—with the *caveat* that the Eviews estimates have been calculated by using the theoretical values as starting values. Figure 6 illustrates this for the case by $\sigma_u^2 = 1, \sigma_a^2 = 0.01, \sigma_b^2 = 0.001$, and $T = 60$. The distributions of the estimates for the weights are practically indistinguishable.

		<i>EV</i>	<i>VC</i>	<i>EV</i>	<i>VC</i>	<i>EV</i>	<i>VC</i>	<i>EV</i>	<i>VC</i>
(γ_a, γ_b)		(1000, 1000)		(100, 1000)		(1000, 100)		(100, 100)	
$T = 40$	$MSE_{10}(\gamma_a)$	29.9	24.6	26.7	24.8	30.5	23.4	27.2	24.2
	$MSE_{10}(\gamma_b)$	25.1	22.6	25.8	23.4	12.7	11.8	13.2	11.9
$T = 60$	$MSE_{10}(\gamma_a)$	29.2	24.3	19.2	18.3	29.5	23.4	20.0	18.4
	$MSE_{10}(\gamma_b)$	21.6	20.3	21.4	20.0	6.9	6.6	7.1	6.9
$T = 80$	$MSE_{10}(\gamma_a)$	26.4	22.3	12.1	11.7	27.1	21.4	13.0	12.2
	$MSE_{10}(\gamma_b)$	16.8	15.5	17.0	15.7	2.9	3.1	3.3	3.3
$T = 100$	$MSE_{10}(\gamma_a)$	23.3	19.9	8.1	7.6	24.0	19.3	9.2	8.6
	$MSE_{10}(\gamma_b)$	13.2	12.3	13.6	12.5	1.3	1.3	1.6	1.5

Table 1: A comparison between the Kalman filter as implemented in Eviews (*EV*) and the VC moments estimator (*VC*), based on 1000 replications for each parameter constellation. The EV estimator is obtained by using the true variances as initial values while the VC estimator does not use initial values; yet the VC estimator outperforms the Kalman filter slightly.

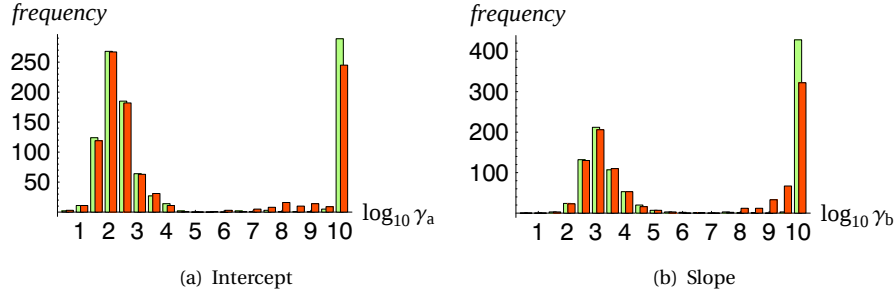


Figure 6: Histograms for the estimated log weights for intercept (left) and slope (right). Each pair of columns gives the frequencies of estimates obtained by the Eviews Kalman filter (left column) and the VC moments estimator (right column). The distributions are nearly identical. The estimated weights cluster around their theoretical values ($\log_{10} \gamma_a = 2$ and $\log_{10} \gamma_b = 3$). (The case depicted here is not well conditioned, being characterised by $\sigma_u^2 = 1$, $\sigma_a^2 = 0.01$, $\sigma_b^2 = 0.001$, and $T = 60$, 1000 replications. In roughly a third of the cases, one or the other weight is estimated exceeding 10^6 which implies that the corresponding parameter is estimated as invariant over time. Yet even here, both methods produce practically identical results.)

12 Concluding comments

The VC estimator discussed here is specifically designed to deal with linear models where the coefficients are following a random walk. The main advantage over the Kalman filter can be seen in its mathematical and descriptive transparency, its ease of use and its free availability. Our simulations suggest that the VC moment estimator and the Eviews Kalman filter perform very similar, provided the starting values of the Eviews filter are provided correctly. For practical purposes this is not feasible, though, and our experience suggests that the performance of the Eviews filter is quite unsatisfactory if proper initial values are not supplied. In our simulations we have used the true parameter values as starting values, though. Yet even under these circumstances, the VC estimator seems to perform slightly better. Our conclusion is that VC filter is preferable for studying the specific class of models for which it is designed. The Kalman filter remains, however, applicable to a much larger class of models. In this sense the method described here does not provide a substitute but rather a complement to existing approaches.

References

- AKAIKE, H. 1989, “Bayesian Modeling for Time Series Analysis,” in: R. S. MARIANO, K. S. LIM, and K. D. LING (eds.), *Advances in Statistical Analysis and Statistical Computing*, pp. 59–69, Greenwich (Conn.) and London: JAI Press.
- ATHANS, M. 1974, “The Importance of Kalman Filtering Methods for Economic Systems,” *Annals of Economic and Social Measurement*, 3, 24–49.
- COOLEY, T. F. and E. F. PRESCOTT 1973, “An Adaptive Regression Model,” *International Economic Review*, 14, 364–71.
- EViews 2005, “EViews 5.1 standard edition,” online at www.eviews.com.
- HAMILTON, J. D. 1994, *Time Series Analysis*, Princeton University Press, Princeton.
- KEYNES, J. M. 1939, “Professor Tinbergen’s Method,” *Economic Journal*, 49, 558–68.
- 1973, *The General Theory and After, Part II: Defense and Development*, vol. xiv of *The Collected Works of John Maynard Keynes*, Macmillan, London.
- LUCAS, R. 1976, “The Econometric Policy Evaluation: A Critique,” in: K. BRUNNER and A. H. MELTZER (eds.), *Phillips Curve and Labor Markets*, pp. 19–46, North Holland, Amsterdam.
- LUCAS, R. and T. J. SARGENT 1979, “After Keynesian Macroeconomics,” *Quarterly Review of the Federal Reserve Bank of Minneapolis*, 3, online at <http://minneapolisfed.org/research/qr/qr321.pdf>.
- LUDSTECK, J. 2004, “VC Package for Mathematica,” online at <http://library.wolfram.com/infocenter/MathSource/5195/>.
- MARSHALL, A. 1949, *Principles of Economics*, 8th ed., London: Macmillan, online at <http://www.econlib.org/library/Marshall/marPtoc.html>, reprint of 8th edition 1920, first edition 1890.
- SCHLICHT, E. 1973, “Forecasting Markov Chains. A Theoretical Foundation for Exponential Smoothing,” Discussion paper 13, University of Regensburg, online at http://www.semverteilung.vwl.uni-muenchen.de/mitarbeiter/es/paper/schlicht-exponential_smoothing.pdf.
- 1977, *Grundlagen der ökonomischen Analyse*, Rowohlt, Reinbek, online at <http://epub.ub.uni-muenchen.de/archive/00000891/>.

- 1985, *Isolation and Aggregation in Economics*, Springer-Verlag Berlin-Heidelberg-New York, online at <http://epub.ub.uni-muenchen.de/archive/00000003/>.
- 2005a, “VC - A Program for Estimating Time-Varying Coefficients,” online at <http://epub.ub.uni-muenchen.de/archive/00000684/>.
- 2005b, “VCC - A Program for Estimating Time-Varying Coefficients. Console Version With Source Code in C,” online at <http://epub.ub.uni-muenchen.de/archive/00000719/>.