

Evans, George W.; Honkapohja, Seppo; Williams, Noah

**Working Paper**

## Generalized stochastic gradient learning

CESifo Working Paper, No. 1576

**Provided in Cooperation with:**

Ifo Institute – Leibniz Institute for Economic Research at the University of Munich

*Suggested Citation:* Evans, George W.; Honkapohja, Seppo; Williams, Noah (2005) : Generalized stochastic gradient learning, CESifo Working Paper, No. 1576, Center for Economic Studies and ifo Institute (CESifo), Munich

This Version is available at:

<https://hdl.handle.net/10419/19040>

**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.*

# GENERALIZED STOCHASTIC GRADIENT LEARNING

GEORGE W. EVANS  
SEPPO HONKAPOHJA  
NOAH WILLIAMS

CESIFO WORKING PAPER NO. 1576  
CATEGORY 5: FISCAL POLICY, MACROECONOMICS AND GROWTH  
OCTOBER 2005

*An electronic version of the paper may be downloaded*  
• *from the SSRN website:* [www.SSRN.com](http://www.SSRN.com)  
• *from the CESifo website:* [www.CESifo-group.de](http://www.CESifo-group.de)

# GENERALIZED STOCHASTIC GRADIENT LEARNING

## Abstract

We study the properties of generalized stochastic gradient (GSG) learning in forward-looking models. We examine how the conditions for stability of standard stochastic gradient (SG) learning both differ from and are related to E-stability, which governs stability under least squares learning. SG algorithms are sensitive to units of measurement and we show that there is a transformation of variables for which E-stability governs SG stability. GSG algorithms with constant gain have a deeper justification in terms of parameter drift, robustness and risk sensitivity.

JEL Code: C62, C65, D83, E10, E17.

Keywords: adaptive learning, E-stability, recursive least squares, robust estimation.

*George W. Evans*  
*Department of Economics*  
*University of Oregon*  
*Eugene, OR 97403-1285*  
*USA*  
*gevans@uoregon.edu*

*Seppo Honkapohja*  
*Faculty of Economics*  
*University of Cambridge*  
*Sidgwick Avenue*  
*Cambridge CB3 9DD*  
*United Kingdom*  
*smsh4@cam.ac.uk*

*Noah Williams*  
*Department of Economics*  
*Princeton University*  
*Fisher Hall*  
*Princeton, NJ 08544-1021*  
*USA*  
*noahw@princeton.edu*

We thank Chryssi Giannitsarou for useful comments. Honkapohja gratefully acknowledges financial support from ESRC grant RES-000-23-1152.

# 1 Introduction

Over the past decade or two there has been a significant amount of macroeconomic research studying the implications of adaptive learning. This literature replaces rational expectations with the assumption that economic agents are boundedly rational but employ statistical forecasting techniques, which may allow the possibility of the economy finding a rational expectations equilibrium in the long run. A large part of this literature has assumed that economic agents used versions of recursive least squares (RLS) algorithms in estimating the parameters required for making forecasts. Occasionally, alternatives to RLS have been considered. The stochastic gradient (SG) algorithm has emerged as a convenient alternative to RLS.

Recursive least squares is simply a recursive, “on-line” implementation of classical least squares. It requires updating both the parameter estimates and an estimate of the second moment matrix of the regressors. The parameters are updated in a way to satisfy the least squares orthogonality condition, and the second moment matrix weights the different elements of the vector of observations. By contrast, the stochastic gradient algorithm (which is also known as least mean squares) does not use information on the second moments of the data but instead uses a fixed weighting matrix, namely the identity matrix.<sup>1</sup> In this paper we also study a generalized stochastic gradient rule which allows the weight to be an arbitrary positive definite (but fixed) matrix.

In classical econometric setups least squares is known to be a consistent and asymptotically efficient estimation method and is clearly the most widely used estimation method in practical forecasting. However, SG estimators, though inefficient, are also consistent and are somewhat simpler to compute than RLS estimators. In fact, the main advantage of the SG algorithm appears to be not so much its plausibility as a model of the adaptive learning rule followed by economic agents, but rather its ease of implementation in simulations that involve studies of global aspects of learning in relatively complicated models. For examples of the latter, see (Bullard and Eusepi 2005) and (Evans and Honkapohja 2005). The fact that SG learning has been found to work well in complex environments suggests that it has certain robustness properties that are absent in RLS.

Our main goal in the current paper is to examine the properties of SG algorithms in self-referential models. In doing so, we are led to a natural extension of the SG algorithm, which we call generalized SG (GSG) learning. We explore several dimensions of the performance and justification of this learning rule.

The first issue concerns the convergence conditions for SG learning. In the literature on adaptive learning, an attraction of RLS is that its convergence properties are relatively easy to compute. In contrast to the classical statistical framework, macroeconomic models with learning are self-referential, i.e. the evolution of the endogenous variables is influenced by the learning process itself. This has the consequence that it is not a foregone conclusion that estimators will be consistent, as the feedback from the learning process to the evolution of

---

<sup>1</sup>For a brief introduction to SG learning and further references, see Section 3.5 of (Evans and Honkapohja 2001).

the state variables may lead the overall system to fail to converge to an equilibrium.

Under RLS it can be shown that a rational expectations equilibrium (REE) is locally stable under learning if what are known as expectational stability (E-stability) conditions hold. A question that has emerged is how general is the class of algorithms for which convergence is governed by E-stability. For example, it is well-known that some mild generalizations of RLS that weight different data points unequally are covered by the E-stability principle. Some recent papers have examined the relationship between E-stability and convergence of SG learning. In a simple class of models E-stability provides exactly the conditions for convergence of SG learning. However, the literature has shown that E-stability does not always imply convergence of SG learning; for an economic example see (Giannitsarou 2005).<sup>2</sup> Further references are given in Section 2.4.

These points raise a number of questions. First, what are the conditions for convergence of SG learning for a general class of linear models? It turns out that, in contrast to RLS, the stability properties of SG learning in general depend on the moment matrix of the regressors. Are there easily stated conditions under which convergence of SG learning is independent of this moment matrix? In particular, when does E-stability govern convergence of SG learning in forward-looking models?

Secondly, the issue of parameter drift has recently received increasing attention. (Sargent and Williams 2005) examine a model in which parameter drift takes the form of a random walk using an approximate Kalman filter, which is the Bayesian estimator. For particular priors on the parameter drift, this is equivalent to “constant gain” recursive least squares. In the current paper we show that for any given prior, an appropriate constant gain GSG algorithm approximates the Bayesian learning rule. This is a more compelling justification for GSG rules than those based on computational cost and numerical performance.

Thirdly, we also show that constant gain GSG learning has a dramatically different justification as a robust learning rule. If agents are uncertain about the true data generating process, they may want to employ a learning rule that performs well across a number of alternative models. GSG rules have precisely this property, as they are robust optimal predictors. Finally, we show that the constant gain GSG rule is optimal if agents are “risk-sensitive,” having greater risk aversion of a particular form: while least squares minimizes the expected sum of squared errors, the GSG rule minimizes the expected exponential of the sum of squared errors. Connections between robustness and risk sensitivity are well-established in the control theory literature (dating from (Jacobson 1973) and (Whittle 1990)), and they emerge again here.

This paper is structured as follows. After formally introducing the economic model and the GSG algorithm, Section 2 is devoted to the relationships between E-stability and the convergence conditions for SG learning. In particular, we present a strengthening of the E-stability condition that is sufficient to guarantee stability of SG learning in all cases. We also present an economic example that illustrates the conditions. In Section 3 we raise a disadvantage of SG estimators. One of the less appreciated advantages of RLS is that it

---

<sup>2</sup>(Sargent and Williams 2005) show that E-stability also need not imply convergence when agents learn via a Kalman filter.

is unit free, i.e. invariant to changes in measurement units. In contrast, SG estimators fail to possess this natural property. This provides one motivation for considering the more general class of GSG estimators (one of which is asymptotically equivalent to least squares). Section 4 turns to the further justification of GSG estimators in terms of parameter drift and robustness in the face of unknown model uncertainty. Some simple simulation exercises illustrate the robustness of SG learning. Section 5 concludes.

## 2 Generalized SG Learning

### 2.1 The Economic Model

We use the basic multivariate linear forward-looking model

$$\begin{aligned} y_t &= \alpha + AE_t^* y_{t+1} + Bw_t + \eta_t, \\ w_t &= Fw_{t-1} + e_t, \end{aligned} \tag{1}$$

where  $y_t$  is  $n \times 1$ , the  $k \times 1$  observed exogenous variables  $w_t$  are assumed to follow a known vector autoregression (VAR), and the unobserved shock  $\eta_t$  is white noise. The innovation  $e_t$  has zero mean and covariance matrix  $\Sigma_e$ .  $F$  is assumed to be invertible with roots inside the unit circle and the asymptotic covariance matrix  $\lim_{t \rightarrow \infty} Ew_t w_t' = M_w$  is positive definite and solves the Lyapunov equation:

$$M_w = FM_w F' + \Sigma_e.$$

$E_t^* y_{t+1}$  denotes the expectations held by private agents, which under learning can differ from rational expectations (RE). This model has a unique RE solution of the form  $y_t = \bar{a} + \bar{b}w_t$ . This solution is often called the “fundamentals” or minimal state variable (MSV) solution.<sup>3</sup>

Under learning agents have a “perceived law of motion” (PLM) of the form  $y_t = a + bw_t$  and estimate the parameters  $a$  and  $b$  econometrically. Thus at time  $t$  agents have the estimated PLM:

$$E_t^* y_t = a_t + b_t w_t,$$

which implies the forecast function

$$E_t^* y_{t+1} = a_t + b_t F w_t.$$

To simplify the analysis we have assumed that  $F$  is known, but it would be straightforward to allow  $F$  also to be estimated, and our results would be in essence unaffected. Any given PLM induces an “actual law of motion” (ALM) that gives the temporary equilibrium value of  $y_t$ . This is obtained by substituting  $E_t^* y_{t+1}$  into (1). For PLM estimates  $a_t, b_t$  we obtain

$$y_t = \alpha + Aa_t + (Ab_t F + B)w_t + \eta_t.$$

---

<sup>3</sup>For simplicity, our main points are made within the purely forward-looking model (1). Appendix C shows how to extend the analysis of convergence of GSG learning to models with lagged endogenous variables.

Introducing the notation  $z'_t = (\mathbf{1}', w'_t)$  for the state variables, where  $\mathbf{1}$  refers to a column vector of 1s, and the notation  $\varphi'_t = (a_t, b_t)$  for the parameters, we can summarize the PLM at  $t$  as  $y_t = \varphi'_t z'_t$  and the ALM at  $t$  as  $y_t = T(\varphi_t)' z'_t + \eta_t$ , where

$$T(\varphi)' = (\alpha + Aa, AbF + B). \quad (2)$$

The MSV RE solution is given by the fixed point of  $T$ , i.e.  $\bar{\varphi}' = (\bar{a}, \bar{b})$ , where  $\bar{a} = (I - A)^{-1}\alpha$  and  $\bar{b} = A\bar{a}F + B$ .

## 2.2 GSG Algorithm

We begin by introducing the Generalized SG algorithm for estimating and updating  $a_t, b_t$ . This is given by:

$$\varphi_t = \varphi_{t-1} + \gamma_t \Gamma z_{t-1} (y_{t-1} - \varphi'_{t-1} z_{t-1})', \quad (3)$$

where in the decreasing gain case  $\gamma_t \rightarrow 0$  as  $t \rightarrow \infty$  and in the constant gain case  $\gamma_t = \gamma > 0$  for all  $t$ . The usual assumption in the decreasing gain case is that  $\gamma_t = 1/t$ . Here  $\Gamma$  is a positive definite weighting matrix. The standard SG algorithm is the special case  $\Gamma = I$ . It is also worth noting that setting  $\Gamma = M_z^{-1}$ , where  $M_z = \lim_{t \rightarrow \infty} E z_t z'_t = \text{diag}(I, M_w)$  delivers an algorithm that is asymptotically equivalent to RLS, as discussed below. Here “diag” denotes a block diagonal matrix.

We focus in this section on the decreasing gain case in which convergence to RE is possible over time. Later, in Section 4 we examine constant gain versions. There we note that the stability conditions are the same, but the sense of convergence differs. Substituting in the ALM we can write:

$$\varphi_t = \varphi_{t-1} + \gamma_t \Gamma z_{t-1} [z'_{t-1} (T(\varphi_{t-1}) - \varphi_{t-1}) + \eta'_t], \quad (4)$$

which is formally a stochastic approximation or stochastic recursive algorithm (SRA). The general techniques for analyzing SRA's yield conditions for convergence of  $\varphi_t$  to a fixed value  $\bar{\varphi}$  as  $t \rightarrow \infty$  under GSG learning. Convergence of  $\varphi_t$  depends, in particular, on the properties of the mapping  $T(\varphi)$ . A well-known method for obtaining the convergence conditions is based on a study of stability of an ordinary differential equation that is associated with the SRA.<sup>4</sup>

For the system (4) it is straightforward to verify that  $\varphi_t \rightarrow \bar{\varphi}$  globally with probability one if  $\bar{\varphi}$  is a locally stable equilibrium of the associated differential equation

$$\frac{d\varphi}{d\tau} = \Gamma M_z (T(\varphi) - \varphi), \quad (5)$$

where  $\tau$  is notional or virtual time.<sup>5</sup> Since both  $\Gamma$  and  $M_z$  are positive definite, their product is nonsingular, which implies that the only equilibrium of the differential equation is the REE  $\bar{\varphi}$ . Standard results for SRAs also imply that  $\bar{\varphi}$  is the only possible point of convergence.

<sup>4</sup>See (Evans and Honkapohja 2001), especially Chapters 6 and 7 for a discussion of SRAs and the study of their convergence properties.

<sup>5</sup>Global convergence applies because equation (9) is linear here. Thus  $\bar{\varphi}$  is in fact globally asymptotically stable if it is locally so.

Local stability conditions for  $\varphi(\tau) \rightarrow \bar{\varphi}$  are given by the linearization of the matrix differential equation (5), giving

$$\frac{d \text{vec } \varphi'}{d\tau} = (\Gamma M_z \otimes I)(DT' - I) \text{vec } \varphi',$$

where “vec” refers to the vectorization of a matrix and  $DT'$  is the  $n(k+1) \times n(k+1)$  Jacobian matrix, of the vectorized  $T'$  map given in (2), evaluated at the fixed point  $\bar{\varphi}$ .<sup>6</sup> Local stability of the differential equation requires that all eigenvalues of

$$(\Gamma M_z \otimes I)(DT' - I) \tag{6}$$

have negative real parts.

We now turn to the classic SG and RLS algorithms and examine the connections between their respective convergence conditions. Later we will come back to analysis of the GSG algorithm.

## 2.3 Classic SG and RLS Algorithms

The classic SG algorithm sets  $\Gamma = I$ . The associated differential equation (5) then takes the explicit form:

$$\frac{d\varphi'}{d\tau} = \begin{pmatrix} \alpha + (A - I)a & (AbF - b + B)M_w \end{pmatrix} \tag{7}$$

since

$$M_z = \begin{pmatrix} I & 0 \\ 0 & M_w \end{pmatrix}.$$

The stability conditions for  $\varphi$  thus imply conditions for  $(a, b)$ . For the  $a$  component we get the stability condition that the eigenvalues of  $A - I$  must have negative real parts. Vectorizing the equation for matrix  $b$  yields

$$\begin{aligned} \frac{d \text{vec}(b)}{d\tau} &= (M_w \otimes I) \text{vec}(AbF - b + B) = \\ &= (M_w \otimes I)[((F' \otimes A) - I) \text{vec}(b) + \text{vec}(B)]. \end{aligned}$$

The relevant coefficient matrix is  $(M_w \otimes I)((F' \otimes A) - I)$  and local stability requires that all of its eigenvalues must have negative real parts. In what follows we say that a matrix is *stable* if all of its eigenvalues have negative real parts.

**SG-stability conditions:**  $A - I$  and  $(M_w \otimes I)((F' \otimes A) - I)$  are stable matrices.

Note that the condition that  $A - I$  is stable is equivalent to the condition that every eigenvalue of  $A$  has real part less than unity.

---

<sup>6</sup>For an  $m \times n$  matrix  $X$ ,  $\text{vec}X$  is the  $mn \times 1$  vector that stacks in order the columns of  $X$ . For the vectorization and matrix differential results see the summary in Section 5.7 of (Evans and Honkapohja 2001). For a full discussion, see (Magnus and Neudecker 1988).

In the literature on adaptive learning, the RLS algorithm has been more widely used and we briefly review its main features. The RLS algorithm is usually written in recursive form as:

$$\begin{aligned}\varphi_t &= \varphi_{t-1} + \gamma_t R_t^{-1} z_{t-1} (y_{t-1} - b'_{t-1} z_{t-1})' \\ R_t &= R_{t-1} + \gamma_t (z_{t-1} z'_{t-1} - R_{t-1}),\end{aligned}\tag{8}$$

where again the standard decreasing gain assumption is that  $\gamma_t = 1/t$ . The associated differential equation takes the form  $d\varphi/d\tau = S^{-1}M_z[T(\varphi) - \varphi]$  and  $dS/d\tau = M_z - S$ , where a timing change  $S_{t-1} = R_t$  must be made to (8) before the formal analysis.<sup>7</sup> Clearly, local convergence of RLS learning requires that the RE solution  $\bar{\varphi}$  (or  $\bar{\varphi}'$ ) be locally asymptotically stable under the differential equation:

$$\frac{d}{d\tau}\varphi' = T(\varphi)' - \varphi'.\tag{9}$$

Equation (9) has been widely studied in the literature.  $\bar{\varphi}$  is said to be *expectationally stable* (E-stable) if  $\bar{\varphi}$  is a locally asymptotically stable fixed point of (9). The conditions for E-stability are:

**E-stability conditions:**  $A - I$  and  $F' \otimes A - I$  are stable matrices.

We remark that the eigenvalues of  $F' \otimes A - I$  are  $f_k \lambda_i - 1$ , where  $f_k$  and  $\lambda_i$  are eigenvalues of  $F$  and  $A$ , respectively. This follows since the eigenvalues of the Kronecker product of two matrices consist of the products of the eigenvalues of each matrix.

Note that in contrast to SG-stability, the E-stability conditions do not depend on  $M_w$ . Furthermore, using the preceding remark, with a strengthening of the condition on  $A$  it is possible to obtain conditions that are also independent of  $F$ . Specifically, if every eigenvalue of  $A$  has modulus less than one then E-stability holds for all permissible  $F$ .

## 2.4 Relationships between E-stability and SG-stability

The relationship between SG and E-stability in specific models has been discussed in the previous literature. The two conditions are identical in multivariate cobweb-type models in which past expectations of current variables, but not expectations of future variables, appear in the structural model; see (Evans and Honkapohja 1998). In models with dependence on expectations of future values of the endogenous variables, an exact correspondence between E-stability and convergence of SG learning no longer holds. This was conjectured by (Barucci and Landi 1997), though they did not give any explicit examples. (Heinemann 2000) suggested an example, though it was in the context of a non-fundamental REE. Recently, (Giannitsarou 2005) has provided an economic example with lagged endogenous variables in which E-stability of the fundamental REE does not imply convergence of SG learning.

---

<sup>7</sup>See pp. 232-235 of (Evans and Honkapohja 2001) for more details.

Thus, SG- and E-stability conditions are not always the same. In fact, for our forward-looking model (1), simple examples show that neither implies the other. See Appendix A for numerical examples showing that the phenomenon can arise in purely forward-looking models with a single endogenous variable and two exogenous variables. These are the simplest examples that can be provided, since with a single exogenous variable it is immediate that E-stability and SG-stability are equivalent.

### 2.4.1 Stability Conditions

We next consider various cases under which E-stability implies SG-stability. The first results give additional conditions that are sufficient for SG-stability for all admissible  $M_w$ . We first define two stability concepts.<sup>8</sup>

**Definition 1** *A matrix  $C$  is H-stable if all the eigenvalues of  $HC$  have negative real parts whenever  $H$  is a positive definite matrix.*

**Definition 2** *A matrix  $C$  is D-stable if all the eigenvalues of  $DC$  have negative real parts whenever  $D$  is a positive diagonal matrix.*

The definition makes clear that a sufficient condition for convergence of SG learning is:

**Proposition 3** *Assume E-stability and assume further that the matrix  $(F' \otimes A) - I$  is H-stable. Then SG-stability holds for all admissible  $M_w$ .*

The property of H-stability is quite restrictive. A sufficient condition for H-stability of a matrix  $C$  is that  $C$  is negative quasi-definite, i.e. that  $C + C'$  is negative definite, i.e. has negative eigenvalues.<sup>9</sup> Note that if  $M_w$  is diagonal, then the set of sufficient conditions is that (i)  $A - I$  is a stable matrix and (ii)  $(F' \otimes A) - I$  is D-stable. There exist various necessary or sufficient condition for D-stability, but a full characterization is apparently not available (this is in contrast to H-stability).<sup>10</sup>

**Corollary 4** *Assume E-stability. If in addition  $F$  is symmetric with positive eigenvalues and  $A - I$  is negative quasi-definite then SG-stability holds for all admissible  $M_w$ .*

Proofs of this and other results are in Appendix B. The conditions in the Corollary can be convenient to apply, but they are much stronger than E-stability of the MSV REE.

The next case of uncorrelated exogenous variables arises in applications, as illustrated below:

---

<sup>8</sup>(Honkapohja and Mitra 2005) show that H-stability provides sufficient stability conditions when there is structural and learning heterogeneity.

<sup>9</sup>See, for example, (Arrow and McManus 1958). They refer to H-stability as S-stability. Necessary and sufficient conditions for H-stability are given in (Carlson 1968).

<sup>10</sup>See (Arrow 1974) and (Johnson 1974).

**Proposition 5** *When  $F$  and  $\Sigma_e$  are diagonal, E-stability and SG-stability are equivalent.*

There are some further special cases in which E-stability guarantees convergence of SG learning. For models with a scalar endogenous variable we have:

**Proposition 6** *Assume that  $n = 1$  with  $|A| < 1$ . If the largest singular value of  $F$  is not greater than one, then E-stability implies SG-stability.*

We recall that the largest singular value of  $F$  is equal to the largest eigenvalue of  $FF'$ .

#### 2.4.2 Economic Example

As an example we consider the bivariate New Keynesian model of monetary policy, which is widely used in current discussions of monetary policy.<sup>11</sup> The key equations of the model take the form:

$$x_t = c_x + E_t^* x_{t+1} - \sigma^{-1}(r_t - E_t^* \pi_{t+1}) + g_t, \quad (10)$$

$$\pi_t = c_\pi + \kappa x_t + \mathcal{B} E_t^* \pi_{t+1} + u_t. \quad (11)$$

Here  $x_t$  is the output gap,  $\pi_t$  is the inflation rate and  $r_t$  is the nominal interest rate. The parameters  $\sigma, \kappa > 0$  and  $0 < \mathcal{B} < 1$ .  $c_x$  and  $c_\pi$  are intercepts, which are from the log-linearization of the exact model.  $w_t' = (g_t, u_t)$  consists of observable shocks to the output gap and inflation, respectively. The stochastic process for  $w_t$  has the form given in (1). The first equation is the IS curve that comes from the Euler equation for consumer optimality and the second equation is the forward-looking Phillips curve based on Calvo price stickiness.

The model is completed by specification of an interest rate rule. A wide variety of different rules have been studied in the literature.<sup>12</sup> One possibility is the standard Taylor rule:

$$r_t = c_r + \phi_\pi \pi_t + \phi_x x_t, \quad (12)$$

where  $c_r$  denotes an intercept. The parameters satisfy  $\phi_\pi, \phi_x > 0$ . (Bullard and Mitra 2002) show that the E-stability condition under the standard Taylor rule is

$$\kappa(\phi_\pi - 1) + (1 - \mathcal{B})\phi_x > 0. \quad (13)$$

Alternatively, (Evans and Honkapohja 2003b) consider optimal discretionary policy and show that the expectations-based interest rate rule

$$r_t = c_r + \left(1 + \frac{\sigma \kappa \mathcal{B}}{\alpha + \kappa^2}\right) E_t^* \pi_{t+1} + \sigma E_t^* x_{t+1} + \sigma g_t + \frac{\sigma \kappa}{\alpha + \kappa^2} u_t, \quad (14)$$

where  $\alpha$  is the weight on output gap in a quadratic loss function of the policy-maker and  $c_r$  is an intercept, always leads to E-stability of the REE. If the two shocks  $g_t$  and  $u_t$  in model (10)-(11) are uncorrelated, Proposition 5 applies.

<sup>11</sup>See e.g. (Clarida, Gali, and Gertler 1999), (Svensson 2003), and (Woodford 2003) for details and analysis.

<sup>12</sup>The issue of stability under learning has been examined by (Bullard and Mitra 2002) and (Evans and Honkapohja 2003b) among others. (Evans and Honkapohja 2003a) review the literature.

**Proposition 7** *Assume that  $g_t$  and  $u_t$  are independent stationary AR(1) processes.*

- (i) *Under the Taylor rule (12) the REE is SG-stable if condition (13) holds, and*
- (ii) *The REE is SG-stable when optimal discretionary policy employs the expectations-based rule (14).*

We next consider SG-stability further under more general assumptions about the shocks  $g_t$  and  $u_t$ . For brevity, we restrict attention to the case where the policy-maker employs the Taylor rule (12). Introducing the notation  $y_t = (x_t, \pi_t)'$ , equations (10), (11) and (12) can be combined to yield the bivariate system of the form (1) with:

$$A = \frac{1}{\sigma + \phi_x + \kappa\phi_\pi} \begin{pmatrix} \sigma & 1 - \beta\phi_\pi \\ \kappa\sigma & \kappa + \beta(\sigma + \phi_x) \end{pmatrix}.$$

We omit the explicit form of  $B$  as it does not affect the stability conditions.

We give numerical examples of the above results using the calibration of the model due to (Rotemberg and Woodford 1997) and widely employed in (Woodford 2003).

**Calibration:**  $\beta = 0.99, \sigma = 0.157, \kappa = 0.024$ .

Suppose first that the policy parameters take on values  $\phi_\pi = 1.05$  and  $\phi_x = 0.2$ . The E-stability conditions on  $A - I$  hold since both eigenvalues of  $A - I$  are in the interval  $(-1, 0)$ . Furthermore,  $A - I$  is negative quasi-definite as the eigenvalues of  $(A - I) + (A - I)'$  are  $-1.186$  and  $-0.0174$ . Assuming that the coefficient matrix  $F$  of the vector of shocks in (1) is symmetric with positive eigenvalues, the sufficient conditions given in Corollary 4 are met and therefore SG-learning is convergent under the specified policy parameter values.

As a second numerical example, we postulate  $\phi_\pi = 1.1$  and  $\phi_x = 0.1$ . E-stability on  $A - I$  continues to hold, as the eigenvalues of  $A - I$  are in the interval  $(-1, 0)$ , but  $A - I$  fails to be negative quasi-definite since the eigenvalues of  $(A - I) + (A - I)'$  are  $-0.987$  and  $0.0599$ . However, consider the following specification of the exogenous shocks:

$$F = \begin{pmatrix} 0.99 & 0 \\ 0 & 0.98 \end{pmatrix}, \Sigma_e = \begin{pmatrix} 0.1 & 0.05 \\ 0.05 & 0.049 \end{pmatrix}.$$

Directly checking the SG stability conditions, we find that the eigenvalues of  $(M_w \otimes I)((F' \otimes A) - I)$  are  $-2.505$ ,  $-0.268$ ,  $-0.218$ , and  $-0.0273$ . Thus SG learning is convergent for this exogenous process even though the sufficient condition of Corollary 4 fails.

These numerical examples illustrate both the applicability and limitations of the preceding stability conditions.

### 3 SG, RLS and Scaling Invariance

#### 3.1 Invariance to Scaling

SG algorithms suffer from a disadvantage relative to RLS, which has not received attention. The SG algorithm is not scale invariant, and thus the resulting estimates are affected by the choice of units. This is demonstrated as follows.

For simplicity, we consider a univariate case. Suppose we are estimating the regression model:

$$y_t = \beta' z_t + \eta_t$$

by least squares, where  $\beta$  and  $z_t$  are  $p \times 1$  column vectors. Our discussion here will initially be in terms of the standard (non-self-referential) regression model, since the point holds generally, but it also applies to the model (1) with learning. The RLS estimate using data through  $t-1$  is given by (8). Suppose we now change units of the regressors so that  $\tilde{z}_t = Dz_t$ , where  $D = \text{diag}(k_1, \dots, k_p)$ . Here  $\text{diag}$  denotes a diagonal matrix and we assume  $k_i > 0$ . Then:

$$y_t = \tilde{\beta}' \tilde{z}_t + \eta_t,$$

where  $\tilde{\beta} = D^{-1}\beta$ . Let  $\tilde{\varphi}_t$  be the RLS estimate of  $\tilde{\beta}$  based on a regression of  $y_t$  on  $\tilde{z}_t$ . Then RLS is *scale invariant* in the sense that:

$$\tilde{\varphi}_t = D^{-1}\varphi_t.$$

To see this, pre-multiply the RLS equation for  $\varphi_t$  by  $D^{-1}$  and for the  $R_t$  equation, premultiply by  $D$  and postmultiply by  $D' = D$ . Defining  $\tilde{R}_t = DR_tD'$  we get:

$$\begin{aligned} \tilde{\varphi}_t &= \tilde{\varphi}_{t-1} + \gamma_t \tilde{R}_t^{-1} \tilde{z}_{t-1} (y_{t-1} - \tilde{\varphi}_{t-1}' \tilde{z}_{t-1}) \\ \tilde{R}_t &= \tilde{R}_{t-1} + \gamma_t (\tilde{z}_{t-1} \tilde{z}_{t-1}' - \tilde{R}_{t-1}). \end{aligned}$$

But this is exactly RLS applied to a regression of  $y_t$  on  $\tilde{z}_t$ .

In contrast, SG estimation is not scale invariant. The SG algorithm for a regression of  $y_t$  on  $z_t$  is:

$$\varphi_t = \varphi_{t-1} + \gamma_t z_{t-1} (y_{t-1} - \varphi_{t-1}' z_{t-1}).$$

Multiplying through by  $D^{-1}$  we get that  $\tilde{\varphi}_t = D^{-1}\varphi_t$  satisfies:

$$\tilde{\varphi}_t = \tilde{\varphi}_{t-1} + \gamma_t D^{-2} \tilde{z}_{t-1} (y_{t-1} - \tilde{\varphi}_{t-1}' \tilde{z}_{t-1}). \quad (15)$$

But SG estimation based on a regression of  $y_t$  on  $\tilde{z}_t$  is instead:

$$\hat{\varphi}_t = \hat{\varphi}_{t-1} + \gamma_t \tilde{z}_{t-1} (y_{t-1} - \hat{\varphi}_{t-1}' \tilde{z}_{t-1}),$$

and clearly  $\tilde{\varphi}_t \neq \hat{\varphi}_t$ .

Note that the same argument applies to transformations of variable  $\tilde{z}_t = Dz_t$  for  $D$  positive definite: RLS is invariant to such transformations while SG is not.

### 3.2 The Forward Expectations Model and the Cholesky Decomposition

Return now to the multivariate forward looking model (1) with SG learning. Consider a change of variables based on the Cholesky decomposition:

$$M_w = QQ'.$$

This is always possible for a positive definite matrix, resulting in a matrix  $Q$  which is triangular and nonsingular.<sup>13</sup> Letting  $L = Q^{-1}$  we have

$$LM_wL' = I.$$

Transforming independent variables to

$$\tilde{w}_t = Lw_t,$$

the RE solution becomes  $y_t = \bar{a} + \hat{b}\tilde{w}_t$  where  $\hat{b} = \bar{b}L^{-1}$ .

Under SG learning with the transformed independent variables  $\tilde{w}_t$  the PLM becomes

$$\begin{aligned} y_t &= a + \tilde{b}\tilde{w}_t + \eta_t, \text{ where} \\ \tilde{w}_t &= \tilde{F}\tilde{w}_{t-1} + \tilde{e}_t, \text{ with } \tilde{F} = LFL^{-1} \text{ and } \tilde{e}_t = Le_t. \end{aligned}$$

Note that

$$E\tilde{w}_t\tilde{w}_t' = I.$$

Thus we have transformed the independent explanatory variables to orthogonal variables with unit variances. Clearly  $\tilde{w}_t$  has the same information content as  $w_t$ , and thus they are equally good for forecasting. Furthermore, note that  $\tilde{F}$  has the same eigenvalues as  $F$  because  $F$  and  $\tilde{F}$  are similar matrices.

The SG-stability conditions for the transformed specification are that  $A - I$  and  $(M_{\tilde{w}} \otimes I)((\tilde{F}' \otimes A) - I)$  are stable matrices. But  $M_{\tilde{w}} \otimes I = I$  and hence the SG-stability conditions are that  $A - I$  and  $\tilde{F}' \otimes A - I$  are stable matrices. Since  $\tilde{F}$  and  $F$  have the same eigenvalues it follows that the SG-stability conditions of the transformed model reduce precisely to the E-stability conditions. We have therefore shown:

**Proposition 8** *There exists a transformation of variables  $\tilde{w}_t = Lw_t$ , with  $L$  positive definite, under which SG-stability is equivalent to E-stability.*

The result indicates that for this model any deviation of the stability condition for SG learning from E-stability is simply a reflection of how the independent variables are measured and the sensitivity of SG estimation to scaling. We remark that with the model extended to include lagged endogenous variables, this result continues to hold, but the matrix  $L$  becomes a function of the second moment matrix of the equilibrium joint process  $(w_t, y_{t-1})$ .<sup>14</sup>

### 3.3 GSG Algorithms and Scaling

The preceding discussion suggests considering a more general class of SG-type algorithms. Note that a change of measurement units for  $z_t$  in the standard SG algorithm led to (15),

<sup>13</sup>For the Cholesky decomposition see, e.g., (Hamilton 1994) pp. 91-2.

<sup>14</sup>In a real-time learning rule with lagged endogenous variables, agents could update the Cholesky transformation each period as part of their recursive estimation.

which is a special case of the general GSG algorithm (3). More generally, a change of measurements of  $z_t$  in a generalized SG algorithm will yield another member of this class of algorithms. For example, as previously noted, choosing  $\Gamma = M_z^{-1}$  leads to an algorithm which is asymptotically approximately the same as the RLS algorithm: the estimated moment matrix  $R_t$  of RLS is being replaced by its limiting value  $M_z$ .

In the context of the forward-looking model (1) the condition for convergence to the RE solution of generalized SG learning is that the eigenvalues of the matrix (6) have negative real parts. Equivalently, the stability condition is that

$$(\Gamma M_z \otimes I) \begin{pmatrix} I \otimes A - I & 0 \\ 0 & F' \otimes A - I \end{pmatrix} \quad (16)$$

have negative real parts.

In the case of  $\Gamma = M_z^{-1}$  the condition becomes equivalent to E-stability of the RE solution. More generally, the analogues of the results on the connections between E-stability and SG-stability given in Section 2 hold provided  $\Gamma$  and  $M_z$  commute.<sup>15</sup> In particular if the independent variables  $z_t$  are transformed, as in the previous section, so that  $M_z = I$  and  $F$  is replaced by  $\tilde{F}$ , then the analogues of the Section 2 results hold for the generalized SG algorithm provided  $\Gamma$  takes the block diagonal form  $\text{diag}(I, \hat{\Gamma})$  where  $\hat{\Gamma}$  is  $k \times k$ .

The formulation of the class of generalized SG algorithms raises the question of the appropriate choice of  $\Gamma$ . We now turn to this issue.

## 4 GSG Learning with Constant Gain

As discussed in the introduction, most studies on adaptive learning in economics have considered the RLS learning rule, which has a well-known statistical motivation. This learning rule minimizes the sum of squared residuals and is a maximum likelihood estimator when the shocks are normally distributed and the model is correctly specified. Some recent studies have also focused on a so-called constant gain version of recursive least squares, in which the gain sequence  $\gamma_t$  is taken to be a small constant  $\gamma > 0$ .<sup>16</sup> This rule can be derived by minimizing the discounted sum of squared residuals, with the discount factor  $1 - \gamma$ . As discussed below, this rule is also closely connected to the Kalman filter when the true coefficients follow a random walk, and thus it has an approximate Bayesian interpretation.

As also noted in the Introduction, in the applied literature constant gain versions of SG learning are also in use. As mentioned above, SG requires no matrix inversion and hence has lower computational cost and deals better with singularities or near-singularities. In this

<sup>15</sup>The eigenvalues of  $\Gamma M_w$  are positive and real, but  $\Gamma M_w$  need not in general be symmetric.

<sup>16</sup>For constant gain learning the basic references are (Evans and Honkapohja 1993), (Sargent 1999), Chapter 14 of (Evans and Honkapohja 2001), (Williams 2001), and (Cho, Williams, and Sargent 2002). Constant gain rules are increasingly used in applied work, for example see (Bichi and Marimon 2001), (Bullard and Cho 2005), (Bullard and Duffy 2004), (Bullard and Eusepi 2005), (Cho and Kasa 2002), (Evans and Honkapohja 2005), (Kasa 2004), (Marcet and Nicolini 2003), (McGough 2005), (Orphanides and Williams 2005a) and (Orphanides and Williams 2005b).

section we provide more formal arguments justifying constant gain GSG rules when agents allow for parameter drift and use Bayesian updating. The gain matrix  $\Gamma$  is then tied to the particular form of the parameter drift.

We also show that constant gain GSG learning has a dramatically different justification as a robust optimal learning rule. In particular, it provides good guaranteed estimation performance in misspecified models. Moreover, the GSG rule is also optimal if agents are “risk-sensitive,” having greater risk aversion of a particular form: the GSG rule minimizes the expected exponential of the sum of squared errors. We then turn to some simple simulation exercises which document the robustness of SG in practice.

## 4.1 Bayesian Interpretation of GSG

We follow (Sargent and Williams 2005), who consider Kalman filter estimation when parameter drift is modeled according to a random walk hypermodel. With a specific prior on the covariance matrix for the parameter drift, it is possible to obtain an algorithm closely related to constant gain RLS estimation. We show here that an alternative prior on the form of parameter drift leads to a constant gain GSG algorithm.

### 4.1.1 Derivation of GSG

In particular, we suppose that an agent believes that the data are generated by the drifting coefficients model:

$$y_t = \beta'_{t-1} z_t + \eta_t \quad (17)$$

$$\beta_t = \beta_{t-1} + \Lambda_t \quad (18)$$

where  $\eta$  and  $\Lambda$  are viewed as mean zero Gaussian shocks with  $E\eta_t^2 = \sigma^2$  and  $\text{cov}(\Lambda_t) = V \ll \sigma^2 I$ . Again, our initial discussion is presented in terms of a standard time-varying parameter regression model, but it is also applicable to the self-referential model.

In this section we assume that  $y_t$  is a scalar. This is merely for notational simplicity and does not affect the analysis. The agent’s estimator is  $\varphi_t \equiv \hat{\beta}_{t|t-1}$ , the optimal estimate of  $\beta_t$  conditional on information up to date  $t - 1$ . It is well known that the (Bayes) optimal estimates in this linear model are provided by the Kalman filter. The Kalman filtering equations are:

$$\varphi_{t+1} = \varphi_t + \frac{P_t}{1 + z'_t P_t z_t} z_t (y_t - \varphi'_t z_t) \quad (19)$$

$$P_{t+1} = P_t - \frac{P_t z_t z'_t P_t}{1 + z'_t P_t z_t} + \sigma^{-2} V. \quad (20)$$

Here  $\text{cov}(\varphi_t - \beta_t) \equiv \sigma^2 P_t$ .

(Benveniste, Metivier, and Priouret 1990) note that for large  $t$  (20) is well approximated by:

$$P_{t+1} = P_t - P_t M_z P_t + \sigma^{-2} V,$$

where  $M_z = Ez_t z_t'$ .<sup>17</sup> Using this approximation (and assuming  $1/(1 + z_t' P_t z_t) \approx 1$ ), the Kalman filter equations simplify to:

$$\varphi_{t+1} = \varphi_t + P_t z_t (y_t - \varphi_t' z_t) \quad (21)$$

$$P_{t+1} = P_t - P_t M_z P_t + \sigma^{-2} V. \quad (22)$$

Furthermore, we have:

**Lemma 9** *Under (22),  $P_t$  locally converges to the unique positive definite matrix  $P$  that solves the equation:*

$$P M_z P = \sigma^{-2} V. \quad (23)$$

The prior belief  $V$  on the form of the parameter drift influences the learning rule, as it influences the speed and direction along which parameters should be updated. In particular, suppose we normalize  $V$  by writing:

$$V = \gamma^2 \sigma^2 \Omega,$$

where  $\gamma$  controls the overall speed of the parameter drift and  $\Omega$  specifies the direction of the drift (e.g. we might normalize by setting  $\det(\Omega) = 1$ ). Since the  $P_t$  recursion (22) converges, the limit  $P$  satisfies:

$$P M_z P = \gamma^2 \Omega \text{ or } (\gamma^{-1} P) M_z (\gamma^{-1} P) = \Omega.$$

Therefore, letting  $\Gamma = \gamma^{-1} P$  we have:

$$\Gamma M_z \Gamma = \Omega = \gamma^{-2} \sigma^{-2} V, \quad (24)$$

and asymptotically the parameter estimates satisfy:

$$\varphi_{t+1} = \varphi_t + \gamma \Gamma z_t (y_t - \varphi_t' z_t), \quad (25)$$

which is the constant gain GSG algorithm. Note that (25) has the same asymptotic behavior as (19)-(20) under the assumed form of  $V$ , since  $P_t$  converges to  $\gamma \Gamma$ , although the transient responses from arbitrary initial conditions may differ.

Thus the choice of the gain matrix  $\Gamma$  in the GSG rule is closely tied to the prior  $V$  on the parameter drift. (Sargent and Williams 2005) apply results from (Benveniste, Metivier, and Priouret 1990) to show that if  $V = \gamma^2 \sigma^2 M_z^{-1}$ , then the Kalman filter is closely related to a constant gain RLS algorithm, as  $\Gamma = M_z^{-1}$ .<sup>18</sup> Alternatively, suppose that instead of being proportional to the ratio of the observation noise variance to the covariance matrix of the regressors ( $M_z$ ), the parameter innovation covariance matrix is proportional to the product of the two:  $V = \gamma^2 \sigma^2 M_z$ . In this case  $\Gamma = I$  and the classic SG rule results. More generally, the prior on  $V$  will lead to a particular choice of the optimal matrix  $\Gamma$  in the GSG rule.

<sup>17</sup>We note that in the model with lagged endogenous variables, considered in Appendix C,  $M$  would depend on  $b$  in equations (22) and (23).

<sup>18</sup>The two rules have the same limits, but the transient phases differ. In the application of (Sargent and Williams 2005), the Kalman filter converges faster.

### 4.1.2 Convergence of Bayesian Learning

Returning to the economic model (1), it follows that if agents are updating their estimates according to the approximate Bayesian learning rule (21)-(22) then for small  $\gamma$  local stability is determined by (16). The convergence conditions for corresponding decreasing gain algorithms carry over to the versions with small constant gain, though convergence is now weak convergence to a stochastic process near the REE, as discussed e.g. in Chapters 7 and 14 of (Evans and Honkapohja 2001) and (Cho, Williams, and Sargent 2002). We thus have the following result.

**Proposition 10** *(i) The REE  $\bar{\varphi}$  is locally stable for sufficiently small  $\gamma > 0$  under the approximate Bayesian learning rule (21)-(22) if (16), with  $\Gamma$  defined in (24), is a stable matrix.*

*(ii) The approximate Bayesian learning rule (21)-(22) is invariant to a change of variables to  $\tilde{z}_t = Dz_t$  for any positive definite matrix  $D$ .*

The invariance of Bayesian learning to a transformation of variables also allows us to make use of some of the stability results from Section 2.

Recalling the results of Section 2.4, one might ask whether there are restrictions on the economic model that guarantee local stability of Bayesian learning for all priors on the parameter drift that are sufficiently small. Combining the idea of Proposition 3 with Proposition 10 we obtain the following result.

**Proposition 11** *Let  $LM_wL' = I$  and let  $\tilde{F} = LFL^{-1}$ . If*

$$\begin{pmatrix} I \otimes A - I & 0 \\ 0 & \tilde{F}' \otimes A - I \end{pmatrix} \quad (26)$$

*is  $H$ -stable then under the approximate Bayesian learning rule (21)-(22), the REE  $\bar{\varphi}$  is locally stable for all priors  $V$  on the parameter drift, where  $V$  is sufficiently small.*

Note that Proposition 11 gives a condition for local stability, for all sufficiently small priors  $V$ , that holds regardless of how the variables  $w_t$  are measured. Finally, we remark that the condition given is a strengthening of E-stability, since the latter is equivalent to stability of the matrix (26).

## 4.2 Constant Gain GSG Learning and Robustness

In the previous section, we showed that the GSG rule is an (approximate) optimal predictor for a particular model of parameter variation. However, the form of parameter variation is quite particular, and its appropriateness in any given application is an open issue. In this section, we provide a motivation that is in some ways more general in that it encompasses a range of different model specifications. More particularly, if agents do not know the correct specification of the model that they estimate, then they may want to choose a rule which

performs well across a range of alternatives. Here we show that the GSG rule is such a robust optimal prediction rule. Our results here follow (Hassibi, Sayed, and Kailath 1996). Robust control methods have recently been applied to a range of economic problems (see (Hansen and Sargent 2004) for an overview), and it is interesting to see that the SG rule has a robust interpretation.

In particular, we suppose that agents believe now that the true coefficients are constant over time, but they are uncertain about the data generating process. We represent this by a variation on (17)-(18) which now takes the form:

$$\begin{aligned} y_t &= \beta'_{t-1} z_t + \eta_t \\ \beta_t &= \beta_{t-1}, \quad \beta_0 = \beta. \end{aligned} \tag{27}$$

But now, instead of the  $\eta_t$  shocks having a Gaussian distribution, agents treat  $\eta_t$  as an approximation error without a specified probability distribution. The  $\eta_t$  shocks are introduced as a means of capturing the possible misspecification of the model, and they may be both autocorrelated and correlated with the state  $z_t$ .

This set-up is particularly interesting in the context of the self-referential model with learning, since under robust estimation, agents explicitly allow for possible misspecification. This contrasts, for example, with standard least-squares learning formulations in which agents ignore a transitory misspecification. There agents assume that the true regression parameters are constant over time, while in reality, under learning, the parameters are time-varying, though they converge over time to the (constant) REE values.

The key assumption in (27) is that agents do not have a (unique) prior probability distribution over these shocks, and so cannot form expectations over them. Agents choose a sequence of predictors  $\varphi_t$  to minimize the prediction errors:

$$e_t = \beta' z_t - \varphi'_{t-1} z_t, \tag{28}$$

but they acknowledge the potential misspecification error. In particular, agents treat (27) with  $\eta_t \equiv 0$  as a benchmark model, but consider a set of perturbations in a neighborhood of this model. As they cannot evaluate the likelihood of potential perturbations, they guard against the worst case in the set of possibilities.

In particular, instead of minimizing the expected squared errors as in the Kalman filter case, they now solve a minimax problem. At date 0 agents have an initial estimate  $\varphi_{-1}$  of the true value  $\beta$ , with prior precision  $(\gamma\Gamma)^{-1}$ , where  $\Gamma$  is a symmetric, positive definite, nonsingular matrix and  $\gamma > 0$ . (Our use of the same notation as above is not coincidental.) Then the agent's problem is:

$$\min_{\{\varphi_s\}} \max_{\{\eta_s\}, \beta} \sum_{s=0}^t |e_s|^2$$

subject to (27), (28), and:

$$\sum_{s=0}^t |\eta_s|^2 + \frac{1}{\gamma} (\beta - \varphi_{-1})' \Gamma^{-1} (\beta - \varphi_{-1}) \leq \mu. \tag{29}$$

Here  $\mu > 0$  measures the size of the set of the alternative models, which are represented by different values of the parameter  $\beta$  and the shocks  $\eta_s$  satisfying (29). As is standard, we can convert the problem from a constrained to a penalized one by putting a Lagrange multiplier  $\theta > 0$  on the constraint (29). Then we can re-write the problem as:

$$\min_{\{\varphi_s\}} \max_{\{\eta_s\}, \beta} \sum_{s=0}^t (|e_s|^2 - \theta |\eta_s|^2) - \frac{\theta}{\gamma} (\beta - \varphi_{-1})' \Gamma^{-1} (\beta - \varphi_{-1}), \quad (30)$$

subject to (27) and (28), where we leave off the inessential term in  $\mu$ . Notice that  $\theta$  and  $\mu$  are inversely related, so we can use  $\theta$  as a measure of the size of the set of alternatives, which hence is a measure of robustness.

As  $\theta$  increases to infinity, perturbations are penalized more, and the size of the set of alternatives shrinks ( $\mu \rightarrow 0$ ) to just the baseline model. There is also a lower bound  $\underline{\theta}$  for  $\theta$  which makes the problem well-posed, and this “maximally robust” critical value is the square of the so-called  $H_\infty$  norm of the system, see (Hansen and Sargent 2004). This is the largest set of uncertainty  $\mu$  that the problem can tolerate, and also has an interpretation as what is known as an induced norm. Loosely speaking, the  $H_\infty$  norm of a system represents the maximum factor by which errors in inputs get translated into errors in outputs.

The robust estimation problem (30) is a special type of a robust control problem, and in turn is equivalent to a  $H_\infty$  estimation problem, see (Hassibi, Sayed, and Kailath 1996). The solution is known to have the following form:

$$\varphi_{t+1} = \varphi_t + K_t z_t (y_t - z_t' \varphi_t) \quad (31)$$

$$\begin{aligned} K_t &= \frac{(P_t^{-1} - \theta^{-1} z_t z_t')^{-1}}{1 + z_t' (P_t^{-1} - \theta^{-1} z_t z_t')^{-1} z_t} \\ P_{t+1}^{-1} &= P_t^{-1} + (1 - \theta^{-1}) z_t z_t', \end{aligned} \quad (32)$$

with  $P_{-1} = \gamma \Gamma$ . Note the similarities between these equations and the Kalman filter algorithm in (19)-(20) with  $V = 0$ . In particular, as  $\theta \rightarrow +\infty$  we see that they coincide.

While the robust rule collapses to the Kalman filter as the level of robustness decreases, it is more interesting in this case to consider the maximally robust learning rule with  $\theta = \underline{\theta}$ . (Hassibi, Sayed, and Kailath 1996) show that if  $\lim_{T \rightarrow \infty} \sum_{t=0}^T z_t' z_t = +\infty$  and  $\gamma \Gamma > \sup_t z_t z_t'$  (i.e. the difference is a positive definite matrix), then  $\underline{\theta} = 1$ .<sup>19</sup> Recalling our discussion above, if  $\underline{\theta}$  were greater than one, then the learning rule would magnify the effect of modeling errors on estimation errors. But here the maximally robust learning rule allows for no such magnification, and hence performs well in the face of misspecification.

Under these conditions, setting  $\theta = \underline{\theta} = 1$ , we see from (32) that  $P_t = \gamma \Gamma$  for all  $t$ . This

---

<sup>19</sup>(Hassibi, Sayed, and Kailath 1996) set  $\Gamma = I$ , but allowing for more general weighting matrices  $\Gamma$  is straightforward.

in turn implies that:

$$\begin{aligned} K_t &= \frac{(\gamma^{-1}\Gamma^{-1} - z_t z_t')^{-1}}{1 + z_t'(\gamma^{-1}\Gamma^{-1} - z_t z_t')^{-1} z_t} \\ &= \gamma\Gamma, \end{aligned}$$

where the last equality follows from the matrix inversion lemma. Thus the “gain matrix”  $K_t$  in the maximally robust learning rule (31) is constant over time, and thus this rule is the generalized constant gain stochastic gradient rule (25) from above. We summarize this discussion as follows:

**Proposition 12** *Given prior precision  $(\gamma\Gamma)^{-1}$  on  $\beta$ , the GSG algorithm (25) is the maximally robust learning rule.*

In the context of the self-referential economic model (1), of course, there is also the issue of local stability of the REE. Since, according to Proposition 12, the maximally robust learning rule takes the form of the GSG algorithm (25), the earlier stability results apply. Thus the REE is locally stable if the matrix (16) is stable.

### 4.3 Constant Gain GSG Learning and Risk Sensitivity

The previous section showed that the constant gain GSG learning rule was the (maximally) robust optimal predictor. This derivation was completely deterministic and relied on minimizing the worst case performance of the predictor over a certain class of alternative models. In this section we briefly discuss a different interpretation of these results in a stochastic setting with enhanced risk aversion, known as risk-sensitivity.<sup>20</sup> Once again we follow (Hassibi, Sayed, and Kailath 1996).

Consider again the state space model (27), where now  $\beta$  and  $\eta$  are Gaussian random variables with means  $\varphi_{-1}$  and 0 and variances  $\gamma\Gamma$  and  $I$  respectively. Then instead of minimizing the expected sum of squared errors as in the Kalman filter case, suppose that agents solve the following:

$$\min_{\{\varphi_s\}} 2\theta \log E \exp \left( \frac{1}{2\theta} \sum_{s=0}^t |e_s|^2 \right) \quad (33)$$

subject to (27) and (28). This exponential adjustment of the objective function increases risk aversion, and hence (33) is known as a risk-sensitive optimization problem (see (Whittle 1990) for a monograph on problems of this type). This can also be thought of as a particular choice of an undiscounted recursive utility objective as in (Epstein and Zin 1989).

While being motivated as an enhanced adjustment to risk instead of robustness against unknown disturbances, there are well-established results linking the solutions of robust and

---

<sup>20</sup> Applications of risk-sensitivity in economics include (Tallarini 2000) and (Anderson 2004). See (Hansen and Sargent 2004) for further discussion.

risk-sensitive control problems.<sup>21</sup> In particular, as shown by (Hassibi, Sayed, and Kailath 1996), the risk sensitive optimal filter solving (33) is identical to the robust optimal filter (31)-(32) above. Thus for the maximally robust level of  $\theta = \underline{\theta} = 1$ , or what in the risk sensitive formulation (Whittle 1990) calls “the point of the onset of neurotic breakdown,” the risk-sensitive optimal predictor is again the constant gain GSG rule (25).

## 4.4 Some Simulation Exercises

In this section we provide the results of some simulation exercises which illustrate the robustness of constant gain GSG learning. We first study a simple case of pure estimation, where there is no feedback from agents’ beliefs. We then turn to a self-referential model, focusing on the performance of SG as compared to RLS in the inflation model of (Sargent 1999) and (Cho, Williams, and Sargent 2002).

### 4.4.1 An Estimation Example

In the first exercise, we examine the relative performance of RLS and two variants of GSG learning in a pure estimation problem. We consider five different data generating processes that are variations on (17)-(18). In each case (17) holds, but we assume different forms of parameter drift. In particular, for each model we let the exogenous variables  $z_t$  be bivariate with evolution:

$$z_t = \begin{bmatrix} 0.9 & 0.5 \\ 0 & 0.5 \end{bmatrix} z_{t-1} + \begin{bmatrix} 0.5 & 0.05 \\ 0.05 & 0.5 \end{bmatrix} e_t$$

where  $e_t$  is bivariate and distributed i.i.d. standard normal. Thus in each case,  $\beta_t$  is a bivariate as well, and we consider the following forms of parameter variation:

- **M1:** Constant coefficients,  $\beta_t = \beta_{t-1} = \beta$ .
- **M2:** Time varying coefficients as in (18), with  $V = \gamma^2 \sigma^2 M_z^{-1}$ , the prior consistent with RLS and GSG with  $\Gamma = M_z^{-1}$ .
- **M3:** Time varying coefficients as in (18), with  $V = \gamma^2 \sigma^2 M_z$ , the prior consistent with classic SG with  $\Gamma = I$ .
- **M4:** Time varying coefficients as in (18), with  $V = \gamma^2 \sigma^2 I$ , the prior consistent with GSG with  $\Gamma = M_z^{-0.5}$ .
- **M5:** Time varying coefficients with a structural break. For the first 250 periods  $\beta_t$  is stationary and follows:

$$\beta_t = \begin{bmatrix} 0.90 & -0.1 \\ 0.49 & 0.5 \end{bmatrix} \beta_{t-1} + \gamma \sigma \begin{bmatrix} 1 & 1 \\ -1 & 1 \end{bmatrix} \Lambda_t,$$

---

<sup>21</sup>As noted above, these connections go back to (Jacobson 1973), but the explicit formulation given here was established by (Glover and Doyle 1988). See (Whittle 1990), (Hassibi, Sayed, and Kailath 1996), and (Hansen and Sargent 2004) for further discussions.

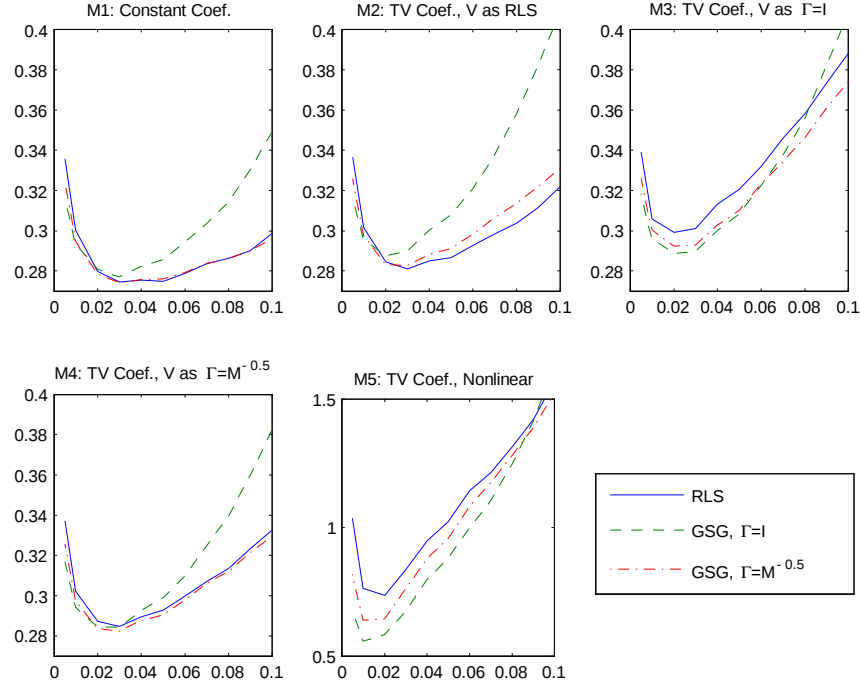


Figure 1: Mean squared forecast error under RLS and two GSG specifications. Data generated constant coefficients (M1) and different specifications of random coefficients (M2-M5). Medians of 1000 simulations of 500 periods, showing gain vs. MSE.

with  $\Lambda_t \sim N(0, I)$ . For the next 250 periods  $\beta_t$  follows a random walk as in (18) where  $\Lambda_t \sim N(0, V)$  and  $V = 5\gamma^2 I$ .

In each of the five models we study the performance of three constant gain learning rules: RLS, standard SG with  $\Gamma = I$ , and generalized SG with  $\Gamma = M_z^{-0.5}$ . Models 1-4 each rationalize different optimal learning rules, with M1 favoring a decreasing gain RLS (which we do not include as it performed badly in the other models), M2 favoring constant gain RLS (or GSG with  $\Gamma = M_z^{-1}$ ), M3 favoring classic SG with  $\Gamma = I$ , and M4 favoring GSG with  $\Gamma = M_z^{-0.5}$ . M5 is included to analyze the effect of misspecification and the performance of the rules under structural breaks.

In Figure 1 we plot the results of a small simulation study. We set the standard deviation of  $\eta_t$  to  $\sigma = 0.5$ , and set  $\beta_0 = [1, 1]'$ . We then consider 1000 simulations of 500 observations each. For each run we initialize each learning rule at a common value, drawing each component of the initial vector independently from a normal distribution with mean 1 and standard deviation 0.5. For RLS, we also set the initial value of the moment matrix to  $R_0 = M_z$ . The figure then plots the median mean squared forecast error (MSE) from the 1000 simulations for different values of the gain  $\gamma$ .

Several points are worth mention. First, in all of the models the learning rules converge rather quickly, as the faster convergence benefits of a larger gain setting are mostly offset by

the additional volatility which results in larger MSE, except at the very low end. Thus in all cases the minimum MSE occurs at a gain setting near  $\gamma = 0.02$ . For the constant coefficient M1, all of the rules perform relatively well, with relatively small fluctuations around the true value. Notice also that, as expected, the rules which are optimal for the particular model of time variation performed best, at least for a range of gain settings. Thus RLS outperforms the others in M2 (except at the very low end), classic SG does best in M3 (except at the very high end) and the GSG rule does best in M4. Thus there is some benefit to the choice of the particular  $\Gamma$ , but in many cases the differences are not large.

Some of our findings also illustrate the robustness of constant gain GSG rules. Thus we also see that for small gain settings the performance of the GSG rules are nearly as good as RLS in M2, while both GSG rules are always superior (except at the very high end) to RLS in M3. This illustrates the robustness of GSG, although admittedly in many cases the differences are not very large. More noticeably, the GSG rules perform much better than RLS in M5, where all of the learning rules are misspecified. We see that for low gain settings the classic SG rule performed quite well and was noticeably superior to RLS, particularly for gains below 0.05. The other generalized SG rule was in between RLS and SG. One other result to note is that the classic SG rule had relatively poor performance for high gain values for models M1-M4, but improved substantially as the gain falls.

Recalling that gain  $\gamma$  also influences the volatility of parameter variation, this suggests that for relatively volatile parameters, a GSG rule such as the one shown here may be the best choice in trading off robustness and performance. For cases with less volatility, the classic SG rule with small gain may be a better choice. In either case, our simulations confirm in practice our theoretical results on the robustness of GSG learning in statistical models. We next see whether this performance extends to a misspecified self-referential model.

#### 4.4.2 The *Conquest* Model

In this section we analyze the implications of different learning rules in the model of adaptive learning in monetary policy due to (Sargent 1999). Building on work by (Sims 1988), Sargent showed how a government adaptively fitting a Phillips curve model recurrently sets inflation near the optimal level, although later inflation returns to the time-consistent sub-optimal outcome. The escape dynamics leading the economy away from the equilibrium, were analyzed by (Cho, Williams, and Sargent 2002) and (Williams 2001). However, these papers focused on RLS. The impact of different prior beliefs in a Kalman filter setting was studied by (Sargent and Williams 2005), but they did not explicitly study the impacts of different types of learning rules as we do here.

Very briefly, the true economy in the model is governed by an expectations-augmented Phillips curve which relates unemployment and surprise inflation. The government controls inflation directly up to a random shock, but it bases its policy on a misspecified model which omits the role of expectations and supposes that inflation and unemployment are directly linked. Letting  $U_t$  be the unemployment rate,  $\pi_t$  the inflation rate,  $\hat{\pi}_t$  expected inflation,

and  $x_t$  the government's policy we have:

$$U_t = u_t - \theta_t(\pi_t - \hat{\pi}_t) + \sigma_1 W_{1t} \quad (34)$$

$$\pi_t = x_t + \sigma_2 W_{2t} \quad (35)$$

$$\hat{\pi}_t = x_t. \quad (36)$$

Here  $u_t$  is the natural unemployment rate,  $\theta_t$  is the true slope of the Phillips curve, and  $W_{1t}$  and  $W_{2t}$  are independent standard normal shocks. We extend the previous analyses of this model by allowing  $\beta_t = [u_t, \theta_t]'$  to be time varying. Equation (34) is an expectations-augmented natural rate Phillips curve; (35) states that the government sets inflation up to a random shock; (36) imposes rational expectations for the public.

The government does not know (34) but instead bases its policy on the estimated model:

$$U_t = b_{0t} + b_{1t}\pi_t + \eta_t, \quad (37)$$

where the estimates  $b_t$  are updated over time. The government seeks to minimize the loss  $E(U_t^2 + \pi_t^2)$  and each period re-optimizes based on its latest estimates. This leads to the policy rule:

$$x_t = x(b_t) \equiv \frac{-b_{0t}b_{1t}}{1 + b_{1t}^2}.$$

In the case of constant coefficients ( $\beta_t = [u, \theta]'$ ), (Cho, Williams, and Sargent 2002) show that there is a unique self-confirming equilibrium where  $b = [b_0, b_1]' = \bar{b} \equiv [u(1 + \theta^2), -\theta]'$ . In the notation above, we have  $z_t = [1, \pi_t]$  and so the moment matrix is:

$$M_z(b) = \begin{bmatrix} 1 & x(b) \\ x(b) & x(b)^2 + \sigma_2^2 \end{bmatrix}.$$

As mentioned above, we consider two different updating rules: constant gain RLS as in the previous work on this model, and standard SG learning with  $\Gamma = I$ .<sup>22</sup> We also consider two models for the true economy:

- **M1:** Constant coefficients:  $u_t = u$ ,  $\theta_t = \theta$ .
- **M2:** Time varying coefficients as in (18), where  $\beta_t = [u_t, \theta_t]'$  with  $V = \gamma^2 \sigma^2 M_z(\bar{b})^{-1}$ , the prior consistent with RLS.

Note that the misspecification of the government's model, coupled with the self-referential nature of the model makes neither learning rule optimal here. The prior in M2 is consistent with RLS only in the absence of the self-referential feedback which is present here.

We then run another simulation study, whose results are shown in Figures 2 and 3. We run 1000 simulations of 1000 periods each for a variety of gain settings. We take the

---

<sup>22</sup>SG learning with  $\Gamma = M_z^{-1}$  was similar to RLS for small gains, but had stability problems for larger gain settings, so we do not report it here.

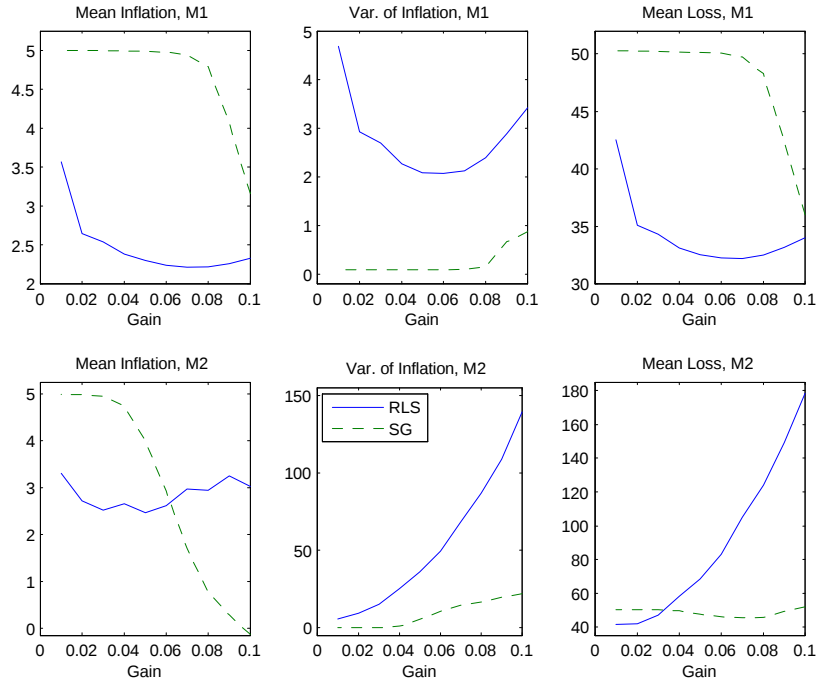


Figure 2: Mean and variance of inflation and loss under classic SG (dashed line) and RLS (solid line) learning with constant coefficients (Model 1) and random coefficients (Model 2). Means of 1000 simulations of 1000 periods.

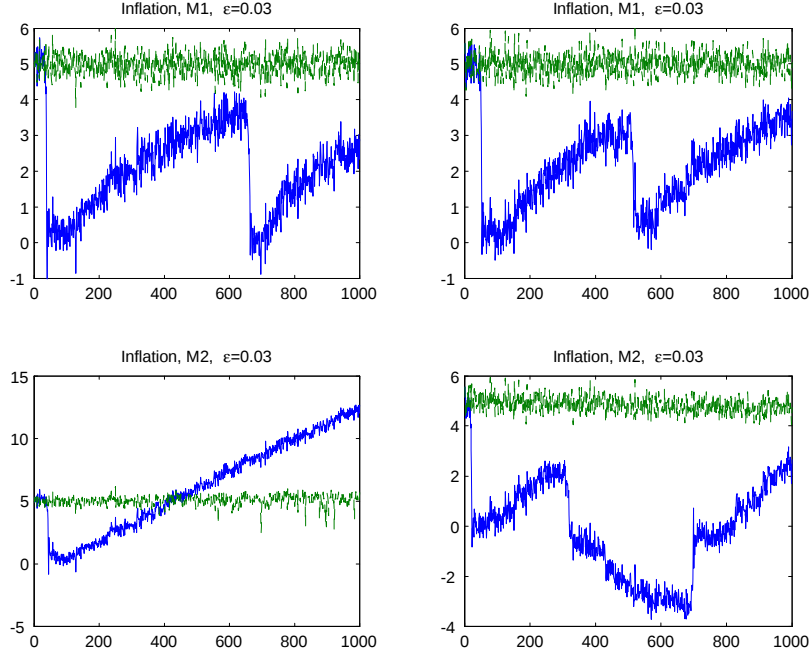


Figure 3: Two simulated time paths of inflation under SG (dashed line) and RLS (solid line) learning with constant coefficients (Model 1) and random coefficients (Model 2).

parameterization of the model from (Sargent 1999), and we initialize the government's beliefs at the self-confirming values  $\bar{b}$ .<sup>23</sup> For M2, we set the initial values  $\beta_0$  equal to the constant values from M1. Figure 2 plots the mean and variance of inflation and the loss averaged over the 1000 simulations in M1 and M2 under the two learning rules, while Figure 3 plots the time series from two representative simulation runs in M1 and M2 for a gain of  $\gamma = 0.03$ .

Several points are worth noting about the figures. First, as Figure 3 makes clear, we find that escapes happen under RLS but not under SG learning.<sup>24</sup> The plots of inflation under RLS in M1 look just like those in (Sargent 1999) and (Cho, Williams, and Sargent 2002), exhibiting the recurrent escapes from the high time-consistent inflation rate (5 in this case) to the low socially optimal rate (zero). But under SG learning the inflation rate always remains near the self-confirming level. Thus SG is robust, as it remains near its limit in this misspecified model. However, escapes are beneficial in this model, as they lead to the socially optimal level of inflation. Thus the good statistical properties of SG learning lead to losses in welfare here. This is clear from Figure 2 where we see that the mean rate of inflation is much lower under RLS, although the escapes make the variance higher than under SG. But the lower mean dominates and leads to the lower losses seen in the figure.<sup>25</sup>

<sup>23</sup>In particular, we set  $\theta = 1$ ,  $u = 5$ ,  $\sigma_1 = \sigma_2 = \sigma = 0.3$ .

<sup>24</sup>This behavior was also noted by (Evans and Honkapohja 2001). There can be escapes with GSG. It seems that in this model a non-diagonal  $\Gamma$  is necessary for escapes.

<sup>25</sup>The losses also depend on the 2nd moment of unemployment, but unemployment is exogenous here and

However with time-varying coefficients in M2, the robustness properties of SG learning lead to improvements in welfare as well as in estimation, at least for a range of gain settings. Again we see from Figure 2 that the mean inflation rate is lower under RLS than SG for small gain settings, although the mean drops for SG under larger gains. However the variance of inflation under RLS is now significantly higher than under SG (note the units in the figure). Thus for gain settings larger than 0.03 the large variance more than offsets the lower mean and leads to larger losses. As Figure 3 makes clear, for small gain settings the behavior of SG is nearly the same in M2 as in M1. The inflation rate fluctuates in a neighborhood of the initial self-confirming equilibrium throughout the sample. RLS again experiences large fluctuations similar to escapes, but instead of slowly returning to the equilibrium, the inflation rate drifts substantially.<sup>26</sup> The two simulated runs we show here illustrate that under RLS inflation may rise well above the equilibrium rate, or may fall well below zero. These large fluctuations lead to welfare losses relative to the stable inflation under SG learning.

## 5 Conclusions

We have examined convergence of adaptive learning when economic agents employ a stochastic gradient algorithm, which is sometimes used because of its ease of use in numerical simulations of global aspects of dynamics. The conditions for stability of SG learning differ from but are related to E-stability, which governs stability under least squares learning. We developed several sufficient conditions under which E-stability of the REE implies convergence of SG learning. It was shown that sensitivity with respect to units of measurement is a disadvantage of SG algorithms. However, there is a transformation of variables for which E-stability governs SG stability.

Important advantages of SG over RLS algorithms emerge when parameter drift is considered. First, a generalized SG algorithm under constant gain has an approximate Bayesian interpretation when the covariance matrix of the priors of parameter drift is proportional to the product of the observation noise variance and the covariance matrix of the regressors. (RLS algorithms have a similar property when the covariance matrix of the prior is instead proportional to the inverse of the covariance matrix of the regressors.) The SG rule can therefore adapt to relatively large parameter variation.

Second, it was shown that the generalized SG algorithm is a maximally robust optimal prediction rule when there is parameter uncertainty. In other words, generalized SG rules perform well across a range of alternative specifications. The maximal robustness property is also related to risk-sensitivity in optimal filtering: the risk-sensitive optimal filter is identical to the robust optimal filter, which in turn is just the constant gain generalized SG rule.

Our results suggest two conclusions. First, though SG stability is in general distinct from E-stability, the connections suggest that E-stability continues to be a natural property to

---

hence is independent of the learning rule.

<sup>26</sup>These are not true escapes, since the coefficients in the true model are varying over time. Similarly, we speak rather loosely of the equilibrium, since it also varies with the true model.

check in studies of convergence of learning. Second, the versatility of SG learning make it a viable model for learning, especially in settings where there is a lot of parameter variation or risk of misspecification. The constant gain GSG class handles parameter drift (Bayesian), deals with model misspecification (robust), and estimates parameters with lower risk (risk-sensitive).

## Appendix

### A Numerical Examples

**Example 1:** (E-stability does not imply SG-stability) For  $n = 1$  and  $k = 2$  select

$$A = -1.8800, F = \begin{pmatrix} -0.9390 & -0.6979 \\ 0.8722 & 0.0828 \end{pmatrix}, \Sigma_e = \begin{pmatrix} 1.0520 & -0.5164 \\ -0.5164 & 0.2581 \end{pmatrix}.$$

These yield

$$(M_w \otimes I)((F' \otimes A) - I) = \begin{pmatrix} -0.2503 & -1.3829 \\ 0.9012 & 0.3256 \end{pmatrix}.$$

This solution is E-stable but it is not convergent under SG learning since the eigenvalues of  $(M_w \otimes I)((F' \otimes A) - I)$  are  $0.0377 \pm 1.7086i$ .

**Example 2:** (SG-stability does not imply E-stability) For  $n = 1$  and  $k = 2$  select

$$A = -1.9022, F = \begin{pmatrix} -1.1281 & 0.7252 \\ -0.4944 & 0.0117 \end{pmatrix}, \Sigma_e = \begin{pmatrix} 0.5361 & 0.5760 \\ 0.5760 & 1.1807 \end{pmatrix}.$$

These yield eigenvalues  $-0.0838 \pm 0.3493i$  for  $(M_w \otimes I)((F' \otimes A) - I)$  and eigenvalues  $0.0618 \pm 0.3493i$  for  $F' \otimes A - I$ .

To find the counterexamples we simply conducted a random search over  $A, F$  and  $\Sigma_e$  under the required constraints. The Matlab routine is available on request.

### B Proofs of Results

**Proof of Corollary 4:** Since  $A - I$  is negative quasi-definite,  $(A + A')/2 - I$  has negative eigenvalues and  $(A + A')/2$  has roots less than one. It follows that  $F \otimes ((A + A')/2)$  has roots less than one and thus

$$\begin{aligned} F \otimes (A + A') - 2I &= (F \otimes A - I) + (F \otimes A' - I) \\ &= (F \otimes A - I) + (F \otimes A - I)' \end{aligned}$$

has negative roots, i.e.  $F \otimes A - I$  is negative quasi-definite. Thus  $F' \otimes A - I$  is H-stable and the result follows.

**Proof of Proposition 5:** We first prove that E-stability implies SG-stability. We can write the differential equation associated with SG learning for the matrix multiplying the exogenous variables as

$$\dot{b} = (AbF - b)M_w,$$

where we have dropped the inessential constant term involving  $B$ . Under our assumptions  $M_w$  is diagonal. Next, we write the coefficient matrix  $A$  in real Jordan canonical form:  $A = S\Lambda S^{-1}$ , where  $\Lambda$  is an upper block triangular matrix. The diagonal blocks are either  $1 \times 1$  blocks, consisting of real eigenvalues, or  $2 \times 2$  blocks of the form  $\begin{pmatrix} a & -b \\ b & a \end{pmatrix}$  for nonreal eigenvalues of the form  $a \pm bi$ . Multiplying we get

$$S^{-1}\dot{b} = (\Lambda S^{-1}bF - S^{-1}b)M_w,$$

which, defining  $q = S^{-1}b$ , is

$$\dot{q} = (\Lambda qF - q)M_w.$$

We then vectorize to get

$$\begin{aligned} \text{vec}(\dot{q}) &= (M_w F \otimes \Lambda - M_w \otimes I) \text{vec}(q) \\ &= (M_w \otimes I)(F \otimes \Lambda - I) \text{vec}(q). \end{aligned}$$

Now the matrix  $F \otimes \Lambda - I$  is block diagonal, i.e.

$$\begin{pmatrix} f_1\Lambda - I & 0 & \cdots & 0 \\ 0 & f_2\Lambda - I & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & f_K\Lambda - I \end{pmatrix},$$

where moreover  $f_i\Lambda - I$  are upper triangular matrices. Also  $M_w \otimes I$  is a diagonal matrix, so that we get

$$(M_w \otimes I)(F \otimes \Lambda - I) = \begin{pmatrix} m_1(f_1\Lambda - I) & 0 & \cdots & 0 \\ 0 & m_2(f_2\Lambda - I) & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & m_K(f_K\Lambda - I) \end{pmatrix}.$$

Each of the matrices  $m_k(f_k\Lambda - I)$  is upper block triangular with either diagonal elements of the form

$$m_k(f_k\lambda_i - 1),$$

where  $m_k > 0$  and where  $f_k \lambda_i - 1$  is negative by E-stability, or  $2 \times 2$  blocks of the form

$$m_k \left( f_k \begin{pmatrix} a & -b \\ b & a \end{pmatrix} - I_2 \right),$$

which again has eigenvalues with negative real parts by E-stability.

To prove that SG-stability implies E-stability, we note that  $m_k(f_k \lambda_i - 1) < 0$  for real eigenvalues of  $A$  and negativity of eigenvalues of  $m_k \left( f_k \begin{pmatrix} a & -b \\ b & a \end{pmatrix} - I_2 \right)$  for a complex pair of eigenvalues of  $A$  clearly also imply  $f_k \lambda_i - 1 < 0$  and negativity of eigenvalues of  $f_k \begin{pmatrix} a & -b \\ b & a \end{pmatrix} - I_2$ , respectively as  $m_k > 0$  for all  $k$ . The latter are just the E-stability conditions.

**Proof of Proposition 6:** We need to show that the eigenvalues of  $M_w(AF' - I)$  have negative real parts. We use results on the field of values of matrices given in (Horn and Johnson 1991). Let  $\mathcal{F}(N)$  denote the field of values of a matrix  $N$ , which is defined as  $\mathcal{F}(N) = \{z^* N z \mid z \in C^n \text{ with } |z| = 1\}$ . Here  $z^*$  denotes the complex conjugate of  $z$ . By assumption  $F'$  is a “contraction” in the sense used by Horn and Johnson. By p. 155 of (Horn and Johnson 1991) we can write  $F'$  as a finite sum  $F' = \sum c_i U_i$  where  $0 < c_i < 1$ , with  $\sum c_i = 1$ , and where  $U_i$  are unitary matrices. Since unitary matrices are normal,  $\mathcal{F}(U_i)$  is equal to the convex hull of the spectrum of  $U_i$ , which we denote  $\sigma(U_i)$ , see p. 11 of (Horn and Johnson 1991). Thus  $\mathcal{F}(U_i)$  is a subset of the unit disk since the eigenvalues of  $U_i$  lie exactly on the unit circle (see p. 71 of (Horn and Johnson 1985)). By the properties of fields of values given on pp. 9-10 of (Horn and Johnson 1991) we have

$$\mathcal{F}(F') \subset \sum c_i \mathcal{F}(U_i) \subset \text{unit disk}.$$

Hence  $\mathcal{F}(AF' - I) = A\mathcal{F}(F') - 1$  lies in the left half-plane of the complex plane. Next we note that  $\mathcal{F}(M_w)$  is a subset of the positive reals since  $M_w$  is symmetric positive definite. (This can be verified by direct computation). Finally, we use the result that  $\sigma(CD) \subset \mathcal{F}(C)\mathcal{F}(D)$  if  $D$  is positive semidefinite (p. 67 of (Horn and Johnson 1991)). Thus  $\sigma(M(AF' - I))$  lies in the negative half-plane.

**Proof of Lemma 9:** Equation (23) is an algebraic Riccati equation. Existence of a unique positive definite solution  $P$  follows from Theorem 3.7 of (Kwakernaak and Sivan 1972). The time invariant regulator associated with (23) takes the form  $\dot{x}(t) = u(t)$  and  $z(t) = x(t)$ . Thus in their general set-up we are setting  $A = 0$ ,  $B = D = I$ ,  $R_3 = \sigma^{-2}V$  and  $R_2^{-1} = M_z$ . The result requires that the associated regulator is stabilizable and detectable. Stabilizability is in turn implied by complete controllability (see their Theorem 1.27), which follows from their Theorem 1.23. Detectability is implied by complete reconstructability (see their Theorem 1.36), which follows from their Theorem 1.32. This establishes the existence of a unique  $P$ .

To show local convergence we linearize and vectorize (22), yielding

$$\text{vec } dP_{t+1} = (I - PM \otimes I - I \otimes PM) \text{vec } dP_t,$$

where  $dP_t$  refers to the deviation from the steady state  $P$ . The eigenvalues of the coefficient matrix are given by  $1 - 2\lambda$  where the  $\lambda$  are the eigenvalues of  $PM$ . This follows from Theorem 4.4.5 of (Horn and Johnson 1991) concerning the eigenvalues of the Kronecker sum of two matrices. Next, note that Theorem 7.6.3 of (Horn and Johnson 1985) implies that the eigenvalues of  $PM$  are positive. Finally, for  $V$  sufficiently small,  $P$  is small and the eigenvalues of  $PM$  can be made small. Thus  $V$  sufficiently small implies that all eigenvalues of the coefficient matrix have modulus strictly less than one. Local stability follows.

**Proof of Proposition 10:** (i) Given the positive definite matrix  $\Omega$ , Lemma 9 shows that for  $\gamma > 0$  sufficiently small, and hence for  $V$  sufficiently small, the difference equation (22) has a unique positive definite fixed point  $P$  and  $P$  is locally stable. It follows that the evolution of  $\varphi_t$  can be approximated by (25). Finally, as noted in Section 3.3, local stability of  $\bar{\varphi}$  under (25) is determined by (16).

(ii) Under the transformation  $\tilde{z}_t = Dz_t$  the model being estimated becomes

$$\begin{aligned} y_t &= \tilde{\beta}'_{t-1} \tilde{z}_t + \eta_t \\ \tilde{\beta}_t &= \tilde{\beta}_{t-1} + \tilde{\Lambda}_t, \end{aligned}$$

where  $\tilde{\beta}_t = D^{-1}\beta_t$  and  $\tilde{\Lambda}_t = D^{-1}\Lambda_t$  with  $\text{cov}(\tilde{\Lambda}_t) = D^{-1}VD^{-1}$ . The corresponding estimator of  $\tilde{\beta}_t$  is given by

$$\begin{aligned} \tilde{\varphi}_{t+1} &= \tilde{\varphi}_t + \tilde{P}_t \tilde{z}_t (y_t - \tilde{\varphi}'_t \tilde{z}_t) \\ \tilde{P}_{t+1} &= \tilde{P}_t - \tilde{P}_t M_{\tilde{z}} \tilde{P}_t + \sigma^{-2} \tilde{V}, \end{aligned}$$

where  $M_{\tilde{z}} = \lim_{t \rightarrow \infty} E \tilde{z}_t \tilde{z}'_t = DM_z D$  and  $\tilde{V} = D^{-1}VD^{-1}$ . The initial priors will also be related by  $\tilde{\varphi}_0 = D^{-1}\varphi_0$  and  $\tilde{P}_0 = D^{-1}P_0 D^{-1}$ . It is easily seen that the  $\tilde{\varphi}_t, \tilde{P}_t$  system is equivalent to (21)-(22) with  $\varphi_t = D\tilde{\varphi}_t$  and  $P_t = D\tilde{P}_t D$ .

**Proof of Proposition 11:** By (ii) of Proposition 10, the Bayesian learning rule is invariant to a transformation of variables. Letting  $\tilde{w}_t = Lw_t$  we have  $\tilde{w}_t = \tilde{F}\tilde{w}_{t-1}$  and  $M_{\tilde{z}} = I$ . Since (26) is H-stable then (16) is stable for all  $\Gamma$  and hence for all  $V$ . Stability of the REE then follows by (ii) of Proposition 10.

## C The Model with Lagged Endogenous Variables

Suppose that the model is extended to include lagged endogenous variables:

$$\begin{aligned} y_t &= \alpha + AE_t^* y_{t+1} + Bw_t + Cy_{t-1}, \\ w_t &= Fw_{t-1} + e_t. \end{aligned} \tag{38}$$

The MSV REE now takes the form

$$y_t = a + bw_t + cy_{t-1}, \quad (39)$$

where the REE values  $(\bar{a}, \bar{b}, \bar{c})$  solve the equations

$$\begin{aligned} (I - Ac - A)a &= \alpha \\ C\bar{c}^2 - c + C &= 0 \\ (I - Ac)b - AbF &= B. \end{aligned}$$

Under learning we assume that agents see current shocks but not the current value of the endogenous variable. The PLM of the agents is (39) and the ALM is easily computed to be

$$y_t = \alpha + A(I + c)a + (A(bF + cb) + B)w_t + Ac^2y_{t-1}$$

and thus the  $T$ -map is

$$T(\varphi)' = (\alpha + A(I + c)a, A(bF + cb) + B, Ac^2).$$

Defining  $z'_t = (\mathbf{1}', w'_t, y'_{t-1})$  and  $\varphi' = (a, b, c)$ , the formal details of the GSG algorithm and the differential equation are analogous to (4). In particular, the associated differential equation is

$$\frac{d\varphi'}{d\tau} = (T(\varphi) - \varphi)'M_z(\varphi)\Gamma$$

where  $M_z(\varphi) = \lim E z_t(\varphi) z'_t(\varphi)$  and where  $z_t(\varphi)$  denotes the stochastic process (39) with fixed  $\varphi$ . Fixed points of  $T(\varphi)$  correspond to REE and we require that the REE process of interest be stationary, i.e. the eigenvalues of  $\bar{c}$  lie inside the unit circle.

Linearizing the transposed differential equation, taking the differential and evaluating at the fixed point  $\bar{\varphi}$ , we get

$$\begin{aligned} d[(T(\varphi) - \varphi)'M_z(\varphi)]|_{\varphi=\bar{\varphi}} &= (T(\bar{\varphi}) - \bar{\varphi})'dM_z(\bar{\varphi})\Gamma + (dT(\bar{\varphi}) - d\bar{\varphi})'M_z(\bar{\varphi})\Gamma \\ &= (dT(\bar{\varphi}) - d\bar{\varphi})'M_z(\bar{\varphi})\Gamma, \end{aligned}$$

since the first term  $(T(\bar{\varphi}) - \bar{\varphi})'dM_z(\bar{\varphi})\Gamma$  is zero. Next, letting  $dT' = [dT'_a, dT'_b, dT'_c]$  and vectorizing we have

$$dvec[(T(\varphi) - \varphi)'M_z(\varphi)\Gamma]|_{\varphi=\bar{\varphi}} = (\Gamma M_z \otimes I)(DT' - I)dvec\varphi',$$

where  $M_z$  and  $DT'$  are evaluated at  $\bar{\varphi}$ . The linearized differential equation is thus

$$\frac{dvec\varphi'}{d\tau} = (\Gamma M_z \otimes I)(DT' - I)vec\varphi'. \quad (40)$$

From this equation we see that we get the same formal structure as in the forward-looking model, but now the multiplying matrix  $M_z$  involves the moments of the exogenous and endogenous variables. It is possible to compute explicitly the expression for  $DT'$ :

$$DT' = \begin{pmatrix} I \otimes A + I \otimes A\bar{c} & 0 & \bar{a}' \otimes A \\ 0 & F' \otimes A + I \otimes A\bar{c} & \bar{b}' \otimes A \\ 0 & 0 & I \otimes A\bar{c} + \bar{c}' \otimes A \end{pmatrix}.$$

We conclude that the SG-stability condition is that all eigenvalues of  $(\Gamma M_z \otimes I)(DT' - I)$  have negative real parts.

## References

- ANDERSON, E. W. (2004): “The Dynamics of Risk-Sensitive Allocations,” *Journal of Economic Theory*, forthcoming.
- ARROW, K. J. (1974): “Stability Independent of Adjustment Speed,” in (Horwich and Samuelson 1974), pp. 181–201.
- ARROW, K. J., AND M. MCMANUS (1958): “A Note on Dynamic Stability,” *Econometrica*, 26, 448–454.
- BARUCCI, E., AND L. LANDI (1997): “Least Mean Squares Learning in Self-Referential Stochastic Models,” *Economics Letters*, 57, 313–317.
- BENVENISTE, A., M. METIVIER, AND P. PRIOURET (1990): *Adaptive Algorithms and Stochastic Approximations*. Springer-Verlag, Berlin.
- BERNANKE, B., AND M. WOODFORD (eds.) (2005): *The Inflation-Targeting Debate*. University of Chicago Press, Chicago, USA.
- BICHI, G.-I., AND R. MARIMON (2001): “Global Stability of Inflation Target Policies with Adaptive Agents,” *Macroeconomic Dynamics*, 5, 148–179.
- BULLARD, J., AND I.-K. CHO (2005): “Escapist Policy Rules,” *Journal OF Economic Dynamics and Control*, *rm* forthcoming.
- BULLARD, J., AND J. DUFFY (2004): “Learning and Structural Change in Macroeconomic Data,” manuscript.
- BULLARD, J., AND S. EUSEPI (2005): “Did the Great Inflation Occur Despite Policymaker Commitment to a Taylor Rule?,” *Review of Economic Dynamics*, 8, 324–359.
- BULLARD, J., AND K. MITRA (2002): “Learning About Monetary Policy Rules,” *Journal of Monetary Economics*, 49, 1105–1129.
- CARLSON, D. (1968): “A New Criterion for H-Stability of Complex Matrices,” *Linear Algebra and Applications*, 1, 59–64.
- CHO, I.-K., AND K. KASA (2002): “Learning Dynamics and Endogenous Currency Crises,” mimeo.
- CHO, I.-K., N. WILLIAMS, AND T. J. SARGENT (2002): “Escaping Nash Inflation,” *Review of Economic Studies*, 69, 1–40.
- CLARIDA, R., J. GALI, AND M. GERTLER (1999): “The Science of Monetary Policy: A New Keynesian Perspective,” *Journal of Economic Literature*, 37, 1661–1707.

- EPSTEIN, L. G., AND S. E. ZIN (1989): “Substitution, Risk Aversion and the Temporal Behavior of Consumption and Asset Returns: A Theoretical Framework,” *Econometrica*, 57, 937–969.
- EVANS, G. W., AND S. HONKAPOHJA (1993): “Adaptive Forecasts, Hysteresis and Endogenous Fluctuations,” *Federal Reserve Bank of San Francisco Economic Review*, 1993(1), 3–13.
- (1998): “Stochastic Gradient Learning in the Cobweb Model,” *Economics Letters*, 61, 333–337.
- (2001): *Learning and Expectations in Macroeconomics*. Princeton University Press, Princeton, New Jersey.
- (2003a): “Adaptive Learning and Monetary Policy Design,” *Journal of Money, Credit and Banking*, 35, 1045–1072.
- (2003b): “Expectations and the Stability Problem for Optimal Monetary Policies,” *Review of Economic Studies*, 70, 807–824.
- (2005): “Policy Interaction, Expectations and the Liquidity Trap,” *Review of Economic Dynamics*, 8, 303–323.
- GIANNITSAROU, C. (2005): “E-Stability Does Not Imply Learnability,” *Macroeconomic Dynamics*, 9, 276–287.
- GLOVER, K., AND J. DOYLE (1988): “State-Space Formulae for All Stabilizing Controllers that Satisfy an H-infinity Norm Bound and Relations to Risk-Sensitivity,” *Systems & Control Letters*, 11, 167–172.
- HAMILTON, J. D. (1994): *Time Series Analysis*. Princeton University Press, Princeton, NJ.
- HANSEN, L. P., AND T. J. SARGENT (2004): “Misspecification in Recursive Macroeconomic Theory,” Book manuscript, University of Chicago and NYU.
- HASSIBI, B., A. H. SAYED, AND T. KAILATH (1996): “H-infinity Optimality of the LMS Algorithm,” *IEEE Transactions on Signal Processing*, 44, 267–280.
- HEINEMANN, M. (2000): “Convergence of Adaptive Learning and Expectational Stability: The Case of Multiple Rational Expectations Equilibria,” *Macroeconomic Dynamics*, 4, 263–288.
- HONKAPOHJA, S., AND K. MITRA (2005): “Learning Stability in Economies with Heterogeneous Agents,” manuscript.
- HORN, R., AND C. JOHNSON (1985): *Matrix Analysis*. Cambridge University Press, Cambridge.

- (1991): *Topics in Matrix Analysis*. Cambridge University Press, Cambridge.
- HORWICH, G., AND P. A. SAMUELSON (eds.) (1974): *Trade, Stability and Macroeconomics*. Academic Press, New York.
- JACOBSON, D. (1973): “Optimal Stochastic Linear Systems with Exponential Performance Criteria and Their Relation to Deterministic Games,” *IEEE Transactions on Automatic Control*, AC-18, 124–131.
- JOHNSON, C. R. (1974): “Sufficient Conditions for D-Stability,” *Journal of Economic Theory*, 9, 53–62.
- KASA, K. (2004): “Learning, Large Deviations, and Recurrent Currency Crises,” *International Economic Review*, 45, 141–173.
- KWAKERNAAK, H., AND R. SIVAN (1972): *Linear Optimal Control Systems*. Wiley-Interscience, New York.
- MAGNUS, J., AND H. NEUDECKER (1988): *Matrix Differential Calculus*. Wiley, New York.
- MARCET, A., AND J. P. NICOLINI (2003): “Recurrent Hyperinflations and Learning,” *American Economic Review*, 93, 1476–1498.
- MCGOUGH, B. (2005): “Shocking Escapes,” *Economic Journal*, forthcoming.
- ORPHANIDES, A., AND J. C. WILLIAMS (2005a): “Imperfect Knowledge, Inflation Expectations, and Monetary Policy,” in (Bernanke and Woodford 2005), chap. 5, pp. 201–234.
- (2005b): “Inflation Scares and Forecast-Based Monetary Policy,” *Review of Economic Dynamics*, 8, 498–527.
- ROTEMBERG, J. J., AND M. WOODFORD (1997): “An Optimization-Based Econometric Framework for the Evaluation of Monetary Policy,” *NBER Macroeconomics Annual*, 12, 297–346.
- SARGENT, T. J. (1999): *The Conquest of American Inflation*. Princeton University Press, Princeton NJ.
- SARGENT, T. J., AND N. WILLIAMS (2005): “Impacts of Priors on Convergence and Escapes from Nash Inflation,” *Review of Economic Dynamics*, 8, 360–391.
- SIMS, C. A. (1988): “Projecting Policy Effects with Statistical Models,” *Revista de Analisis Economico*, 3, 9–20.
- SVENSSON, L. E. (2003): “What is Wrong with Taylor Rules? Using Judgement in Monetary Policy through Targeting Rules,” *Journal of Economic Literature*, 41, 426–477.

- TALLARINI, T. D. (2000): “Risk-Sensitive Real Business Cycles,” *Journal of Monetary Economics*, 45, 507–32.
- WHITTLE, P. (1990): *Risk Sensitive Optimal Control*. Wiley, New York.
- WILLIAMS, N. (2001): “Escape Dynamics in Learning Models,” Ph.D. Dissertation, University of Chicago.
- WOODFORD, M. (2003): *Interest and Prices: Foundations of a Theory of Monetary Policy*. Princeton University Press, Princeton, NJ.

# CESifo Working Paper Series

(for full list see [www.cesifo-group.de](http://www.cesifo-group.de))

---

- 1517 Kira Boerner and Silke Uebelmesser, Migration and the Welfare State: The Economic Power of the Non-Voter?, August 2005
- 1518 Gabriela Schütz, Heinrich W. Ursprung and Ludger Wößmann, Education Policy and Equality of Opportunity, August 2005
- 1519 David S. Evans and Michael A. Salinger, Curing Sinus Headaches and Tying Law: An Empirical Analysis of Bundling Decongestants and Pain Relievers, August 2005
- 1520 Michel Beine, Paul De Grauwe and Marianna Grimaldi, The Impact of FX Central Bank Intervention in a Noise Trading Framework, August 2005
- 1521 Volker Meier and Matthias Wrede, Pension, Fertility, and Education, August 2005
- 1522 Saku Aura and Thomas Davidoff, Optimal Commodity Taxation when Land and Structures must be Taxed at the Same Rate, August 2005
- 1523 Andreas Haufler and Søren Bo Nielsen, Merger Policy to Promote ‘Global Players’? A Simple Model, August 2005
- 1524 Frederick van der Ploeg, The Making of Cultural Policy: A European Perspective, August 2005
- 1525 Alexander Kemnitz, Can Immigrant Employment Alleviate the Demographic Burden? The Role of Union Centralization, August 2005
- 1526 Baoline Chen and Peter A. Zadrozny, Estimated U.S. Manufacturing Production Capital and Technology Based on an Estimated Dynamic Economic Model, August 2005
- 1527 Marcel Gérard, Multijurisdictional Firms and Governments’ Strategies under Alternative Tax Designs, August 2005
- 1528 Joerg Breitschdel and Hans Gersbach, Self-Financing Environmental Mechanisms, August 2005
- 1529 Giorgio Fazio, Ronald MacDonald and Jacques Mélitz, Trade Costs, Trade Balances and Current Accounts: An Application of Gravity to Multilateral Trade, August 2005
- 1530 Thomas Christiaans, Thomas Eichner and Ruediger Pethig, A Micro-Level ‘Consumer Approach’ to Species Population Dynamics, August 2005
- 1531 Samuel Hanson, M. Hashem Pesaran and Til Schuermann, Firm Heterogeneity and Credit Risk Diversification, August 2005

- 1532 Mark Mink and Jakob de Haan, Has the Stability and Growth Pact Impeded Political Budget Cycles in the European Union?, September 2005
- 1533 Roberta Colavecchio, Declan Curran and Michael Funke, Drifting Together or Falling Apart? The Empirics of Regional Economic Growth in Post-Unification Germany, September 2005
- 1534 Kai A. Konrad and Stergios Skaperdas, Succession Rules and Leadership Rents, September 2005
- 1535 Robert Dur and Amihai Glazer, The Desire for Impact, September 2005
- 1536 Wolfgang Buchholz and Wolfgang Peters, Justifying the Lindahl Solution as an Outcome of Fair Cooperation, September 2005
- 1537 Pieter A. Gautier, Coen N. Teulings and Aico van Vuuren, On-the-Job Search and Sorting, September 2005
- 1538 Leif Danziger, Output Effects of Inflation with Fixed Price- and Quantity-Adjustment Costs, September 2005
- 1539 Gerhard Glomm, Juergen Jung, Changmin Lee and Chung Tran, Public Pensions and Capital Accumulation: The Case of Brazil, September 2005
- 1540 Yvonne Adema, Lex Meijdam and Harrie A. A. Verbon, The International Spillover Effects of Pension Reform, September 2005
- 1541 Richard Disney, Household Saving Rates and the Design of Social Security Programmes: Evidence from a Country Panel, September 2005
- 1542 David Dorn and Alfonso Sousa-Poza, Early Retirement: Free Choice or Forced Decision?, September 2005
- 1543 Clara Graziano and Annalisa Luporini, Ownership Concentration, Monitoring and Optimal Board Structure, September 2005
- 1544 Panu Poutvaara, Social Security Incentives, Human Capital Investment and Mobility of Labor, September 2005
- 1545 Kjell Erik Lommerud, Frode Meland and Odd Rune Straume, Can Deunionization Lead to International Outsourcing?, September 2005
- 1546 Robert Inklaar, Richard Jong-A-Pin and Jakob de Haan, Trade and Business Cycle Synchronization in OECD Countries: A Re-examination, September 2005
- 1547 Randall K. Filer and Marjorie Honig, Endogenous Pensions and Retirement Behavior, September 2005
- 1548 M. Hashem Pesaran, Til Schuermann and Bjoern-Jakob Treutler, Global Business Cycles and Credit Risk, September 2005

- 1549 Ruediger Pethig, Nonlinear Production, Abatement, Pollution and Materials Balance Reconsidered, September 2005
- 1550 Antonis Adam and Thomas Moutos, Turkish Delight for Some, Cold Turkey for Others?: The Effects of the EU-Turkey Customs Union, September 2005
- 1551 Peter Birch Sørensen, Dual Income Taxation: Why and how?, September 2005
- 1552 Kurt R. Brekke, Robert Nuscheler and Odd Rune Straume, Gatekeeping in Health Care, September 2005
- 1553 Maarten Bosker, Steven Brakman, Harry Garretsen and Marc Schramm, Looking for Multiple Equilibria when Geography Matters: German City Growth and the WWII Shock, September 2005
- 1554 Paul W. J. de Bijl, Structural Separation and Access in Telecommunications Markets, September 2005
- 1555 Ueli Grob and Stefan C. Wolter, Demographic Change and Public Education Spending: A Conflict between Young and Old?, October 2005
- 1556 Alberto Alesina and Guido Tabellini, Why is Fiscal Policy often Procyclical?, October 2005
- 1557 Piotr Wdowinski, Financial Markets and Economic Growth in Poland: Simulations with an Econometric Model, October 2005
- 1558 Peter Egger, Mario Larch, Michael Pfaffermayr and Janette Walde, Small Sample Properties of Maximum Likelihood Versus Generalized Method of Moments Based Tests for Spatially Autocorrelated Errors, October 2005
- 1559 Marie-Laure Breuillé and Robert J. Gary-Bobo, Sharing Budgetary Austerity under Free Mobility and Asymmetric Information: An Optimal Regulation Approach to Fiscal Federalism, October 2005
- 1560 Robert Dur and Amihai Glazer, Subsidizing Enjoyable Education, October 2005
- 1561 Carlo Altavilla and Paul De Grauwe, Non-Linearities in the Relation between the Exchange Rate and its Fundamentals, October 2005
- 1562 Josef Falkinger and Volker Grossmann, Distribution of Natural Resources, Entrepreneurship, and Economic Development: Growth Dynamics with Two Elites, October 2005
- 1563 Yu-Fu Chen and Michael Funke, Product Market Competition, Investment and Employment-Abundant versus Job-Poor Growth: A Real Options Perspective, October 2005
- 1564 Kai A. Konrad and Dan Kovenock, Equilibrium and Efficiency in the Tug-of-War, October 2005

- 1565 Joerg Breitung and M. Hashem Pesaran, Unit Roots and Cointegration in Panels, October 2005
- 1566 Steven Brakman, Harry Garretsen and Marc Schramm, Putting New Economic Geography to the Test: Free-ness of Trade and Agglomeration in the EU Regions, October 2005
- 1567 Robert Haveman, Karen Holden, Barbara Wolfe and Andrei Romanov, Assessing the Maintenance of Savings Sufficiency Over the First Decade of Retirement, October 2005
- 1568 Hans Fehr and Christian Habermann, Risk Sharing and Efficiency Implications of Progressive Pension Arrangements, October 2005
- 1569 Jovan Žamac, Pension Design when Fertility Fluctuates: The Role of Capital Mobility and Education Financing, October 2005
- 1570 Piotr Wdowinski and Aneta Zglinska-Pietrzak, The Warsaw Stock Exchange Index WIG: Modelling and Forecasting, October 2005
- 1571 J. Ignacio Conde-Ruiz, Vincenzo Galasso and Paola Profeta, Early Retirement and Social Security: A Long Term Perspective, October 2005
- 1572 Johannes Binswanger, Risk Management of Pension Systems from the Perspective of Loss Aversion, October 2005
- 1573 Geir B. Asheim, Wolfgang Buchholz, John M. Hartwick, Tapan Mitra and Cees Withagen, Constant Savings Rates and Quasi-Arithmetic Population Growth under Exhaustible Resource Constraints, October 2005
- 1574 Christian Hagist, Norbert Klusen, Andreas Plate and Bernd Raffelhueschen, Social Health Insurance – the Major Driver of Unsustainable Fiscal Policy?, October 2005
- 1575 Roland Hodler and Kurt Schmidheiny, How Fiscal Decentralization Flattens Progressive Taxes, October 2005
- 1576 George W. Evans, Seppo Honkapohja and Noah Williams, Generalized Stochastic Gradient Learning, October 2005