

NBER WORKING PAPER SERIES

LEARNING, LARGE DEVIATIONS AND RARE EVENTS

Jess Benhabib
Chetan Dave

Working Paper 16816
<http://www.nber.org/papers/w16816>

NATIONAL BUREAU OF ECONOMIC RESEARCH
1050 Massachusetts Avenue
Cambridge, MA 02138
February 2011

We thank Chryssi Giannitsarou, In-Koo Cho, John Duany, George Evans, Boyan Jovanovic, Tomasz Sadzik, Benoit de Saporta, and Tom Sargent for helpful comments and suggestions. The usual disclaimer applies. The views expressed herein are those of the author and do not necessarily reflect the views of the National Bureau of Economic Research.

© 2011 by Jess Benhabib and Chetan Dave. All rights reserved. Short sections of text, not to exceed two paragraphs, may be quoted without explicit permission provided that full credit, including © notice, is given to the source.

Learning, Large Deviations and Rare Events
Jess Benhabib and Chetan Dave
NBER Working Paper No. 16816
February 2011, Revised March 2011
JEL No. D83,D84

ABSTRACT

We examine the asymptotic distribution of estimated coefficients and endogenous variables in a dynamic self-referential model when agents learn adaptively using a constant gain stochastic gradient algorithm. The model environment can represent a number of economic models, including asset pricing models, that have been studied recently in the adaptive learning framework. The asymptotic distributions of forecasts and endogenous variables are characterized using techniques from linear recursions with multiplicative noise and large deviations, and are shown to exhibit fat tails.

Jess Benhabib
Department of Economics
New York University
19 West 4th Street, 6th Floor
New York, NY 10012
and NBER
jess.benhabib@nyu.edu

Chetan Dave
New York University
Department of Economics
19 W. 4th Street, 6FL
New York, NY, 10012
cdave@nyu.edu

1. Introduction

A large literature has recently focussed on studying how rational expectations equilibria (REE) can be attained in an economy where agents use adaptive learning mechanisms. In their seminal works Evans and Honkapohja¹ replace expectations with regressions to study in detail how ‘learning’ leads to convergence to REE in dynamic stochastic macroeconomic models. Such ‘adaptive learning’ algorithms assume that agents form expectations by conducting regressions on data available to them (within the model), with the most commonly used regressions being of the recursive least squares variety. Sargent (1999) and Cho et al. (2002) delve deeper into the notion of recursive least squares learning and consider environments in which agents place heavier emphasis on recent observations to draw inferences about model parameters, using ‘constant gain’ learning algorithms. Under such least squares constant gain learning algorithms, uncertainty about estimated parameters persists, and can fuel ‘escape’ dynamics in which a sequence of rare and unusual shocks propel agents away from the REE.² Characterizing the limiting probabilities of such escape dynamics and large deviations from REE is the subject of our paper.

Our context is a simple but widely employed univariate linear expectational difference equation that characterizes equilibrium dynamics for a number of models, for example, asset pricing or overlapping generations models and others.³ We demonstrate that the constant gain learning algorithm, specialized for expository clarity to its stochastic gradient version

¹See in particular Evans and Honkapohja (2001), and also Marcet and Sargent (1989) and Woodford (1990).

²See for example Williams (2009) and Evans, Honkapohja and Williams (2010) for an excellent discussion of constant gain-stochastic learning gradient algorithms. Such algorithms are special versions of the early Robbins and Munro (1951) learning algorithms and simplifications of the Kalman filter.

³See Evans and Honkapohja (1999, 2001) and Carceles-Poveda and Giannitsarou (2007) for an overview.

(see Evans, Honkapohja and Williams (2010)), yields a recursion where occasional large deviations or ‘rare events’ can induce a limiting ‘fat tailed’ power law distribution for the estimated coefficients, and therefore for the endogenous variables that they affect.⁴ In the asset price model interpretation of the reduced form, the REE relates dividends to asset prices. Under adaptive learning, our results show that the ratios of asset prices to dividends can significantly deviate from their REE values.

The remainder of the paper is structured as follows. Section 2 specifies the model while Section 3 demonstrates the use of large deviation theory with random linear recursions characterizations of learning algorithms. Section 3 discusses some relevant comparative statics with the parameter governing the power law as a function of changes in model parameter values. Section 4 concludes.

2. The Model

Our focus is on univariate models whose reduced form is given by

$$p_t = \delta E_t(p_{t+1}) + \gamma d_t, \quad \delta \in (0, 1). \quad (1)$$

Here d_t denotes an exogenous Markov chain on $(\mathbb{R}, \mathcal{R})$ where $\mathbb{R}$ is the real line and $\mathcal{R}$ its Borel subsets:

$$d_t = \rho d_{t-1} + \varepsilon_t, \quad |\rho| < 1, \quad t = 1, 2, \dots \quad (2)$$

⁴By fat-tailed distributions, we mean distributions for which some higher order moments do not exist.

in which ε_t is an *i.i.d.* random variable with compact support $[-a, a]$, $a > 0$, and a non-singular distribution function F .⁵ The linear expectational difference equation in (1) is a widely studied reduced form related to several linearized rational expectations models in economics, such as a linearized approximation to the Euler equation of an asset pricing model with a single asset and CRRA preferences.⁶

A widely held assumption is that agents form expectations as⁷

$$E_t(p_{t+1}) = \phi_{t-1} d_t. \quad (3)$$

When inserted into (1), this assumption yields what is known as the actual law of motion (ALM)

$$p_t = \underbrace{(\delta\phi_{t-1} + \gamma)\rho d_{t-1}}_{=T(\phi_{t-1})} + \underbrace{(\delta\phi_{t-1} + \gamma)}_{=V(\phi_{t-1})=\frac{T(\phi_{t-1})}{\rho}} \varepsilon_t, \quad (4)$$

that drives the dynamics of the endogenous variable (here, p_t) as a function of the exogenous process. The perceived law of motion (PLM) corresponding to the above ALM is

$$p_t = \phi_{t-1} d_{t-1} + \xi_t \quad (5)$$

where ξ_t is a regression error the agent employs to estimate the ϕ_{t-1} parameter of the

⁵ F is non-singular with respect to the Lebesgue measure if there exists a function $f \in R_+$, $\int_R f(t)dt > 0$, such that $F(dt) \geq f(t)dt$.

⁶ For a model with a single asset and CRRA preferences parameterized by θ the Euler equation

$$P_t = E_t \left\{ \delta \left(\frac{D_{t+1}}{D_t} \right)^{-\theta} (P_{t+1} + D_{t+1}) \right\}$$

can be linearized around its' non-stochastic steady state to yield the reduced form (1) provided that the lower case variables in (1) are interpreted as logarithmic deviations from steady state and $\gamma \equiv (1 - \delta - \theta)\rho + \theta$.

⁷ See footnote 8.

PLM via recursive least squares (RLS) or any other adaptive learning algorithm.⁸ Equating coefficients of (4) and (5), the rational expectations equilibrium (REE) value for ϕ is a constant,

$$\phi^{REE} = \frac{\gamma\rho}{1 - \delta\rho} \quad (6)$$

for all $\delta\rho \neq 1$, a condition that we maintain. The focus in the adaptive learning literature is on the ability of agents to learn ϕ^{REE} using the data available to them.⁹

We focus on the case in which agents employ a constant gain stochastic gradient learning algorithm to update ϕ_t as

$$\phi_t = \phi_{t-1} + g d_{t-1}(p_t - d_{t-1}\phi_{t-1}), \quad g \in (0, 1) \quad (7)$$

where the parameter g is referred to as the gain parameter.¹⁰

⁸It is standard in the adaptive learning literature to assume that, since p_{t+1} and the forecast ϕ_t are simultaneous, that agents do not know ϕ_t in forming $E(p_{t+1})$ and use ϕ_{t-1} instead, as in (3).

⁹The $T(\phi_{t-1})$ in (4) is the T -map associated with the ALM. Evans and Honkapohja (2001) use this map to show the expectational stability of the ϕ^{REE} , the fixed point of the T -map.

¹⁰A typical setting is $g = \frac{1}{T}$ where the fixed T is the horizon of time that the agent considers for updating ϕ_t . Constant gain algorithms are particularly useful for examining issues related to structural change. An optimal Bayesian derivation of the constant gain stochastic gradient algorithm under parameter drift, that is when agents also expect ϕ_t to drift according to a random walk, is given by Sargent and Williams (2005), and by Evans, Honkapohja and Williams (2010) in section 2 of their paper. The residual uncertainty in ϕ_t in each period prevents the weight given to recent observations and the optimal gain parameter from going to zero. In our model this yields a particular interpretation of the gain parameter g without altering the derivations and analysis.

3. Characterizing Large Deviations

The SGCG algorithm (7) can be re-written as

$$\phi_t = \phi_{t-1} + gd_{t-1}(p_t - d_{t-1}\phi_{t-1}) = (1 - gd_{t-1}^2)\phi_{t-1} + gd_{t-1}p_t. \quad (8)$$

Inserting the ALM in place of p_t yields an equation whose asymptotics are often analyzed in order to determine the stability of ϕ^{REE} (Carceles-Poveda and Giannitsarou (2007)). The substitution yields

$$\begin{aligned} \phi_t &= (1 - gd_{t-1}^2)\phi_{t-1} + gd_{t-1}[(\delta\phi_{t-1} + \gamma)\rho d_{t-1} + (\delta\phi_{t-1} + \gamma)\varepsilon_t] \\ &= [1 - (1 - \rho\delta)gd_{t-1}^2 + \delta gd_{t-1}\varepsilon_t]\phi_{t-1} + \gamma\rho gd_{t-1}^2 + \gamma gd_{t-1}\varepsilon_t. \end{aligned} \quad (9)$$

Give our interest in applying the results from the theory of large deviations and rare events, we re-write the above as

$$\phi_{t+1} = \lambda_{t+1}\phi_t + \psi_{t+1} \quad (10)$$

$$\lambda_{t+1} = 1 - (1 - \rho\delta)gd_t^2 + \delta gd_t\varepsilon_{t+1} = 1 - gd_t^2 + g\delta d_{t+1}d_t \quad (11)$$

$$\psi_{t+1} = \gamma\rho gd_t^2 + \gamma gd_t\varepsilon_{t+1} = \gamma gd_{t+1}d_t. \quad (12)$$

We note that λ_{t+1} is a random variable, generating multiplicative noise, and can be the source of large deviations and fat tails for the stationary distribution of ϕ_{t+1} . In the rest of the paper we follow the work of Saporta (2005), Roitershtein (2007), Collamore (2009) to

characterize the tail of the distribution of ϕ_{t+1} .¹¹

Let $\mathbb{N} = 0, 1, 2, \dots$. We first note that the stationary $AR(1)$ Markov chain $\{d_t\}_{t \in \mathbb{Z}}$ given by (2) is uniformly recurrent, and has compact support $\left[\frac{-a}{1-\rho}, \frac{a}{1-\rho}\right]$ (see Nummelin (1984), p. 93). We denote the stationary distribution of $\{d_t\}_{t \in \mathbb{N}}$ by π . Since $\{d_t\}_{t \in \mathbb{N}}$ and ε_t for $t = 1, 2, \dots$ are bounded, so are $\{\lambda_t\}_{t \in \mathbb{N}}$ and $\{\psi_t\}_{t \in \mathbb{N}}$. In fact, following the first definition of Roitershtein (2007), $\{\lambda_t, \psi_t\}_{t \in \mathbb{N}}$ constitutes a Markov Modulated Process (MMP): conditional on d_t , the evolution of the random variables $\lambda_{t+1}(d_t, d_{t-1})$ and $\psi_{t+1}(d_t, d_{t-1})$ are given by

$$P(d_t \in A, (\lambda_t, \psi_t) \in B) = \int_A K(d, dy) G(d, y, B) |_{d=d_{t-1}}, \quad (13)$$

$$G(d, y, \cdot) = P((\lambda_t, \psi_t) \in \cdot) | d_{t-1} = d, d_t = y, \quad (14)$$

where $K(d, dy)$ is the transition kernel of the Markov chain $\{d_t\}_{t \in \mathbb{N}}$.

Next we seek restrictions on the support of the *i.i.d.* noise $\varepsilon_t \in [-a, a]$ to assure that $\{\lambda_t\}_{t \in \mathbb{N}}$ remains positive. We assume¹²:

$$a < \frac{(1-\rho)}{(g(1+\delta(1-2\rho)))^{0.5}} \quad (15)$$

¹¹For an application of these techniques to the distribution of wealth see Benhabib et al. (2011) and to regime switching, Benhabib (2010).

¹²Since at it's stationary distribution $d_t \in \left(\frac{-a}{1-\rho}, \frac{a}{1-\rho}\right)$, $\varepsilon_t \in (-a, a)$,

$$\begin{aligned} \lambda_{t+1} &= 1 - (1 - \rho\delta)gd_t^2 + \delta gd_t \varepsilon_{t+1} \\ &> 1 - g(1 - \rho\delta) \left(\frac{a}{1-\rho}\right)^2 - g\delta \frac{a^2}{1-\rho} \\ &= 1 - g \left(\frac{a}{1-\rho}\right)^2 (1 - \rho\delta + (1-\rho)\delta) \\ &= 1 - g \left(\frac{a}{1-\rho}\right)^2 (1 + \delta(1-2\rho)) \end{aligned}$$

So $\lambda_t > 0$ if $a < \frac{1-\rho}{(g(1+\delta(1-2\rho)))^{0.5}}$.

From (12) it is easy to show that $\lambda_t > 0$ if (15) holds.

Let $S_n = \sum_{t=1}^n \log \lambda_t$. Following Roitershtein (2007) and Collamore (2009)¹³ the tail of the stationary distribution of $\{\phi_t\}_t$ depends on the limit¹⁴

$$\Lambda(\beta) = \lim_{n \rightarrow \infty} \sup \frac{1}{n} \log E \prod_{t=1}^n (\lambda_t)^\beta = \lim_{n \rightarrow \infty} \sup \frac{1}{n} \log E[\exp(\beta S_n)] \quad \forall \beta \in \mathbb{R}. \quad (16)$$

Using results in Roitershtein (2007), we can now prove the following about the tails of the stationary distribution of $\{\phi_t\}_{t \in \mathbb{N}}$:

Proposition 1 *For π -almost every $d_0 \in [-a, a]$, there is a unique positive $\beta < \infty$ that solves $\Lambda(\beta) = 0$, and the following limits exist and are positive:*

$$K_1(d_0) = \lim_{\tau \rightarrow \infty} \tau^\beta P(\phi > \tau | d_0) \quad \text{and} \quad K_{-1}(d_0) = \lim_{\tau \rightarrow \infty} \tau^\beta P(\phi < -\tau | d_0). \quad (17)$$

Proof. The results follow directly from Roitershtein (2007), Theorem 1.6 if we show the following:

- (i) There exists a β_0 such that $\Lambda(\beta_0) < 0$. First we note that $\Lambda(0) = 0$ for all n . Note

¹³For results on processes driven by finite state Markov Chains see Saporta (2005).

¹⁴ $\lim_{n \rightarrow \infty} \sup \frac{1}{n} \log E[\exp(\beta S_n)]$ is the Gartner Ellis limit that also appears in Large Deviation theory. For an exposition see Hollander (2000).

also that

$$\begin{aligned}
\Lambda'(0) &= \lim_{n \rightarrow \infty} \sup \frac{1}{n} \frac{d \log E \prod_{t=1}^n (\lambda_t)^\beta}{d\beta} \Big|_{\beta=0} \\
&= \lim_{n \rightarrow \infty} \sup \frac{1}{n} \left(E \prod_{t=1}^n (\lambda_t)^\beta \right)^{-1} E \left(\prod_{t=1}^n (\lambda_t)^\beta \prod_{t=1}^n \log(\lambda_t) \right) \Big|_{\beta=0} \\
&= \lim_{n \rightarrow \infty} \sup \frac{1}{n} E \prod_{t=1}^n \log \lambda_t
\end{aligned}$$

For large n , as $\{\lambda_t\}_t$ converges to its stationary distribution ω , we have

$$\Lambda'(0) = \lim_{n \rightarrow \infty} \sup \frac{1}{n} E \prod_{t=1}^n \lambda_t = E_\omega(\log \lambda_t)$$

Note however that

$$E_\omega(\lambda_t) = 1 + g\delta E_\pi(d_t d_{t-1}) - gE_\pi(d_{t-1}^2) \quad (18)$$

$$= 1 + g\delta \rho \frac{\sigma^2}{1 - \rho^2} - g \frac{\sigma^2}{1 - \rho^2} \quad (19)$$

$$= 1 - g \frac{\sigma^2}{1 - \rho^2} (1 - \delta \rho) < 1 \quad (20)$$

Therefore $\Lambda'(0) = E_\omega \log(\lambda_t) < 0$, and there exists $\beta_0 > 0$ such that $\Lambda(\beta_0) < 0$.

(ii) There exists a β_1 such that $\Lambda(\beta_1) > 0$. As in (i) above, we can evaluate, using Jensen's inequality,

$$\Lambda(\beta) = \lim_{n \rightarrow \infty} \sup \frac{1}{n} \log E \prod_{t=1}^n (\lambda_t)^\beta = \lim_{n \rightarrow \infty} \sup \frac{1}{n} \log E[\exp(\beta S_n)] \quad (21)$$

$$= \lim_{n \rightarrow \infty} \sup \log (E[\exp(\beta S_n)])^{\frac{1}{n}} \geq \lim_{n \rightarrow \infty} \sup \log \left(E[\exp(\beta \frac{S_n}{n})] \right) \quad (22)$$

so that at the stationary distribution of $\{\lambda_t\}_{t \in \mathbb{N}}$

$$\Lambda(\beta) \geq \log E_\omega[\exp(\beta \log \lambda_t)] = \log \int_\lambda [\exp(\beta \log \lambda_t)] d\omega(\lambda). \quad (23)$$

As $\beta \rightarrow \infty$ for $\log \lambda < 0$ we have $[\exp(\beta \log \lambda_t)] \rightarrow 0$, but if $P_\omega(\log \lambda > 0) > 0$ at the stationary distribution of $\{\lambda_t\}_t$, then as $\lim_{\beta \rightarrow \infty} \Lambda(\beta) = \log \int_\lambda [\exp(\beta \log \lambda_t)] d\omega(\lambda) \rightarrow \infty$.

Therefore if we can show that $P_\omega(\log \lambda_t > 0) > 0$, it follows that there exists a β_1 for which $\Lambda(\beta_1) > 0$. Since $\Lambda(\beta)$ is convex¹⁵, it follows that there exists a unique κ for which $\Lambda(\kappa) = 0$.

To show that $P_\omega(\lambda > 1) > 0$, define $A = \left\{d \in \left(0, \frac{\mu a \delta}{1 - \rho \delta}\right)\right\}$, $\mu \in (0, 1)$ so that $\frac{\mu a \delta}{1 - \rho \delta} < \frac{a}{1 - \rho}$.

At its stationary distribution $\{d_t\}_{t \in \mathbb{N}}$ is uniformly recurrent over $\left[\frac{-a}{1 - \rho}, \frac{a}{1 - \rho}\right]$ which implies that $P_\pi(d_{t-1} \in A) > 0$. We have $\lambda_t = 1 - \delta g d_{t-1} (\delta^{-1}(1 - \rho \delta) d_{t-1} - \varepsilon_t)$, so for $d_{t-1} \in A$ and $\varepsilon_t \in (\mu a, a]$, it follows that $\lambda_t > 1$. Thus $P_\omega(\lambda_t > 1) = P_\pi(d_{t-1} \in A_t) P(\varepsilon_t \in (\mu a, a]) > 0$.

(iii) The non-arithmeticity assumption required by Roitershtein (2007) (p. 574, (A7))

holds¹⁶: There does not exist an $\alpha > 0$ and a function $G : \mathcal{R} \times \{-1, 1\} \rightarrow \mathbb{R}$ such that

$$P(\log |\lambda_t| \in G(d_{t-1}, \eta) - G(d_t, \eta \cdot \text{sign}(\lambda_t)) + \alpha \mathbb{N}) = 1 \quad (24)$$

and since $\lambda_t > 0$,

$$P(\log \lambda_t \in G(d_{t-1}, \eta) - G(d_t, \eta) + \alpha \mathbb{N}) = 1. \quad (25)$$

¹⁵This follows since the moments of nonnegative random variables are log convex (in β); see Loeve (1977, p. 158).

¹⁶See also Alsmeyer (1997). In other settings $\{\lambda_t\}_t$ may contain additional *i.i.d.* noise independent of the Markov Process $\{d_t\}_t$, in which case the non-arithmeticity is much more easily satisfied.

We have

$$\log \lambda_t = \log(1 - gd_{t-1}^2 + g\delta d_t d_{t-1}) = 1 - (1 - \rho\delta)gd_{t-1}^2 + \delta g d_{t-1} \varepsilon_t = F(d_{t-1}, \varepsilon_t) \quad (26)$$

which contains the cross-partial term $d_t d_{t-1}$. Therefore in general $F(d_{t-1}, \varepsilon_t)$ cannot be represented in separable form as $R(d_{t-1}, \eta) - R(d_t, \eta) + \alpha\mathbb{N} \quad \forall (d_{t-1}, d_t)$ where $d_t = \rho d_{t-1} + \varepsilon_t$. Suppose to the contrary that there is a small rectangle $[D, D^*] \times [E, E^*]$ in the space of (d, ε) , such that $F(d, \varepsilon) = R(d) - R(\rho d + \varepsilon)$, d is in the interior of $[D, D^*]$, and ε is in the interior of $[E, E^*]$, up to a constant from the discrete set $\alpha\mathbb{N}$, which we can ignore for variations if $[D, D^*] \times [E, E^*]$ that are small enough. Now fix d, d' close to one another in the interior of $[D, D^*]$. We must have, for $\varepsilon \in [E + \rho|d - d'|, E^* - \rho|d - d'|]$, that

$$F(d, \varepsilon) - R(d) = -R(\rho d + \varepsilon) = -R(\rho d' + \varepsilon + \rho(d - d')) \quad (27)$$

$$= F(d', \varepsilon + \rho(d - d')) - R(d'), \quad (28)$$

or $F(d, \varepsilon) - F(d', \varepsilon + \rho(d - d')) = R(d) - R(d')$. However the latter cannot hold since the cross-partial term $d_{t-1} \varepsilon_t$ in $F(d_{t-1}, \varepsilon_t) = 1 - (1 - \rho\delta)gd_{t-1}^2 + \delta g d_{t-1} \varepsilon_t$ is non-zero except of a set of zero measure where d or ε are zero.^{17,18}

¹⁷We thank Tomasz Sadzik for suggesting this proof for (iii).

¹⁸We can also avoid possible degeneracies that may occur if λ_t and ψ_t have a specific form of dependence so that

$$P(\phi|\lambda_t \phi + \psi_t = \phi) = 1.$$

Note

$$\begin{aligned} \phi &= \frac{\psi_t}{1 - \lambda_t} = \frac{\gamma \rho g d_t^2 + \gamma g d_t \varepsilon_{t+1}}{1 - (1 - \rho\delta)g d_t^2 + \delta g d_t \varepsilon_{t+1}} \\ &= \frac{\gamma}{\delta} \frac{\delta \rho g d_t^2 + \delta g d_t \varepsilon_{t+1}}{1 - (1 - \rho\delta)g d_t^2 + \delta g d_t \varepsilon_{t+1}} \end{aligned}$$

(iv) The positivity of $K_1(d_0)$ and $K_{-1}(d_0)$ follows from Condition G required by Roiterstein (2007); see his Definition 1.7 and subsequent discussion. This condition holds because $\lambda_t > 0$ for all t , and $\{d_t\}_{t \in \mathbb{N}}$ is uniformly recurrent and therefore also irreducible. ■

The Proposition above characterizes the tail of the stationary distribution of ϕ as a power tail with exponent κ . It follows that the distribution of ϕ has moments only up to the highest integer less than κ , and is a ‘fat tailed’ distribution rather than a Normal Distribution. The results are driven by the fact that the stationary distribution of $\{\lambda_t\}_{t \in \mathbb{N}}$ has a mean less than one but also support above 1 with positive probability. Then large deviations as strings of realizations of λ_t above one, even though they may be rare events, can produce fat tails.

In the asset price model ϕ relates the dividends to assets prices. Under adaptive learning, the results above show how the probability distribution of large deviations, or "escapes" of ϕ from its REE value is characterized by a fat tailed distribution, and will occur with higher likelihood than under a Normal distribution.¹⁹

We now briefly discuss the case where $\{d_t\}_t$ is an $MA(1)$ process. Proposition 1 still applies and we obtain similar results to the $AR(1)$ case. Let

$$d_t = \varepsilon_t + \zeta \varepsilon_{t-1}, \quad |\zeta| < 1, \quad t = 1, 2, \dots \quad (29)$$

Differentiating wrt ε_t , the right side is zero only if $\delta \rho g d_t^2 = 1 - (1 - \rho \delta) g d_t^2$, or $\delta \rho g = 1 - g + g \rho \delta$. This holds only if $g = 1$. So in general, for any d_t , there exists a constant ϕ such that $P(\phi | \lambda_t \phi + \psi_t = \phi) = 1$ only if $g = 1$, which we ruled out by assumption.

¹⁹In the model of Cho, Sargent and Williams (2002), the monetary authority has a misspecified Philips curve and sets inflation policy to optimize a quadratic target. The learning algorithm using a constant gain however is not linear in the recursively estimated parameters (the natural rate and the slope of the Philips curve).

Then at its stationary distribution $d_t \in [-a(1 + \zeta), a(1 + \zeta)]$. Under the PLM

$$p_t = \phi_{0t}\varepsilon_t + \phi_{1t}\varepsilon_{t-1}, \quad (30)$$

after observing ε_t at time t but not ϕ_{1t+1} , the agents expect

$$E_t(p_{t+1}) = \phi_{0t}E_t(\varepsilon_{t+1}) + \phi_{1t}E_t(\varepsilon_t) = \phi_{1t}\varepsilon_t \quad (31)$$

Then the ALM is

$$p_t = \delta\phi_{1t}\varepsilon_t + \theta(\varepsilon_t + \zeta\varepsilon_{t-1}) = [\delta\phi_{1t} + \theta]\varepsilon_t + \theta\zeta\varepsilon_{t-1}$$

and the REE is given by

$$\phi_0 = \theta(1 + \delta\zeta) \quad (32)$$

$$\phi_1 = \theta\zeta. \quad (33)$$

Under the learning algorithm in equation (7) we obtain

$$\phi_{1t} = \phi_{1t-1} + gd_{t-1}(p_t - \phi_{1t-1}d_{t-1}) \quad (34)$$

$$\phi_{1t+1} = \lambda_{t+1}\phi_{1t} + \psi_{t+1} \quad (35)$$

$$\lambda_{t+1} = 1 - gd_t^2 + g\delta\varepsilon_{t+1}d_t \quad (36)$$

$$\psi_{t+1} = g\theta\varepsilon_{t+1}d_t + \theta\zeta gd_t\varepsilon_t \quad (37)$$

It is straightforward to show that at the stationary distribution of $\{\lambda_t\}_t$, $E(\lambda_t) < 1$, and that $P(\lambda_t > 1) > 0$. It is also easy to check that $\lambda_t > 0$ if $a < ((1 + \zeta)(1 + \zeta - \delta))^{-0.5}$. With the latter restriction, it is easy to check that the other conditions in the proof of Proposition 1 are satisfied.

4. Comparative Statics

To explore how κ is related to the underlying parameters of our model, we can simulate the learning algorithm that updates ϕ , and then estimate κ using the Hill (1975) estimator. We can then explore how our estimate of κ from simulated series varies as we vary parameters.

We simulate 100 series for ϕ_t under the $AR(1)$ and $MA(1)$ assumptions for dividends with each series being of length 10000, and average our κ estimates. In the $AR(1)$ case we expect lower κ , or fatter tails, as the support of λ_t that lies above 1 gets larger. Since $\lambda_{t+1} = 1 - (1 - \rho\delta)gd_t^2 + \delta gd_t\varepsilon_{t+1}$, given the stationary distribution of $\{d_t\}_t$ and that of $\{\varepsilon_t\}_t$, the support of λ_t above 1 unambiguously increases if δ increases. Increasing ρ however has an ambiguous effect: while the term $(1 - \delta\rho)$ declines and tends to raise λ_t , the support of the stationary distribution of $\{d_t\}_t$ gets bigger with higher ρ , so that $(1 - \rho\delta)gd_t^2$ can now reduce λ_t and its support above 1 for large realizations of d_t^2 . Finally in our simulations decreasing g tends to shrink the support of λ_t that is above 1 and κ increases with g : as the gain parameter decreases towards zero, the tails of the stationary distribution of $\{\phi_t\}$ get thinner.²⁰

²⁰This of course is in accord with the Theorem 7.9 in Evans and Honkapohja (2001). As the gain $g \rightarrow 0$ and $tg \rightarrow \infty$, $\{\phi_t^g - \varkappa\}/g^{0.5}$ converges to a Gaussian variable where $\varkappa$ is the stable point of the associated ODE describing the mean dynamics.

We use the baseline parameterization, $(\rho, g, \delta, \theta) = (0.95, 0.01, 0.95, 2.5)$ and vary each element of (ρ, g, δ) while keeping the other two at their baseline values. We do not vary θ since it does not affect λ or κ . We choose the support of $\{\varepsilon_t\}$ in each case, as defined by the parameter a , so that the inequality (15) in the $AR(1)$ case is satisfied over the range over which we vary each parameter. The resulting a values corresponding to varying ρ , g , and δ are, respectively, $(0.24, 0.38, 1.03)$. We plot the results in Figure 1.

In the $MA(1)$ case we use the same baseline parametrization, except that now the parameter ρ is replaced with ζ . The a values chosen to satisfy inequality (15), and corresponding to varying ζ , g , and δ are now, respectively, $(0.69, 0.66, 0.68)$. We vary ζ between 0.85 and 0.99, g between 0.01 and 0.30 and δ between 0.85 and 0.99, again with 0.01 increments. As before, θ is not varied since it does not affect κ . We plot the comparative statics for the estimated average κ from our simulations below:

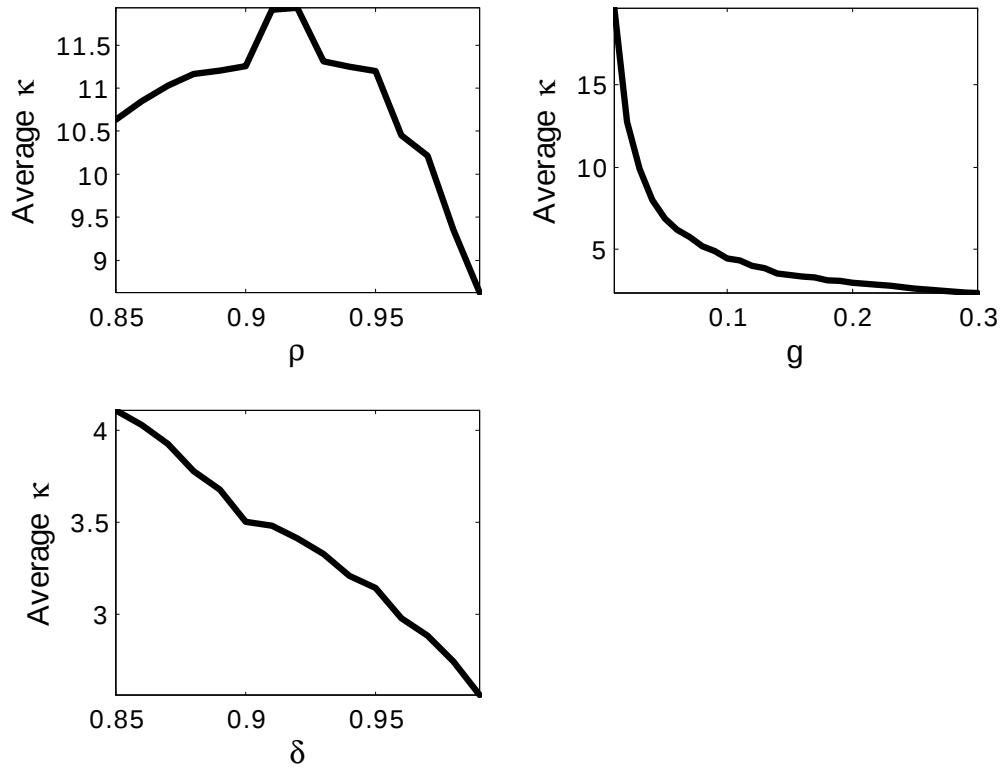


Figure 1. Average κ as a function of model parameters ($AR(1)$ case).

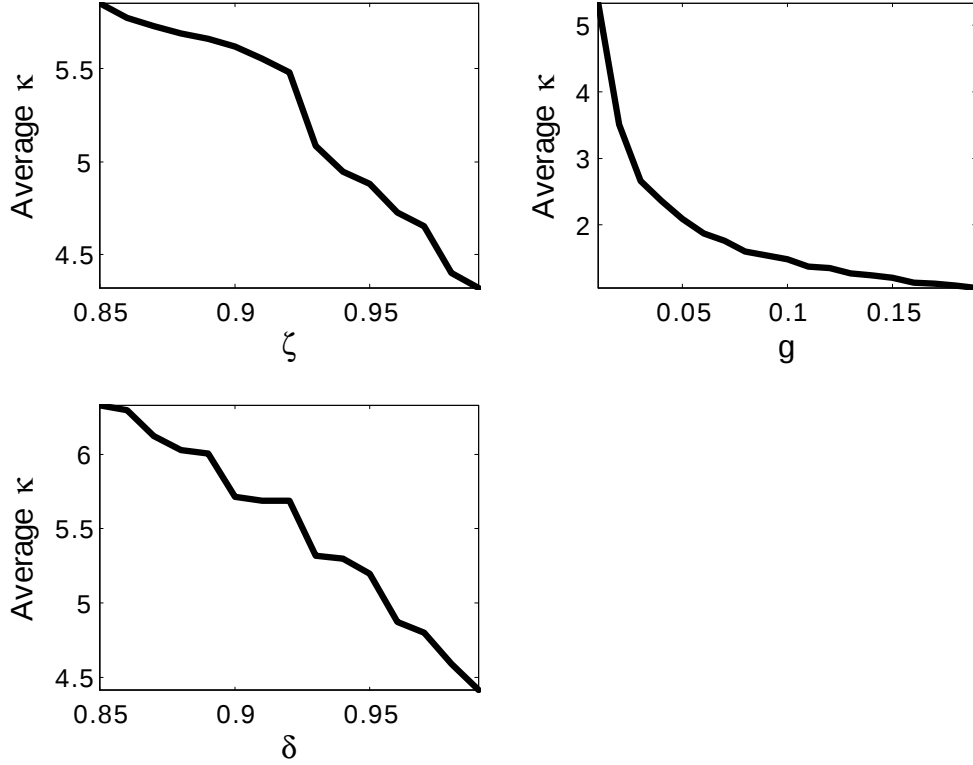


Figure 2. Average κ as a function of model parameters ($MA(1)$ case).

If we focus on the asset pricing interpretation of our model, we find that typically dividend data are exceptionally smooth: publicly traded corporations try to provide a steady stream of dividends to shareholders. Occasionally however under the stress of a rare financial crisis like the one of 2008-2009, dividends paid by some companies may collapse and trigger a large deviation in the forecasts and expectations of adaptive learners. To explore this using detrended dividend data, we can first estimate ρ (or ζ in the $MA(1)$ case), and then use dividend series to calculate the series for $\{\lambda_t\}_t$ and $\{\phi_t\}_t$. Using these series based on actual dividend data, we can then estimate κ . The estimated κ however will be sensitive to the specification of the stochastic process for dividends.

In the $AR(1)$ case, for Citibank dividend data over 1987-2009 we estimate $\rho = 0.8162$

and $\kappa = 20.198$ and for Curtiss-Wright $\rho = 0.2834$ and $\kappa = 10.217$. For Bank of America for data over 1986-2010 we obtain $\rho = 0.9819$ and $\kappa = 70.63$. Conducting the same estimation for monthly S and P 500 series from 1871 to 2010 we estimate $\rho = 0.99664$ and $\kappa = 1858.1$.

However, under the $MA(1)$ assumption for dividends, for Citibank dividend data we estimate $\zeta = 0.5227$ and $\kappa = 6.7582$, for Curtiss-Wright $\zeta = 0.2089$ and $\kappa = 3.3312$. In the case of Bank of America we estimate $\zeta = 0.8858$ and $\kappa = 16.657$. Finally, for the same linearly detrended monthly S and P 500 dividends data employed in the $AR(1)$ case, surprisingly, we estimate instead $\zeta = 0.95892$ and $\kappa = 7.499$.

5. Conclusion

An important and growing literature replaces expectations in dynamic stochastic models not with realizations and unforecastable errors, but with regressions where agents ‘learn’ the rational expectations equilibria. When such agents employ constant gain learning algorithms that put heavier emphasis on recent observations, escape dynamics can propel estimated coefficients away from the REE values. In an asset pricing interpretation of the model, ‘bubbles,’ or asset prices that exhibit large deviations from their REE ratios to dividends, can occur with a frequency associated with a fat tailed power law. The techniques used in our paper generalize to higher dimensions and to finite state Markov chains under certain assumptions,²¹ and can be applied to other more general economic models.

²¹See for example Saporta (2005) and Gosh et al. (2010).

References

- [1] Alsmeyer, G. (1997). The Markov Renewal Theorem and Related Results. *Markov Process Related Fields* 3 103–127
- [2] Benhabib, J. Bisin, A. and S. Zhu, 2011. The Distribution of Wealth and Fiscal Policy in Economies with Finitely Lived Agents. *Econometrica* 79, 123-158.
- [3] Benhabib, J., (2009). A Note Regime Switching, Monetary Policy and Multiple Equilibria, *NBER Working Paper* No. 14770.
- [4] Carceles-Poveda, E., Giannitsarou, C., 2007. Adaptive Learning in Practice. *Journal of Economic Dynamics and Control* 31, 2659-2697.
- [5] Cho, I-K., Sargent, T. J., Williams, N., 2002. Escaping Nash Inflation. *Review of Economic Studies* 69, 1-40.
- [6] Collamore, J. F. , 2009. Random Recurrence Equations and Ruin in a Markov-Dependent Stochastic Economic Environment, *Annals of Applied Probability* 19, 1404–1458.
- [7] Evans, G., Honkapohja, S., 2001. *Learning and Expectations in Macroeconomics*. Princeton University Press.
- [8] Evans, G., Honkapohja, S., and N. Williams, 2010. Generalized Stochastic Gradient Learning, *International Economic Review*, 51, 237-262.
- [9] Goldie, C. M., 1991. Implicit Renewal Theory and Tails of Solutions of Random Equations, *Annals of Applied Probability*, 1, 126–166.

- [10] Ghosh, A.P, Haya, D., Hirpara, H., Rastegar, R., Roitershtein, A.,Schulteis, A., and Suhe, J, (2010), Random Linear Recursions with Dependent Coefficients, *Statistics and Probability Letters* 80, 1597 1605
- [11] Hill, B. M., 1975. A Simple General Approach to Inference about the Tail of a Distribution, *Annals of Statistics*, 3, 1163–1174
- [12] Kesten, H., 1973. Random Difference Equations and Renewal Theory for Products of Random Matrices, *Acta Mathematica*. 131 207–248.
- [13] Marcet, A, and Sargent, T. J., 1989. Convergence of Least Squares Learning Mechanisms in Self-Referential Linear Stochastic Models, *Journal of Economic Theory*, 48, 337-368.
- [14] Nummelin, E., 1984. *General irreducible Markov chains and non-negative operators*. Cambridge Tracts in Mathematics 83, Cambridge University Press.
- [15] Roitershtein, A., 2007. One-Dimensional Linear Recursions with Markov-Dependent Coefficients, *The Annals of Applied Probability*, 17(2), 572-608.
- [16] Robbins, H. and Munro, S., 1951. A Stochastic Approximation Method, *Annals of Mathematical Statistics*, 22, 400-407.
- [17] Saporta, B., 2005. Tail of the stationary solution of the stochastic equation $Y_{n+1} = a_n Y_n + \gamma_n$ with *Markovian Coefficients*, *Stochastic Processes and their Applications*, 115(12), 1954-1978.
- [18] Sargent, T. J., 1999. *The Conquest of American Inflation*. Princeton University Press.

- [19] Sargent, T. J. and Williams, N., 2005. Impacts of Priors on Convergence and Escape from Nash Inflation, *Review of Economic Dynamics*, 8(2), 360-391.
- [20] Williams, N. 2009. Escape Dynamics, Manuscript.
- [21] Woodford, M., 1990. Learning to Believe in Sunspots, *Econometrica*, 58(2), 277-307.