

da Silveira, Rava Azeredo; Sung, Yeji; Woodford, Michael

Working Paper

Optimally Imprecise Memory and Biased Forecasts

CESifo Working Paper, No. 8709

Provided in Cooperation with:

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

Suggested Citation: da Silveira, Rava Azeredo; Sung, Yeji; Woodford, Michael (2020) : Optimally Imprecise Memory and Biased Forecasts, CESifo Working Paper, No. 8709, Center for Economic Studies and Ifo Institute (CESifo), Munich

This Version is available at:

<https://hdl.handle.net/10419/229527>

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.

Optimally Imprecise Memory and Biased Forecasts

Rava Azeredo da Silveira, Yeji Sung, Michael Woodford

Impressum:

CESifo Working Papers

ISSN 2364-1428 (electronic version)

Publisher and distributor: Munich Society for the Promotion of Economic Research - CESifo GmbH

The international platform of Ludwigs-Maximilians University's Center for Economic Studies and the ifo Institute

Poschingerstr. 5, 81679 Munich, Germany

Telephone +49 (0)89 2180-2740, Telefax +49 (0)89 2180-17845, email office@cesifo.de

Editor: Clemens Fuest

<https://www.cesifo.org/en/wp>

An electronic version of the paper may be downloaded

- from the SSRN website: www.SSRN.com
- from the RePEc website: www.RePEc.org
- from the CESifo website: <https://www.cesifo.org/en/wp>

Optimally Imprecise Memory and Biased Forecasts

Abstract

We propose a model of optimal decision making subject to a memory constraint. The constraint is a limit on the complexity of memory measured using Shannon's mutual information, as in models of rational inattention; but our theory differs from that of Sims (2003) in not assuming costless memory of past cognitive states. We show that the model implies that both forecasts and actions will exhibit idiosyncratic random variation; that average beliefs will also differ from rational-expectations beliefs, with a bias that fluctuates forever with a variance that does not fall to zero even in the long run; and that more recent news will be given disproportionate weight in forecasts. We solve the model under a variety of assumptions about the degree of persistence of the variable to be forecasted and the horizon over which it must be forecasted, and examine how the nature of forecast biases depends on these parameters. The model provides a simple explanation for a number of features of reported expectations in laboratory and field settings, notably the evidence of over-reaction in elicited forecasts documented by Afrouzi et al. (2020) and Bordalo et al. (2020a).

JEL-Codes: D840, E030, G410.

Keywords: rational inattention, over-reaction, survey expectations.

Rava Azeredo da Silveira
ENS & University of Basel
Basel / Switzerland
rava@iob.ch

Yeji Sung
Columbia University
New York / NY / USA
ys2992@columbia.edu

Michael Woodford
Columbia University
New York / NY / USA
michael.woodford@columbia.edu

November 12, 2020

We thank Hassan Afrouzi, Ben Hébert, David Laibson, Yueran Ma, and Andrei Shleifer for helpful discussions, and are especially grateful to the authors of Afrouzi et al. (2020) for sharing the data plotted in Figure 7. We also thank the Alfred P. Sloan Foundation, the CNRS through UMR 8023, and the IOB for research support.

The hypothesis of rational expectations (RE) proposes that decisions are based on expectations that make use of all available information in an optimal way: that is, those that would be derived by correct Bayesian inference from an objectively correct prior and the data that has been observed to that date. Yet both in surveys of individual forecasts of macroeconomic and financial variables and in forecasts elicited in experimental settings, beliefs are more heterogeneous than this hypothesis should allow, and forecast errors are predictable on the basis of variables observable by the forecasters, contrary to this hypothesis. In particular, a number of studies have argued that forecasts typically over-react to new realizations of the variable being forecasted. (See Bordalo *et al.*, 2020a, and Afrouzi *et al.*, 2020, for recent examples with extensive references to prior literature.)

A variety of models of expectation formation have been proposed that allow for such over-reaction. The simplest type of model simply posits that the forecasted future value of a variable is a particular linear function of the last few observations of the variable; with an appropriate choice of the coefficients (such as those proposed by Metzler, 1941), a forecasting heuristic of this kind may imply that a recent increase in the variable will be extrapolated into the future, so that further increases are anticipated, regardless of whether the degree of serial correlation of changes in the variable make this an optimal forecast. A classic critique of such proposals, however, is that of Muth (1961): why should decision makers continue to forecast in this way, if their forecasts are systematically biased, as repeated observations should eventually make clear? Moreover, a mechanical heuristic with fixed coefficients is unable to explain how the biases in observed forecasts change depending on the persistence of the process that is forecasted (Afrouzi *et al.*, 2020).

Fuster *et al.* (2010, 2011) propose a more sophisticated model, in which decision makers are assumed to forecast a time series by modeling it as a low-order autoregressive process; the coefficients of their forecasting rule are those implied by the AR(k) model that best fits the autocorrelation function of the actual series, for some fixed bound on k . The authors argue that actual time series often involve long-horizon dependencies, and show that in this case (say, an AR(40) process forecasted by people who consider models with no more than 10 lags), long-horizon forecasts using the best-fitting AR(k) model can significantly over-react to recent trends in the data.

This proposal, however, remains subject to several objections. Why should the restriction to models of the data with a fixed upper bound on k be maintained, even when the available sequence of observations with which to estimate the model becomes unboundedly long? Moreover, even if one grants that a constraint on model complexity requires that no more than some finite number of explanatory variables be stored and used as a basis for forecasts, why must the explanatory variables correspond to the last k observations of the series? In the kind of example in which Fuster *et al.* argue that their proposal predicts over-reaction, more accurate long-horizon forecasts would be possible if the forecast were conditioned on a long moving average of observations, rather than only recent observations; yet tracking a small number of moving averages would seem no more complex than always having access to the last k observations. And above all, the Fuster *et al.* explanation implies that over-reaction should only be observed in the case of variables that are not well-described by an AR(k) process of low enough order. Yet Afrouzi *et al.* (2020) find significant over-reaction in an experiment in which the true data-generating process is an AR(1) process; and in fact, they find the most severe degree of over-reaction (as discussed further below) when the process

to be forecasted is white noise.

Here we offer a different explanation for the pervasiveness of over-reaction. We consider a model in which a decision maker’s forecasts (or more generally, actions with consequences that depend on the future realization of some variable) can be based both on currently observable information and an imperfect memory of past observations. Subject to this constraint on the information that the decision rule can use, we assume that the rule is optimal. Moreover, rather than making an arbitrary assumption about the kind of statistics about past experience that can be recalled with greater or lesser precision, we allow the memory structure to be specified in a very flexible way, and assume that it is optimized for the particular decision problem, subject only to a constraint on the overall complexity of the information that can be stored in (and retrieved from) memory — or more generally, subject to a cost of using a more complex memory structure.

In the limiting case in which the cost of memory complexity is assumed to be negligible, the predictions of our model coincide with those of the rational expectations hypothesis. But when the cost is larger (or the constraint on memory complexity is tighter), our model predicts that forecasts should be both heterogeneous (even in the case of forecasters who observe identical data) and systematically biased. Moreover, the predicted biases include the type of over-reaction to news documented in surveys of forecasts of macroeconomic and financial time series by Bordalo *et al.* (2020a) and in laboratory forecasting experiments by Afrouzi *et al.* (2020). And unlike the model of Fuster *et al.* (2010, 2011), our model predicts that over-reaction to news will be most severe in the case of time series exhibiting little serial correlation.

In seeking to endogenize the information content of the noisy cognitive state on the basis of which people must act, our theory is in the spirit of Sims’s (2003) theory of “rational inattention”; and indeed, we follow Sims in modeling the complexity constraint using information theory. There is nonetheless an important difference between our theory and that of Sims (2003). Sims assumes a constraint on the precision with which new observations of the world can reflect any current or past conditions outside the head of the decision maker, but assumes perfectly precise memory of all of the decision maker’s own past cognitive states, and also assumes that past external states can be observed at any time with the same precision as current conditions. We instead assume (for the sake of simplicity) that the current external state can be observed with perfect precision, but that memory of past cognitive states is subject to an information constraint; and we further assume that the decision maker has no access to external states that occurred in the past, except through (information-constrained) access to her own memory of those past states. These differences are crucial for the ability of our model to explain over-reaction to news.

Other recent papers that explore the consequences of assuming that memory allows only a noisy recollection of past observations include Neligh (2019) and Afrouzi *et al.* (2020). While these authors also assume that some aspects of memory structure are optimized for a particular decision problem, the classes of memory structures that they consider are different than the one that we analyze here.

Both of these papers assume that successive observations of the external state are individually encoded (possibly with variable precision) and stored in memory at the time of the observation; the precision of the record of each observation then evolves over time in a way that is exogenously specified, rather than optimized; and finally, when memory is accessed

later to make a decision, the nature of the signal about the contents of memory may also be endogenized. (Neligh focuses on endogenizing the precision of the initial encoding of individual observations; Afrouzi *et al.* instead focus on endogenizing the precision of what is retrieved from memory.) Both papers take it as given that memory is a vector of separately encoded values of individual observations, and allow no re-encoding of the contents of memory between the time of an initial observation and its retrieval for use in a decision. Our concern is instead with the way in which it is optimal for memory to be constantly re-encoded as time passes, in order to make the most efficient use of a finite limit on the complexity of the stored representations. We comment further on the differences between our framework and those of these other papers below.¹

We present the assumptions of our model and state the optimization problem that we consider in section 1. Section 2 offers a general characterization of the optimal memory structure in our model, showing in particular that it is optimal under our assumptions for the memory state at each point in time to be represented by a single real number, a random variable the mean of which depends on the entire sequence of previous observations. Section 3 illustrates the model’s implications, discussing quantitative aspects of numerical solutions of the model for particular parameter values. We emphasize the failure of beliefs ever to converge to those associated with a rational expectations equilibrium, and show that instead, there are perpetual stationary fluctuations in subjective beliefs similar (though not identical) to those predicted by models of “constant-gain learning” (Evans and Honkapohja, 2001). Finally, section 4 presents the quantitative predictions of the model for statistics of the kind reported by Afrouzi *et al.* (2020) and Bordalo *et al.* (2020a), showing not only that the model can produce over-reaction to news, but that it can be parameterized so as to predict roughly the degree of over-reaction measured by these authors. Section 5 concludes.

1 A Flexible Model of Imprecise Memory

Here we introduce the class of linear-quadratic-Gaussian decision problems that we study, and specify the nature of a general constraint on the precision of memory. This gives rise to a dynamic optimization problem, the solution to which we study below.

1.1 The forecasting problem

In the kind of decision problem which we consider, a decision maker [DM] observes the successive realizations of a univariate stochastic process y_t (“the external state”), which we assume to be a stationary AR(1) process. We write the law of motion of this process in the form

$$y_t = \mu + \rho(y_{t-1} - \mu) + \epsilon_{yt}, \quad (1.1)$$

where μ is the mean value of the external state, ρ is the coefficient of serial correlation (with $|\rho| < 1$), and $\{\epsilon_{yt}\}$ is an i.i.d. sequence, drawn each period from a Gaussian distribution

¹See section 1.2. In addition to considering a different class of possible memory structures, Neligh (2019) addresses largely distinct questions from those analyzed here. Afrouzi *et al.* (2020) instead adopt the explanation that we propose here for their (previously circulated) experimental findings, but show that similar biases are also predicted by a simpler model of noisy memory than the one that we present here.

$N(0, \sigma_\epsilon^2)$. The variance of the external state (conditional on the value of μ and the other parameters) will therefore equal $\sigma_y^2 \equiv \sigma_\epsilon^2 / (1 - \rho^2)$.

The DM's problem is to produce each period a vector of forecasts z_t , so as to minimize the expected value of a discounted quadratic loss function

$$\mathbb{E} \sum_{t=0}^{\infty} \beta^t (z_t - \tilde{z}_t)' W (z_t - \tilde{z}_t), \quad (1.2)$$

where $0 < \beta < 1$, W is a positive definite matrix specifying the relative importance of accuracy of the different dimensions of the vector of forecasts, and the eventual outcomes that the DM seeks to forecast are functions of the future evolution of the external state,²

$$\tilde{z}_t \equiv \sum_{j=0}^{\infty} A_j y_{t+j},$$

where the coefficients $\{A_j\}$ satisfy $\sum_j |A_j| < \infty$. This formalism allows us to assume that the DM may produce forecasts about the future state at multiple horizons (as is typically true in surveys of forecasters, and also in the experiment of Afrouzi *et al.*, 2020). It also allows us to treat cases in which the DM may choose a vector of actions, the rewards from which are a quadratic function of the action vector and the external state in various periods; the problem of action choice to maximize expected reward in such a case is equivalent to a problem of minimizing a quadratic function of the DM's error in forecasting certain linear combinations of the value of the external state at various horizons.³

To simplify our discussion, we assume that the second moments of the stochastic process for the external state are known (more precisely, that the DM's decision rule can be optimized for particular values of these parameters, that are assumed to be the correct ones), while the first moment is not, so that the DM's decision rule must respond adaptively to evidence about the unknown mean value provided by the DM's observations of the state. Thus the values of the parameters ρ and σ_ϵ^2 are assumed to be known, while μ is not; the parameter μ is assumed to be drawn from a non-degenerate prior distribution, $\mu \sim N(0, \Omega)$. Conditional on the value of μ , the initial lagged state y_{-1} is assumed to be drawn from the prior distribution $N(\mu, \sigma_y^2)$, the ergodic distribution for the external state given a value for μ . When we consider the optimality of a possible decision rule for the DM, we integrate over this prior distribution of possible values for μ and y_{-1} , assuming that the decision rule must operate in the same way regardless of which values happen to be true in a given environment.

Note that in the absence of any memory limitation — and given the assumption (maintained in this paper) of perfect observability of the realizations of y_t — it should be possible eventually for the DM to learn the value of μ to arbitrary precision, so that despite our

²Note that the variables denoted $\tilde{z}_t$ are not quantities the value of which is determined at time t ; the subscript t is used to identify the time at which the DM must produce a forecast of the quantity, not the time at which the outcome will be realized. Thus the best possible forecast of $\tilde{z}_t$ at time t , even with full information, would be given by $\mathbb{E}_t \tilde{z}_t$, which will generally not be the same as the realized value $\tilde{z}_t$.

³For example, in a standard consumption-smoothing problem with quadratic consumption utility, the DM's level of expected utility depends on the accuracy with which “permanent income” is estimated at each point in time. This requires the DM to forecast a single variable $\tilde{z}_t$, for which the coefficient A_j is proportional to β^j for all $j \geq 0$.

assumption that the value of μ must be learned, the optimal decision rule should coincide asymptotically with the full-information rational-expectations prediction. We show, however, that this is not true when the precision of memory is bounded.

In any problem of this form (regardless of the assumed memory limitations, which have yet to be specified), the DM's problem can equivalently be formulated as one of simply choosing an estimate $\hat{\mu}_t$ of the unknown mean μ at each date t , based on the information available at the time that z_t must be chosen. It follows from the law of motion (1.1) that

$$E_t \tilde{z}_t = \sum_{j=0}^{\infty} A_j [\mu + \rho^j (y_t - \mu)],$$

where we use the notation $E_t[\cdot]$ for the expected value conditional on the true state at time t , i.e., the value of μ and the history of realizations $(y_0, \dots, y_t)$, even though not all of this information is available to the DM. Conditioning instead on the coarser information set that represents the DM's cognitive state at time t (and noting that this includes precise awareness of the value of y_t), we similarly find that the optimal estimate of $\tilde{z}_t$ will be given by

$$z_t = \sum_{j=0}^{\infty} A_j [\hat{\mu}_t + \rho^j (y_t - \hat{\mu}_t)], \quad (1.3)$$

where $\hat{\mu}_t$ is the expectation of μ conditional on the DM's information set at time t .

We show in the appendix that the DM's expected loss cannot be reduced by restricting attention to a class of decision rules of the form (1.3), under different possible assumptions about how the estimate $\hat{\mu}_t$ is formed.⁴ In the case of any forecasting rule of that kind, the loss function (1.2) is equal to

$$\alpha \cdot \sum_{t=0}^{\infty} \beta^t MSE_t \quad (1.4)$$

plus a term that is independent of the DM's forecasts, where

$$MSE_t \equiv E[(\hat{\mu}_t - \mu)^2]$$

is the mean squared error in estimating μ , and $\alpha > 0$ is a constant that depends on the coefficients $\{A_j\}$ and W . Thus one can equivalently formulate the DM's problem as one of optimal choice of an estimate $\hat{\mu}_t$ each period, so as to minimize MSE_t .

1.2 Memory structures and their cost

We assume that the memory carried into each period $t \geq 0$ can be summarized by a vector m_t of dimension d_t ; the action chosen in period t (i.e., the choice of $\hat{\mu}_t$) must be a function of the cognitive state specified by $s_t = (m_t, y_t)$. The dimension of the memory state is assumed only to be finite, and is not required to be the same for all t . (The case of *perfect memory* can be accommodated by our notation, by assuming that $d_t = t$, and that the elements of the vector m_t correspond to the values $(y_0, y_1, \dots, y_{t-1})$.) We assume that current external

⁴See Appendix A for details of the argument.

state y_t is perfectly observable,⁵ but that behavior can depend on past states only to the extent that memory provides information about them.

We further suppose that the memory state evolves according to a linear law of motion of the form

$$m_{t+1} = \Lambda_t s_t + \omega_{t+1}, \quad \omega_{t+1} \sim N(0, \Sigma_{\omega,t+1}) \quad (1.5)$$

starting from an initial condition of dimension $d_0 = 0$ (that is, s_0 consists only of y_0). However, the dimension d_{t+1} of the memory that is stored, and the elements of the matrices $\Lambda_t, \Sigma_{\omega,t+1}$ are allowed to be arbitrary; we require only that $\Sigma_{\omega,t+1}$ must be positive semi-definite (though it need not be of full rank).

For example, one type of memory structure that this formalism allows us to consider is an “episodic” memory of the kind assumed by Neligh (2019).⁶ In this case, $d_t = t$, and there is an element of m_t corresponding to each of the past observations y_τ for $0 \leq \tau \leq t - 1$ (generalizing the case of perfect memory just discussed). The memory of y_τ at some later time t is given by $m_{\tau+1,t} = y_\tau + u_{\tau+1,t}$, where $u_{\tau+1,t}$ is a Gaussian noise term, independent of the value of y_τ , and with a variance that is necessarily non-decreasing in t . This can be represented by letting $d_t = t$, Λ_t be the identity matrix of dimension $t + 1$, and $\Sigma_{\omega,t+1}$ a diagonal matrix of dimension $n + 1$ (with non-negative elements, but not necessarily of full rank).

Another type of memory that we can consider is one in which only the n most recent past observations of the external state can be recalled, though these are recalled with perfect precision. The requirement that forecasts be functions of the cognitive state would then require them to be functions of $(y_t, y_{t-1}, \dots, y_{t-n})$ for some finite n , as under the hypothesis of “natural expectations” proposed by Fuster, Hébert, and Laibson (2011). This case would correspond to a specification in which $d_t = n$ for all t ; Λ_t is an $n \times (n + 1)$ matrix, the right $n \times n$ block of which is an identity matrix, and the first column of which consists entirely of zeroes; and $\Sigma_{\omega,t+1} = 0$. Our formalism is much more flexible than either of these cases, however, and neither of those specifications turns out to be optimal.

We limit the precision of memory by further assuming that there is a cost of storing and/or accessing the memory state m_{t+1} , that is an increasing function of the Shannon mutual information between the memory state m_{t+1} and the cognitive state s_t about which it provides information.⁷ If I_t is the mutual information between these two random variables, we assume a memory cost in period t of $c(I_t)$, where $c(I)$ is an increasing, convex function. Two extreme possibilities, both of which we consider further below, are the one in which $c(I)$ is a linear function, $c(I) = \theta \cdot I$ for some $\theta > 0$; and the opposite limiting case in which

⁵We might also assume that the current state is observable only imprecisely, as in the model of Sims (2003); but in the present treatment, we simplify the analysis, and highlight the consequences of imperfect memory, by considering the limiting case in which there is no cost of precise observation of the current external state.

⁶Note however that Neligh’s model is not a special case of ours, because in addition to restricting attention to a more special class of memory structures, he assumes a different cost function for precision than the one we propose below.

⁷Mutual information is a non-negative scalar quantity that can be defined for any joint distribution for (s_t, m_{t+1}) , that measures the degree to which the realized value of either random variable provides information about the value of the other (Cover and Thomas, 2006). This measure is used to determine the relative cost of different information structures in the rational inattention theory of Sims (2003); properties of this measure as an information cost function are discussed in Caplin, Dean and Leahy (2019).

there is a finite upper bound $\bar{I}$ on feasible information transmission, with zero cost for any $I < \bar{I}$. Either of these cases allows us to consider a one-parameter family of possible cost functions, ranging from an arbitrarily loose information constraint (θ near zero, $\bar{I}$ very large) to a prohibitively tight one (θ extremely high, $\bar{I}$ near zero).

The cost $c(I_t)$ can equivalently be viewed as either a cost of storing a memory record with information content I_t (that is then available with perfect precision in period $t + 1$), or a cost of retrieving a signal from memory with information content I_t in period $t + 1$ (while the memory stored in period t is taken to have been a perfect record of the period t cognitive state). These two formulations are identical, given that we assume that only the signal m_{t+1} that is retrieved in period $t + 1$ can be stored for future use; thus only the fidelity with which the retrieved memory m_{t+1} reproduces the cognitive state s_t matters. Under the retrieval-cost interpretation, however, our model remains importantly different from the one proposed by Afrouzi *et al.* (2020), in which memory contains a perfect record of all past observations, but there is a cost of retrieving a precise signal about the contents of memory for use in a decision. That model assumes that past observations can be stored indefinitely with perfect precision, with a limit on the precision of recall becoming relevant only when memory must be consulted; this means that it does not predict “recency bias” as ours does.⁸

The memory structure each period, together with the rule for choosing an estimate $\hat{\mu}_t$ as a function of each period’s cognitive state, are then assumed to be chosen so as to minimize total discounted costs

$$\sum_{t=0}^{\infty} \beta^t [\alpha \cdot MSE_t + c(I_t)], \quad (1.6)$$

taking into account both the cost of less accurate forecasts (1.4) and the cost of greater memory precision. Note that no expectation is needed in this objective, since both MSE_t and I_t are functions of the entire joint probability distribution of possible values for $\mu, m_t, y_t, \hat{\mu}_t$ and m_{t+1} .

2 The Optimal Memory Structure

We turn now to a general characterization of the solution to the dynamic optimization problem just posed.

2.1 Implications of linear-Gaussian dynamics

For any memory structure in the class specified in the previous section, the posterior distribution over possible values of $(\mu, y_{-1}, y_0, \dots, y_{t-1})$ implied by memory state m_t will be a multivariate Gaussian distribution. It is thus fully characterized by specifying a finite set of first and second moments of the posterior associated with the memory state. Moreover, the particular memory state m_t at a given date t can be identified by the associated vector of first moments. For the second moments of the posterior are the same for all possible memory states at any time t : they depend on the matrices $\{\Lambda_\tau, \Sigma_{\omega, \tau+1}\}$ for $\tau < t$, but not on the

⁸See the discussion in section 3.4 below.

history of the external state, or on the history of realizations of the memory noise $\{\omega_{t+1}\}$. In what follows, we therefore use the notation m_t for the vector of posterior means.

Among the state variables about which the memory state may convey information, we are particularly interested in the vector of variables $x_t = (\mu, y_{t-1})'$, which are the states determined prior to period t that are relevant for predicting the external state in periods $\tau \geq t$. Let $\bar{m}_t \equiv E[x_t | m_t]$ be the two elements of the memory state that identify the posterior mean of x_t , and let Σ_t be the 2×2 block of second moments of x_t under this same posterior, so that

$$x_t | m_t \sim N(\bar{m}_t, \Sigma_t).$$

And let us furthermore introduce the vectors

$$e'_1 \equiv [1 \ 0], \quad c' \equiv [1 - \rho \ \rho]$$

to select particular elements of this reduced state vector. Then $e'_1 \bar{m}_t$ is the posterior mean and $e'_1 \Sigma_t e_1$ the posterior variance for μ ; while $c' \bar{m}_t$ is the posterior mean and $c' \Sigma_t c$ the posterior variance for the full-information forecast $E_{t-1} y_t$.

The posterior mean and variance for μ after also observing y_t will then be given by the usual Kalman filter formulas,

$$\hat{\mu}_t \equiv E[\mu | s_t] = e'_1 \bar{m}_t + \gamma_{1t} [y_t - c' \bar{m}_t], \quad (2.1)$$

$$\hat{\sigma}_t^2 \equiv \text{var}[\mu | s_t] = e'_1 \Sigma_t e_1 - \gamma_{1t}^2 [c' \Sigma_t c + \sigma_\epsilon^2], \quad (2.2)$$

where the Kalman gain is equal to⁹

$$\gamma_{1t} = \frac{e'_1 \Sigma_t c}{c' \Sigma_t c + \sigma_\epsilon^2}. \quad (2.3)$$

Since y_t is observed precisely, these formulas completely characterize posterior beliefs in cognitive state s_t about the states x_{t+1} that are relevant for forecasting y_τ for all $\tau > t$. Note that $\hat{\sigma}_t^2$ is necessarily positive (complete certainty about the value of μ cannot be achieved in finite time, even in the case of perfect memory), and must satisfy the upper bound

$$\hat{\sigma}_t^2 \leq \hat{\sigma}_0^2 \equiv \frac{\Omega \sigma_y^2}{\Omega + \sigma_y^2}, \quad (2.4)$$

which corresponds to the degree of uncertainty about μ after observing the external state in the case of no informative memory whatsoever (the DM's situation in period $t = 0$).

Then the average losses from inaccurate forecasting in period t are given by

$$MSE_t = \hat{\sigma}_t^2. \quad (2.5)$$

This determines the value of one of the terms in (1.6) as a function of the posterior uncertainty associated with the memory state each period. We note that the optimal estimate $\hat{\mu}_t$ depends only on $\bar{m}_t$ (not other components of the memory state), and that the average loss implied by this estimate depends only on the posterior uncertainty Σ_t associated with those same two components.

⁹We use a 1 subscript in the notation for this variable because it is the first element of a vector of Kalman gains, defined in the more general formula given in Appendix B.

2.2 The sufficiency of memory of a reduced cognitive state

We further show in the appendix¹⁰ that an optimal memory structure makes the memory state m_{t+1} a function only of the “reduced cognitive state”

$$\bar{s}_t \equiv \begin{bmatrix} \hat{\mu}_t \\ y_t \end{bmatrix} = E[x_{t+1} | s_t]. \quad (2.6)$$

We first note (using (2.1) and the fact that y_t is part of the cognitive state) that the elements of $\bar{s}_t$ are a linear function of s_t . Thus we can choose a representation of the vector s_t in which its elements are made up of two parts, $\bar{s}_t$ and $\underline{s}_t$, where the elements of $\underline{s}_t$ are uncorrelated with those of $\bar{s}_t$. We then observe that

$$\bar{m}_{t+1} = E[\bar{s}_t | m_{t+1}].$$

Hence the only aspect of the memory state that matters for $\bar{m}_{t+1}$, and hence for determining both the optimal estimate $\hat{\mu}_{t+1}$ and the reduced cognitive state $\bar{s}_{t+1}$, will be the information that m_{t+1} contains about $\bar{s}_t$.

To the extent that m_{t+1} conveys any information about the elements of $\underline{s}_t$, this information has no consequences for the DM’s estimates $\hat{\mu}_\tau$ in any periods $\tau \geq t+1$, but it increases the mutual information between s_t and m_{t+1} , and hence the information cost $c(I_t)$. Hence under an optimal information structure, the reduced memory state $\bar{m}_t$ must evolve according to a law of motion of the form

$$\bar{m}_{t+1} = \bar{\Lambda}_t \bar{s}_t + \bar{\omega}_{t+1}, \quad (2.7)$$

where $\bar{\omega}_{t+1} \sim N(0, \Sigma_{\bar{\omega}, t+1})$ is distributed independently of the cognitive state. And in addition, the complete memory state must convey no more information about s_t than what is conveyed by the reduced memory state, so that we can without loss of generality assume that m_{t+1} consists solely of $\bar{m}_{t+1}$ (so that $d_{t+1} = 2$ for all $t \geq 0$).

Finally, the 2×2 matrices $\bar{\Lambda}_t$ and $\Sigma_{\bar{\omega}, t+1}$ must satisfy additional restrictions in order for the reduced memory state defined in (2.7) to be consistent with the normalization

$$E[\bar{s}_t | \bar{m}_{t+1}] = \bar{m}_{t+1}. \quad (2.8)$$

We show in the appendix that this relationship will hold if and only if¹¹

$$\Sigma_{\bar{\omega}, t+1} = (I - \bar{\Lambda}_t) X_t \bar{\Lambda}_t', \quad (2.9)$$

where $X_t \equiv \text{var}[\bar{s}_t]$. Note that (2.6) implies that

$$\text{var}[x_{t+1}] = \text{var}[\bar{s}_t] + \text{var}[x_{t+1} | s_t],$$

from which we see that

$$X_t = X(\hat{\sigma}_t^2) \equiv \begin{bmatrix} \Omega - \hat{\sigma}_t^2 & \Omega \\ \Omega & \Omega + \sigma_y^2 \end{bmatrix}. \quad (2.10)$$

¹⁰See Appendix C for details of the argument.

¹¹See the introductory section of Appendix D for details of the argument.

Thus the matrix X_t depends only on the value of $\hat{\sigma}_t^2$. In addition, (2.4) implies that X_t will be positive semi-definite (p.s.d.), and non-singular (hence positive definite) except in the case that $\hat{\sigma}_t^2 = \hat{\sigma}_0^2$ (the case of a totally uninformative memory state m_t).

In order for it to be possible for (2.9) to hold, the matrix $\bar{\Lambda}_t$ must satisfy certain properties: (a) the matrix $\bar{\Lambda}_t X_t = X_t \bar{\Lambda}_t'$ must be symmetric (so that the right-hand side of (2.9) is also symmetric); and (b) the right-hand side of (2.9) must be a p.s.d. matrix. For any symmetric, positive definite 2×2 matrix X_t , we let $\mathcal{L}(X_t)$ be the set of matrices $\bar{\Lambda}_t$ with these properties. Then in addition to assuming that $\bar{\Lambda}_t \in \mathcal{L}(X_t)$, the variance matrix $\Sigma_{\bar{\omega}, t+1}$ must be given by (2.9).

In the special case in which m_t is completely uninformative, $\hat{\mu}_t$ is proportional to the observation y_t , so that there exists a vector $w \gg 0$ such that $\bar{s}_t = w \cdot y_t$. In this case,

$$X_t = X_0 \equiv [\Omega + \sigma_y^2] w w',$$

and we can show that the requirements stated above are satisfied by a matrix $\bar{\Lambda}_t$ if and only if $\bar{\Lambda}_t w = \lambda_t w$ (w is a right eigenvector), with an eigenvalue satisfying $0 \leq \lambda_t \leq 1$. Since the two elements of $\bar{s}_t$ are perfectly collinear in this case, the only part of the matrix $\bar{\Lambda}_t$ that matters for the evolution of the memory state is the implied vector $\bar{\Lambda}_t w$ (which must be a multiple of w). Thus we can without loss of generality impose the further restriction that if $\hat{\sigma}_t^2 = \hat{\sigma}_0^2$, we will describe the dynamics of the memory state using a matrix $\bar{\Lambda}_t$ of the form

$$\bar{\Lambda}_t = \lambda_t \frac{w w'}{w' w}, \quad (2.11)$$

for some $0 \leq \lambda_t \leq 1$. We now adopt this more restrictive definition of the set $\mathcal{L}(X_0)$ in this special case.¹²

We have now shown that the memory structure for period t is completely determined by a specification of a matrix $\bar{\Lambda}_t \in \mathcal{L}(X_t)$. In any period t , the value of $\hat{\sigma}_t^2$ and hence the matrix X_t will be implied by the choice of memory structure for the periods prior to t . Given a choice of $\bar{\Lambda}_t$, the variance-covariance matrix $\Sigma_{\bar{\omega}, t+1}$ is uniquely determined by (2.9). As shown in the appendix,¹³ this then uniquely determines Σ_{t+1} , indicating the degree of uncertainty implied by the memory state m_{t+1} , which then determines $\hat{\sigma}_{t+1}^2$ using (2.2). The degree of uncertainty about μ in the following period is then given by a function of the form

$$\hat{\sigma}_{t+1}^2 = f(\hat{\sigma}_t^2, \bar{\Lambda}_t),$$

that is uniquely defined for any non-negative $\hat{\sigma}_t^2$ satisfying the bound (2.4) and any $\bar{\Lambda}_t \in \mathcal{L}(X(\hat{\sigma}_t^2))$.

Then given that we start from an initial degree of uncertainty $\hat{\sigma}_0^2$ at time $t = 0$ defined by (2.4), we can define the class of sequences $\{\bar{\Lambda}_t\}$ for all $t \geq 0$ with the property that $\bar{\Lambda}_t \in \mathcal{L}(X_t)$ for all $t \geq 0$; let us call this class $\mathcal{L}^{seq}$. Moreover, for any sequence of transition matrices in $\mathcal{L}^{seq}$, we can uniquely define the sequences of values $\{\Sigma_t, \gamma_{1t}, \hat{\sigma}_t^2, X_t\}$ for all $t \geq 0$

¹²Restricting the set of transition matrices $\bar{\Lambda}_t$ that may be chosen in this way has no consequences for the evolution of the memory state, but it makes equation (2.12) below also valid in the case that $X_t = X_0$, and thus it allows us to state certain conditions below more compactly.

¹³See Appendix D.1 for details of the argument.

implied by it. Thus given any sequence $\{\bar{\Lambda}_t\} \in \mathcal{L}^{seq}$, we can calculate the implied sequence of losses $\{MSE_t\}$ from forecast inaccuracy, using (2.5).

We can also uniquely identify the information cost implied by such a sequence of transition matrices, since as shown in the appendix,¹⁴ the mutual information between s_t and m_{t+1} will be given by

$$I_t = I(\bar{\Lambda}_t) \equiv -\frac{1}{2} \log \det(I - \bar{\Lambda}_t) \quad (2.12)$$

each period. Note that the requirement that $\bar{\Lambda}_t \in \mathcal{L}(X_t)$ implies that

$$0 < \det(I - \bar{\Lambda}_t) \leq 1,$$

so that the quantity (2.12) is well-defined, and necessarily non-negative. As the elements of $\bar{\Lambda}_t$ are made small, so that memory ceases to be very informative about the prior cognitive state, $I - \bar{\Lambda}_t$ approaches the identity matrix, and I_t approaches zero. If $\bar{\Lambda}_t$ is varied in such a way as to make one of its eigenvalues approach 1, $I - \bar{\Lambda}_t$ approaches a singular matrix, and $\Sigma_{\hat{\omega}, t+1}$ must approach a singular matrix as well; this means that in the limit, some linear combination of the elements of $\bar{s}_t$ is a random variable with positive variance that comes to be recalled with perfect precision. In this case, $\det(I - \bar{\Lambda}_t)$ approaches zero, so that I_t grows without bound.

Thus a given sequence of transition matrices $\{\bar{\Lambda}_t\}$ uniquely determines sequences $\{MSE_t, I_t\}$, allowing the value of the objective (1.6) to be calculated. The problem of optimal design of a memory structure can then be reduced to the choice of a sequence $\{\bar{\Lambda}_t\} \in \mathcal{L}^{seq}$ so as to minimize (1.6). This objective is necessarily well-defined for any such sequence, since each of the terms is non-negative; the infinite sum will either converge to a finite value, or will diverge, in which case the sequence in question cannot be optimal.¹⁵

2.3 A recursive formulation

We now observe that if for any point in time t , we know the value of $\hat{\sigma}_t^2$ (which depends on the choices made regarding memory structure in periods $\tau < t$), the set of admissible transition matrices $\{\bar{\Lambda}_\tau\}$ for $\tau \geq t$ specifying the memory structure from that time onward will depend only on the value of $\hat{\sigma}_t^2$, and not any other aspect of choices made about the earlier periods. Moreover, any admissible continuation sequence $\{\bar{\Lambda}_\tau\}$ for $\tau \geq t$ implies unique continuation sequences $\{MSE_\tau, I_\tau\}$ for $\tau \geq t$, so that the value of the continuation objective

$$\sum_{\tau=t}^{\infty} \beta^{\tau-t} [\alpha \cdot MSE_\tau + c(I_\tau)] \quad (2.13)$$

will be well-defined.¹⁶

¹⁴See Appendix D.2 for details of the argument.

¹⁵Note that it is clearly possible to choose memory structures for which the infinite sum converges. For example, if one chooses $\bar{\Lambda}_t = 0$ for all $t \geq 0$ (perfectly uninformative memory at all times), $MSE_t = \hat{\sigma}_0^2$ and $I_t = 0$ each period, and (1.6) will equal the finite quantity $\hat{\sigma}_0^2/(1 - \beta)$.

¹⁶Since a finite value for the continuation objective is always possible (see (2.14) below), it is clear that plans that make (2.13) a divergent series cannot be optimal, and can be excluded from consideration.

We can then consider the problem of choosing an admissible continuation plan $\{\bar{\Lambda}_\tau\}$ for $\tau \geq t$ so as to minimize (2.13), given an initial condition for $\hat{\sigma}_t^2$. (This is simply a more general form of our original problem choosing memory structures for all $t \geq 0$ to minimize (1.6), given the initial condition for $\hat{\sigma}_0^2$ specified in (2.4).) Let $V(\hat{\sigma}_t^2)$ be the lowest achievable value for (2.13), as a function of the initial condition $\hat{\sigma}_t^2$; this function is defined for any value of $\hat{\sigma}_t^2$ satisfying the bound (2.4), and is independent of the date t from which we consider the continuation problem. Note that the lower bound necessarily lies in the interval

$$\alpha \hat{\sigma}_t^2 \leq V(\hat{\sigma}_t^2) \leq \alpha \left[\hat{\sigma}_t^2 + \frac{\beta}{1-\beta} \hat{\sigma}_0^2 \right]. \quad (2.14)$$

(The lower bound follows from the fact that $MSE_t = \hat{\sigma}_t^2$, and all other terms in (2.13) must be non-negative; the upper bound is the value of (2.13) if one chooses $\bar{\Lambda}_\tau = 0$ for all $\tau \geq t$, which is among the admissible continuation plans.)

This value function also necessarily satisfies a Bellman equation of the form

$$V(\hat{\sigma}_t^2) = \min_{\bar{\Lambda}_t \in \mathcal{L}(X(\hat{\sigma}_t^2))} [\alpha \hat{\sigma}_t^2 + c(I(\bar{\Lambda}_t)) + \beta V(f(\hat{\sigma}_t^2, \bar{\Lambda}_t))], \quad (2.15)$$

where $I(\bar{\Lambda}_t)$ is the function defined in (2.12). Thus once we know how to compute the value function for arbitrary values of $\hat{\sigma}_{t+1}^2$, the problem of the optimal choice of a memory structure in any period t can be reduced to the one-period optimization problem stated on the right-hand side of (2.15). This indicates how the memory structure for period t must be chosen to trade off the cost $c(I_t)$ of retaining a more precise memory against the continuation loss $V(\hat{\sigma}_{t+1}^2)$ from having access to a less precise memory in period $t+1$.

Let $\mathcal{F}$ be the class of continuous functions $V(\hat{\sigma}_t^2)$, defined for values of $\hat{\sigma}_t^2$ consistent with (2.4), and consistent with the bounds (2.14) everywhere on this domain. Then (2.15) defines a mapping $\Phi : \mathcal{F} \rightarrow \mathcal{F}$: given any conjectured function $V(\hat{\sigma}_{t+1}^2) \in \mathcal{F}$ that is used to evaluate the right-hand side for any value of $\hat{\sigma}_t^2$, the minimized value of the problem on the right-hand side defines a new function $\tilde{V}(\hat{\sigma}_t^2)$ that must also belong to $\mathcal{F}$. Condition (2.15) states that the value function that defines the minimum achievable continuation loss must be a fixed point of this mapping: a function such that $V = \Phi(V)$.

We can further show that for any function $V \in \mathcal{F}$, the function $\Phi(V)$ defined by the right-hand side of (2.15) is necessarily a monotonically increasing function.¹⁷ It follows that the fixed point $V(\hat{\sigma}_t^2)$ must be a monotonically increasing function. Moreover, we can restrict the domain of the mapping Φ to the subset $\mathcal{F}^*$ of increasing functions.

This then provides us with an approach to computing the optimal memory structure for a given parameterization of our model. First, we find the value function $V(\hat{\sigma}^2) \in \mathcal{F}^*$ that is a fixed point of the mapping Φ , by iterating Φ to convergence. Then, given the value function, we can numerically solve the minimization problem on the right-hand side of (2.15) to determine the optimal transition matrix $\bar{\Lambda}_t$ in any period, once we know the value of $\hat{\sigma}_t^2$ for that period. Solution of this problem also allows us to determine the value of $\hat{\sigma}_{t+1}^2 = f(\hat{\sigma}_t^2, \bar{\Lambda}_t)$, so that the entire sequence of values $\{\hat{\sigma}_\tau^2\}$ for all $\tau \geq t$ can be determined once we know $\hat{\sigma}_t^2$. Finally, we recall that for the initial period $t = 0$, the value of $\hat{\sigma}_0^2$ is given

¹⁷See Appendix E.1 for a proof.

by (2.4); we can thus solve for the entire sequence $\{\hat{\sigma}^2\}$ for all $t \geq 0$ by integrating forward from this initial condition.

Once we have determined the sequence of values $\{\hat{\sigma}_t^2\}$ implied by an optimal memory structure for each period, we can determine the elements of the matrix $X_t = X(\hat{\sigma}_t^2)$ each period, using (2.10). We show in the appendix¹⁸ that the degree of uncertainty at the beginning of any period given the structure of the memory chosen for the previous period is given by

$$\Sigma_{t+1} = \Sigma_0 - X_t \bar{\Lambda}_t'.$$

This in turn allows us to calculate the DM's optimal estimate $\hat{\mu}_t$ each period, as a function of the history of realizations $\{y_\tau\}$ of the external state for all $0 \leq \tau \leq t$ and the history of realizations of the DM's memory noise $\{\tilde{\omega}_{\tau+1}\}$ for all $0 \leq \tau \leq t-1$, using (2.1). The DM's complete vector of forecasts z_t each period is then given by (1.3).

2.4 Optimality of a unidimensional memory state

We can show further that the optimal memory state must have a one-dimensional representation. This simplifies the computational formulation of the optimization problem on the right-hand side of (2.15), and provides further insight into the nature of an optimally imprecise memory. Although the information contained in the cognitive state s_t that is relevant for predicting (at time t) what actions will be desirable for the DM in later periods is two-dimensional (both elements of $\bar{s}_t$ matter, if $\rho > 0$, and except when memory is completely uninformative, these are not perfectly correlated with each other), we find that it is optimal for the DM's memory to include only a noisy record of a single linear combination of the two variables. Moreover, this is true regardless of how small memory costs may be.

There is in fact a fairly simple intuition for the result. Note that in any period t , the Kalman filter (2.1) implies that the optimal estimate of the unknown value of μ will be given by a linear function of elements of the cognitive state of the form

$$\hat{\mu}_t = \zeta_t + \delta_t' \bar{m}_t. \quad (2.16)$$

It follows from this that the only information in the memory state m_t that matters for the estimate $\hat{\mu}_t$ is the single quantity $\delta_t' \bar{m}_t$.

We can establish the optimality of a unidimensional memory in the following way. Consider the optimization problem on the right-hand side of (2.15) in any period t , given the degree of uncertainty $\hat{\sigma}_t^2$ determined by the memory structures chosen in earlier periods. The fact that $V(\hat{\sigma}_{t+1}^2)$ is an increasing function, and that $c(I_t)$ is at least weakly increasing, means that an optimal memory structure must minimize the mutual information I_t given the uncertainty $\hat{\sigma}_{t+1}^2$ that it implies for the following period.¹⁹ Hence the optimal choice for

¹⁸See Appendix D.1 for details of the argument.

¹⁹In the case that $c(I)$ is constant over some interval, reducing I_t need not reduce $c(I_t)$, but it cannot increase it; thus the solution to the problem (2.17) must be among the solutions to the problem on the right-hand side of (2.15), even if it is not a unique solution. In such case, showing that the solution to (2.17) is necessarily a singular matrix suffices to show that we can without any loss impose the further constraint in (2.15) that the matrix $\bar{\Lambda}_t$ must be at most of rank one.

$\bar{\Lambda}_t$ must solve the problem

$$\min_{\bar{\Lambda}_t \in \mathcal{L}(X(\hat{\sigma}_t^2))} I(\bar{\Lambda}_t) \quad \text{s.t.} \quad f(\hat{\sigma}_t^2, \bar{\Lambda}_t) \leq \hat{\sigma}_{t+1}^2, \quad (2.17)$$

for given values of $(\hat{\sigma}_t^2, \hat{\sigma}_{t+1}^2)$. We shall show that whenever $(\hat{\sigma}_t^2, \hat{\sigma}_{t+1}^2)$ are such that the set of matrices satisfying the constraint in (2.17) is non-empty,²⁰ the solution $\bar{\Lambda}_t$ to this problem must be at most of rank one. Thus it must be of the special form

$$\bar{\Lambda}_t = \lambda_t X_t v_t v_t', \quad (2.18)$$

where λ_t is a scalar satisfying $0 \leq \lambda \leq 1$ and v_t is a vector normalized to satisfy $v_t' X_t v_t = 1$. It follows that in each period $\bar{m}_{t+1} = X_t v_t \tilde{m}_{t+1}$, where $\tilde{m}_{t+1}$ is a unidimensional memory state with a law of motion

$$\tilde{m}_{t+1} = \lambda_t v_t' \bar{s}_t + \tilde{\omega}_{t+1}, \quad \tilde{\omega}_{t+1} \sim N(0, \lambda_t(1 - \lambda_t)). \quad (2.19)$$

If $\hat{\sigma}_t^2 = \hat{\sigma}_0^2$, the set $\mathcal{L}(X_0)$ consists only of matrices of the form (2.18), with

$$v_t = \frac{w}{(\Omega + \sigma_y^2)^{1/2}(w'w)}, \quad (2.20)$$

because of (2.11). Hence the asserted result is obviously true in that case. Suppose instead that $\hat{\sigma}_t^2 < \hat{\sigma}_0^2$, and consider any matrix $\bar{\Lambda}_t \in \mathcal{L}(X(\hat{\sigma}_t^2))$ that satisfies the constraint in (2.17). If $\bar{\Lambda}_t$ is not itself of rank one (or lower), we shall show that we can choose an alternative transition matrix of the form (2.18), that is also consistent with the constraint in (2.17), but which achieves a lower value of $I(\bar{\Lambda}_t)$.

Let the alternative transition matrix be given by (2.18), with

$$\lambda_t = \frac{\delta'_{t+1} \bar{\Lambda}_t X_t \bar{\Lambda}_t' \delta_{t+1}}{\delta'_{t+1} X_t \bar{\Lambda}_t' \delta_{t+1}}, \quad v_t = \frac{\bar{\Lambda}_t' \delta_{t+1}}{(\delta'_{t+1} \bar{\Lambda}_t X_t \bar{\Lambda}_t' \delta_{t+1})^{1/2}},$$

where $\delta_{t+1} \equiv e_1 - \gamma_{1,t+1}c$ is the vector introduced in (2.16), and let the matrix $\Sigma_{\bar{\omega},t+1}$ be correspondingly modified in the way specified by (2.9). We show in the appendix²¹ that this specification implies that $0 \leq \lambda_t \leq 1$, so that this alternative matrix also belongs to $\mathcal{L}(X_t)$. Moreover, the new memory structure implies a conditional distribution

$$\delta'_{t+1} \bar{m}_{t+1} | s_t \sim N(\delta'_{t+1} \bar{\Lambda}_t \bar{s}_t, \delta'_{t+1} \Sigma_{\bar{\omega},t+1} \delta_{t+1})$$

that is the same as under the original memory structure. This implies that the optimal estimate $\hat{\mu}_{t+1}$ conditional on the cognitive state s_{t+1} will be the same function of $\bar{m}_{t+1}$ and y_{t+1} in the case of the new memory structure, and that the conditional distribution $\hat{\mu}_{t+1} | s_t, y_{t+1}$ will be the same. It follows that $\hat{\sigma}_{t+1}^2$ will be the same, so that the alternative transition matrix also satisfies the constraint in (2.17).

At the same time, we show in the appendix that the reduction in the complexity of memory cannot increase information costs in any period.²² The new memory structure

²⁰Note that this must be the case if $\hat{\sigma}_{t+1}^2$ is chosen optimally given $\hat{\sigma}_t^2$.

²¹See Appendix E.2 for details of the argument.

²²See Appendix E.2 for details of the argument.

consists effectively of a scalar memory state $\tilde{m}_{t+1}$ in each period, which is a multiple of $d'_{t+1}\tilde{m}_{t+1}$, a particular linear combination of the elements of the memory state under the previous memory structure. Hence the information about $\bar{s}_t$ that is revealed by m_{t+1} under the new memory structure (i.e., that is revealed by $\tilde{m}_{t+1}$) is also information that was revealed by $\tilde{m}_{t+1}$ under the previous memory structure; thus the value of I_t under the previous memory structure must have been at least as large as under the new memory structure. In fact, the only case in which the mutual information will not be reduced by the proposed modification of the memory structure is if all elements of $\tilde{m}_{t+1}$ were multiples of $d'_{t+1}\tilde{m}_{t+1}$; which is to say, only if $\bar{\Lambda}_t$ were already of the special form (2.18).

We conclude, then, that an optimal memory structure must involve a transition matrix in every period of the special form (2.18), so that the memory state each period can be represented by a scalar quantity $\tilde{m}_t$. The choice of memory structure can then be reduced to a problem of choosing, in each period $t \geq 0$, a scalar quantity $0 \leq \lambda_t \leq 1$, and the direction of a vector v_t (the length of which will then be chosen each period so as to ensure that $v'_t X_t v_t = 1$); the values chosen for these quantities then determine the law of motion for the unidimensional memory state $\tilde{m}_{t+1}$, specified by (2.19). This in turn determines the elements of the matrix Σ_{t+1} , and hence the value of the gain coefficient $\gamma_{1,t+1}$ in the Kalman filter formula (2.1) and the value of $\hat{\sigma}_{t+1}^2$, which determines the matrix $X_{t+1} = X(\hat{\sigma}_{t+1}^2)$.

For any value $0 \leq \hat{\sigma}_t^2 < \hat{\sigma}_0^2$, let $\mathcal{V}(\hat{\sigma}_t^2)$ be the set of vectors v_t satisfying $v'_t X(\hat{\sigma}_t^2) v_t = 1$. In the case that $\hat{\sigma}_t^2 = \hat{\sigma}_0^2$, we add the further stipulation that $\mathcal{V}(\hat{\sigma}_0^2)$ consists only of the single vector (2.20). Then given a value for $\hat{\sigma}_t^2$, determined by the memory structures for periods $\tau < t$, the memory structure for period t is specified by a scalar quantity $0 \leq \lambda_t \leq 1$ and a vector $v_t \in \mathcal{V}(\hat{\sigma}_t^2)$. These determine a value for $\hat{\sigma}_{t+1}^2 = f(\hat{\sigma}_t^2, \lambda_t, v_t)$, where now the function f is defined for any values of its arguments satisfying $0 \leq \hat{\sigma}_t^2 \leq \hat{\sigma}_0^2$, $0 \leq \lambda_t \leq 1$, and $v_t \in \mathcal{V}(\hat{\sigma}_t^2)$.

Because of the monotonicity of the value function $V(\hat{\sigma}_{t+1}^2)$, the optimal weight vector v_t in any period must be the one that solves the static optimization problem

$$\bar{f}(\hat{\sigma}_t^2, \lambda_t) \equiv \min_{v_t \in \mathcal{V}(\hat{\sigma}_t^2)} f(\hat{\sigma}_t^2, \lambda_t, v_t). \quad (2.21)$$

In the appendix,²³ we give an explicit algebraic solution for the optimal v_t for any given values $0 \leq \hat{\sigma}_t^2 \leq \hat{\sigma}_0^2$ and $0 < \lambda_t \leq 1$,²⁴ and hence for the function $\bar{f}(\hat{\sigma}_t^2, \lambda_t)$. The latter function is also defined when $\lambda_t = 0$, and easily seen to equal $\bar{f}(\hat{\sigma}_t^2, 0) = \hat{\sigma}_t^2$. Thus we can solve for the dynamics of $\{\hat{\sigma}_t^2\}$ implied by any sequence $\{\lambda_t\}$, by iterating the law of motion

$$\hat{\sigma}_{t+1}^2 = \bar{f}(\hat{\sigma}_t^2, \lambda_t),$$

starting from the initial condition $\hat{\sigma}_0^2$ defined in (2.4).

Moreover, it follows from (2.12) that the mutual information associated with the period t memory structure will be given by

$$I_t = -\frac{1}{2} \log(1 - \lambda_t). \quad (2.22)$$

The Bellman equation (2.15) can therefore be written in the simpler form

$$V(\hat{\sigma}_t^2) = \min_{0 \leq \lambda_t \leq 1} [\alpha \hat{\sigma}_t^2 + c(-(1/2) \log(1 - \lambda_t)) + \beta V(\bar{f}(\hat{\sigma}_t^2, \lambda_t))]. \quad (2.23)$$

²³See Appendix E.3 for details.

²⁴Note that no solution is needed in the case that $\lambda_t = 0$, since in this case v_t is undefined.

3 Features of the Model Solution

Here we provide numerical examples of solutions for an optimal memory structure, under alternative assumptions about both the degree of persistence of the process that must be forecasted and the nature of the information cost function. In reporting our results, it is useful to describe the model solution in terms of scale-invariant quantities — that is, ones that are independent of the value of σ_y , indicating the amplitude of the transitory fluctuations in the external state around its mean. Thus we parameterize the degree of prior uncertainty about μ not in terms a value for Ω (the variance of the prior distribution for μ), but rather by a value for $K \equiv \Omega/\sigma_y^2$ (the variance of the prior distribution for μ/σ_y). We similarly measure the degree of uncertainty about μ conditional on the cognitive state at a given point in time (i.e., after a given amount of experience) not in terms of the value of $\hat{\sigma}_t^2$, but rather by the scaled uncertainty measure $\eta_t \equiv \hat{\sigma}_t^2/\sigma_y^2$.

In terms of this scaled uncertainty measure, an optimal memory structure minimizes the value of

$$\sum_{t=0}^{\infty} \beta^t [\eta_t + \tilde{c}(I_t),]$$

a scaled version of (1.6), where the scaled cost function is defined as $\tilde{c}(I) \equiv c(I)/(\alpha\sigma_y^2)$. (Dividing by α further reduces the numbers of parameters that we need to specify in considering the different possible forms that the optimal memory structure may take, since it is only the relative weights on the two loss terms in the objective (1.6) that matter for the optimal memory structure.) Our scale-invariant model is then completely specified by values for the parameters ρ, β, K and the scaled cost function $\tilde{c}(I)$. In our quantitative analysis, we assume that each “period” of our discrete-time model corresponds to a quarter of a year (the variable to be forecasted is a quarterly time series), and hence set $\beta = 0.99$ (implying a discount rate of 4 percent per annum). We consider a variety of values $0 \leq \rho < 1$ for the assumed degree of serial correlation of the external state, and explore the effects of different assumptions regarding the degree of prior uncertainty and the size of information costs.

3.1 The case of a fixed per-period bound on mutual information

We begin by considering the case in which $\tilde{c}(I) = 0$ for all $I \leq \bar{I}$, but the cost is infinite for any value $I > \bar{I}$. (That is, there is a fixed upper bound on the possible mutual information between s_t and m_{t+1} in each period; but any memory structure consistent with this bound is equally feasible and has the same cost.) Here $\bar{I}$ is some finite positive quantity.

Solution for the optimal memory structure is particularly simple in this case. Because of (2.22), the per-period bound on mutual information can equivalently be written as an upper bound $\lambda_t \leq \bar{\lambda}$, where

$$\bar{\lambda} \equiv 1 - e^{-2\bar{I}} > 0.$$

The optimal memory structure in period t is then characterized by the λ_t that minimizes $\bar{f}(\hat{\sigma}_t^2, \lambda_t)$ subject to this constraint. We show in the appendix²⁵ that the minimizing value of λ_t is necessarily the largest feasible value; hence in the solution to this problem, $\lambda_t = \bar{\lambda}$, the value determined by the per-period information bound.

²⁵See Appendix F.1 for details of the argument.

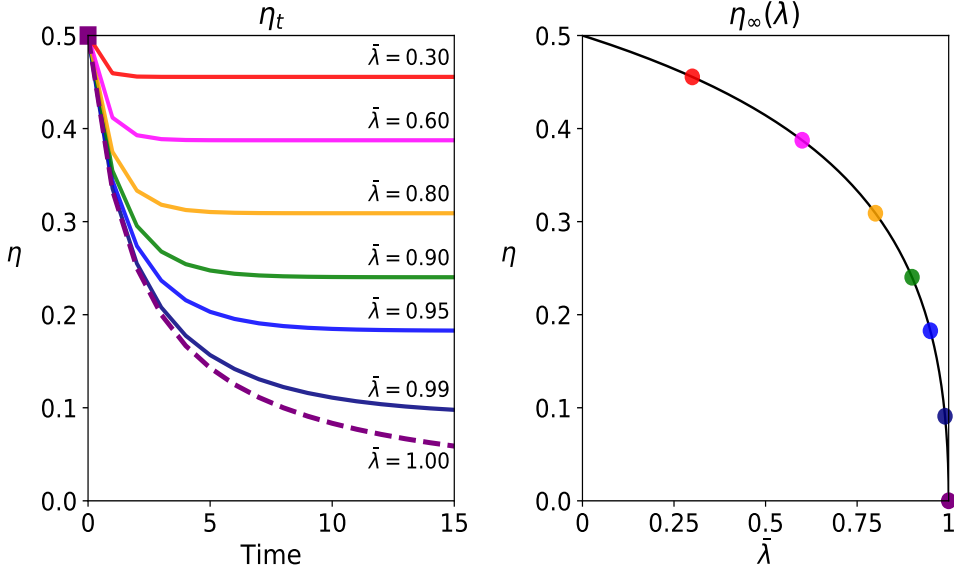


Figure 1: The evolution of scaled uncertainty about μ as the number t of previous (imperfectly remembered) observations grows. The right panel shows the long-run value of scaled uncertainty (to which η_t converges as $t \rightarrow \infty$) as a function of the constraint on the complexity of memory, parameterized by $\bar{\lambda}$.

The dynamics of the uncertainty measure are then given by $\hat{\sigma}_{t+1}^2 = \bar{f}(\hat{\sigma}_t^2, \bar{\lambda})$. In terms of the rescaled variables, the law of motion can be written as

$$\eta_{t+1} = \phi(\eta_t; \bar{\lambda}), \quad (3.1)$$

where $\phi(\eta; \bar{\lambda})$ is a function that is independent of the scale factor σ_y .²⁶

For any value of $\bar{\lambda}$ indicating the tightness of the constraint on the complexity of memory, equation (3.1) indicates how the DM's degree of uncertainty about μ evolves as additional observations of the external state are made. Starting from the initial condition $\eta_0 = K/(K+1)$ implied by (2.4), the law of motion (3.1) can be iterated to obtain a unique solution for the complete sequence of values $\{\eta_t\}$ for all $t \geq 0$. In the limiting case $\bar{\lambda} = 1$ (unlimited memory), the law of motion (3.1) takes the especially simple form

$$\frac{1}{\eta_{t+1}} = \frac{1}{\eta_t} + \frac{1-\rho}{1+\rho}. \quad (3.2)$$

This is simply the standard result for the linear growth in posterior precision under Bayesian updating as additional observations are made; it has the implication that η_t declines monotonically, and converges to zero for large t . Thus in the case of perfect memory, the DM should eventually learn the value of μ with perfect precision, and hence make forecasts of the kind implied by the hypothesis of rational expectations.

When $\bar{\lambda} > 0$, instead, the law of motion (3.1) implies that η_t should decrease initially, as even imprecise memory of the DM's observations of the external state reduces uncertainty

²⁶See Appendix F.1 for an explicit algebraic solution for this function.

to some degree, but that η_t remains bounded away from zero, and converges to a value $\eta_\infty(\bar{\lambda}) > 0$. This is illustrated in Figure 1, which shows the dynamics implied by (3.1) for each of several different values of $\bar{\lambda}$, in the case that $\rho = 0$ and $K = 1$.²⁷ The left panel plots the sequence of values $\{\eta_t\}$ implied by (3.1) for a given value of $\bar{\lambda}$. (Note that the initial value η_0 is the same in each case.) The right panel shows the value η_∞ to which the sequence converges as t grows; this value depends on $\bar{\lambda}$, and the functional relationship between $\bar{\lambda}$ and this limiting degree of uncertainty can be described by a function $\eta_\infty(\bar{\lambda})$, plotted as a smooth curve in the right panel of the figure.

In the case that $\bar{\lambda} = 1$ (shown as a dashed curve in the left panel of Figure 1), the sequence $\{\eta_t\}$ decreases monotonically to zero at the rate predicted by the difference equation (3.2). But for any number of prior observations $t > 0$, the value of η_t remains higher the lower is $\bar{\lambda}$ (that is, the tighter the memory constraint), and the long-run degree of uncertainty about μ , measured by η_∞ , is a decreasing function of $\bar{\lambda}$ as well, as shown by the curve in the right panel of the figure. Because of the limit on the amount of information that can be retained in memory, the DM's uncertainty about the value of μ never falls below a certain level, even in the long run, despite our assumption that the value of μ is fixed for all time. We further observe that the long-run degree of uncertainty η_∞ is larger, the smaller is $\bar{\lambda}$ (that is, the tighter the constraint on memory). In the limit as $\bar{\lambda}$ approaches zero (corresponding to a constraint that memory must be completely uninformative), the long-run degree of uncertainty η_∞ approaches the prior degree of uncertainty $\eta_0 = K/(K + 1)$.

As η_t falls along one of these trajectories, the weight vector v_t that solves the problem (2.21) shifts as well. As η_t converges to the long-run value η_∞ , the optimal weight vector v_t similarly converges to a long-run value v_∞ , indicating the particular linear combination of $\hat{\mu}_t$ and y_t that is imprecisely recorded in memory each period. Associated with this stationary long-run memory structure there will also be a stationary long-run value for the Kalman gain coefficient γ_1 in equation (2.1), and more generally, stationary values for the coefficients of the linear difference equations describing the joint dynamics $\{y_t, \tilde{m}_t\}$ of the external state and the memory state.

These long-run stationary coefficients will depend on the value of $\bar{\lambda}$ (indicating the tightness of the memory constraint) and also on the value of ρ (indicating the degree of persistence of the fluctuations in the external state). Figure 2 indicates how variation in each of these parameters affects several of the long-run stationary coefficients.²⁸ In each panel, a curve shows how the coefficient in question varies as a function of ρ (for values of ρ between 0.0 and 0.9), for a given value of $\bar{\lambda}$; curves of this kind are shown for each of three different values of $\bar{\lambda}$, ranging between $\bar{\lambda} = 0.95$ (in which case memory is relatively precise) and $\bar{\lambda} = 0.30$ (in which case it is much more constrained). All of the curves shown in Figure 2 are again for the case of prior uncertainty $K = 1$.

The upper-right panel of the figure shows the long-run direction vector v_∞ ; the quantity reported on the vertical axis is the long-run value of the ratio v_2/v_1 of the vector's two

²⁷The effects of variation in the parameters ρ and K are illustrated in additional figures shown in Appendix F.1. We use the parameterization $K = 1$ in the figures shown in the text because this value allows a reasonably good fit of the predictions shown in Figure 7 below with the experimental evidence reported by Afrouzi *et al.* (2020).

²⁸See Appendix G.1 for the formulas used to calculate each of the coefficients plotted here as functions of the model parameters.

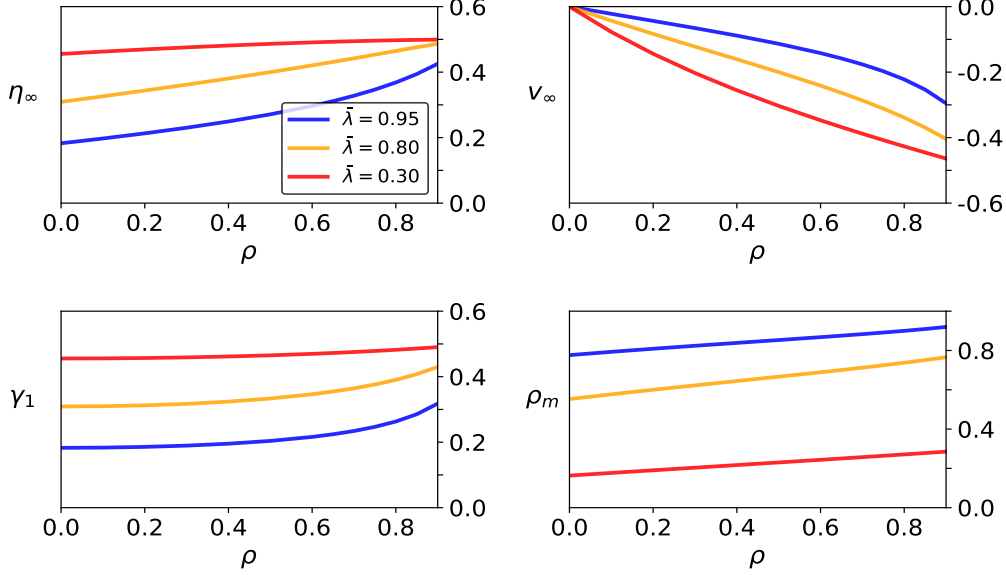


Figure 2: Coefficients describing the optimal memory structure in the long run, as a function of the degree of persistence ρ of the external state, for alternative values of $\bar{\lambda}$. Respective panels show the long-run values for η (measuring uncertainty about μ), the direction vector v (indicating the content of the memory state), the Kalman gain γ_1 (for updating the DM's estimate of μ), and ρ_m (measuring the intrinsic persistence of fluctuations in the memory state).

components.²⁹ Thus a value of -0.3 for this quantity means that the univariate memory state $\tilde{m}_{t+1}$ is (up to a multiplicative factor that does not affect its information content) equal to the value of $\hat{\mu}_t - 0.3y_t$, plus additive Gaussian noise. The figure shows that when $\rho = 0$, the optimal univariate memory state involves $v_2 = 0$; that is, only the current estimate $\hat{\mu}_t$ of the unknown mean is remembered with noise, with the current observation y_t being completely forgotten. This is optimal because when $\rho = 0$, the current value y_t contains no information that is relevant for improving subsequent forecasts of the external state, except to the extent that it helps to improve the DM's estimate of μ (which information is already reflected in the estimate $\hat{\mu}_t$). Instead, when the external state is serially correlated, it is optimal to commit to memory a linear combination of $\hat{\mu}_t$ and the current state y_t ; in the case that $\rho > 0$, the optimal linear combination puts a negative relative weight on y_t , to an extent that is greater the greater the degree of serial correlation, and greater the tighter the constraint on memory.

The upper-left panel of the figure shows the long-run degree of uncertainty about μ , measured by η_∞ . As shown in Figure 1, when $\rho = 0$, η_∞ is a decreasing function of $\bar{\lambda}$. We see in Figure 2 that this is also true when $\rho > 0$. However, for a given memory constraint $\bar{\lambda}$, the long-run value η_∞ is also an increasing function of ρ , with the degree of increase

²⁹This information (together with the value of η_∞ given in the upper left panel) suffices to completely determine the vector v_t , since the vector is normalized so that $v'Xv = 1$. The value of λ (given by the constraint $\bar{\lambda}$), the matrix X (determined by the value of η_∞), and the vector v then completely determine the long-run stationary elements of the matrix $\bar{\Lambda}$ (using (2.18)) and hence also of the matrix Σ_ω (using (2.9)); thus the dynamics of the memory state given by (2.7) are completely specified.

when the external state is highly persistent being particularly notable when memory is more accurate. The greater the serial correlation of the state, the fewer the effective number of independent noisy observations of μ that the DM receives over any finite time period; thus even under perfect Bayesian updating, equation (3.2) indicates that the rate at which precision is increased by each additional observation is smaller the larger is ρ . In the case of perfect memory, the long-run degree of uncertainty about μ is nonetheless zero (there is simply slower convergence to that long-run value when ρ is large); but with moderately imperfect memory, the effective amount of experience that can ever be drawn upon remains bounded, so that the uncertainty about μ remains larger forever when ρ is larger. When memory is even more imperfect, not much more than one observation (the most recent one) can be used in any event, so that the value of η_∞ is in this case less sensitive to the value of ρ .

In the long run, the dynamics of the cognitive state $\bar{s}_t$ and the memory state $\bar{m}_{t+1}$ are described by linear equations with constant coefficients. The lower-left panel of Figure 2 shows the long-run value for the Kalman gain γ_{1t} in (2.1). With imperfect memory, this is always a quantity between 0 and 1, meaning that a higher value of the current state y_t raises the DM's estimate of the value of μ , though by less than the amount of the increase in y_t . For a given value of ρ , the Kalman gain is larger the tighter the constraint on memory; in the limit as $\bar{\lambda} \rightarrow 1$ (perfect memory), the long-run value of this coefficient approaches zero (as the true value of μ is eventually learned), while in the limit as $\bar{\lambda} \rightarrow 0$ (no memory), the value approaches a maximum value that is still less than one (because of the DM's finite-variance prior).

Finally, the lower-right panel reports the long-run value of ρ_m , a measure of the intrinsic persistence of the memory state. The impulse response function for the effect of a memory-noise innovation $\tilde{\omega}_t$ on the subsequent path of the univariate memory state $\tilde{m}_\tau$ is proportional to $(\rho_m)^{\tau-t}$ for all $\tau \geq t$;³⁰ thus the value of ρ_m indicates the rate of exponential decay of the memory state back to its long-run average value. A smaller value of ρ_m means that the contents of memory decay more rapidly; for any value of ρ , the figure shows that ρ_m is smaller, the tighter the memory constraint. At the same time, while a larger value of ρ_m implies that memory persists for a longer time, it also implies that when memory noise creates an erroneous impression of prior experience, this bias in what is recalled about is also corrected more slowly; thus the value of ρ_m is an important determinant of the predicted persistence of forecast bias.

3.2 The case of a linear cost of information

Analysis of the model is more complex when instead the amount of information stored in memory each period can be increased at some finite cost. As an illustration we consider the polar opposite case in which $\tilde{c}(I)$ is a linear function of I , so that the marginal cost of a further increase in the mutual information is independent of how large it already is. Thus we assume that $\tilde{c}(I) = \tilde{\theta} \cdot I$, for some coefficient $\tilde{\theta} > 0$ which parameterizes the cost of memory.

³⁰Here we refer to the difference that the realization of $\tilde{\omega}_t$ makes for the forecasts of $\tilde{m}_\tau$ at different horizons $\tau \geq t$, by an observer who knows the true value of μ and the DM's cognitive state at time $t - 1$, in addition to observing the realization of $\tilde{\omega}_t$. See Appendix G.1 for details of the calculation.

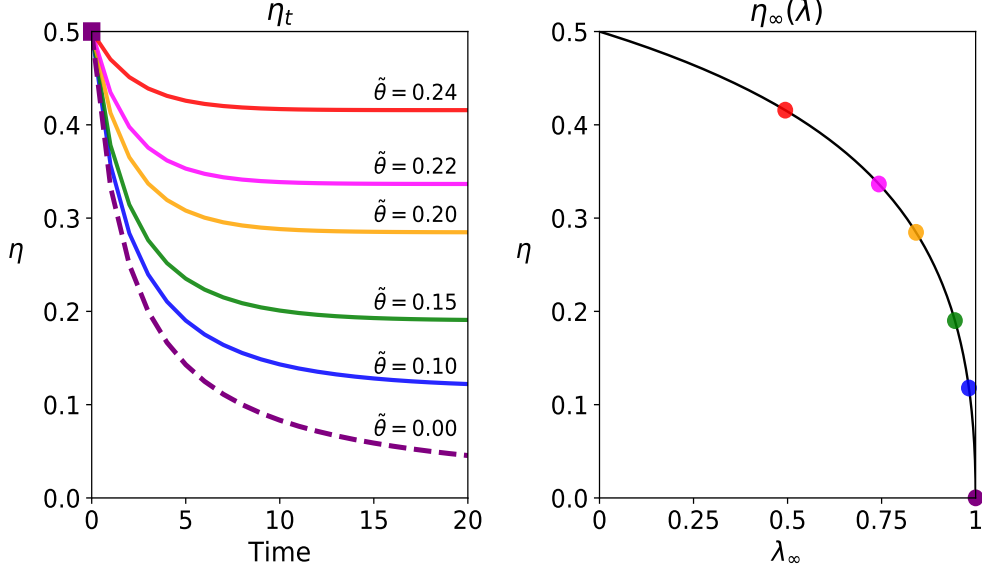


Figure 3: The evolution of scaled uncertainty about μ as the number t of previous (imperfectly remembered) observations grows, now for the case of a linear cost of memory complexity. The right panel shows the long-run value of scaled uncertainty for each value of the cost parameter $\tilde{\theta}$, plotted as a point on the same locus of optimal long-run memory structures as in Figure 1.

In this case, the optimal choice of λ_t in any period will depend on the value of reducing uncertainty in the following period. We note that the value function $V(\hat{\sigma}_{t+1}^2)$ appearing in the Bellman equation (2.23) can be written as $\sigma_y \cdot \tilde{V}(\eta_{t+1})$, where η_{t+1} is the scaled uncertainty measure and the function $\tilde{V}(\eta)$ is independent of the scale factor σ_y (for given values of the parameters K, ρ, β and $\tilde{\theta}$). We can then write the Bellman equation in the scale-invariant form

$$\tilde{V}(\eta_t) = \min_{0 \leq \lambda_t \leq 1} \left\{ \eta_t - \frac{\tilde{\theta}}{2} \log(1 - \lambda_t) + \beta \tilde{V}(\phi(\eta_t; \lambda_t)) \right\}. \quad (3.3)$$

The optimal choice of λ_t in any period will be the value that solves the problem on the right-hand side of (3.3). This problem has a solution $\lambda_t = \lambda^*(\eta_t)$ which depends only on the value of η_t , the degree of uncertainty in period t determined by the memory structures chosen for periods prior to t .

Thus we can solve for the optimal policy function $\lambda^*(\eta_t)$ once we know the value function $\tilde{V}(\eta_{t+1})$, and we can solve numerically for the value function by iterating the Bellman equation (3.3), as discussed further in the appendix.³¹ The policy function $\lambda_t = \lambda^*(\eta_t)$ together with the law of motion

$$\eta_{t+1} = \phi(\eta_t; \lambda_t) \quad (3.4)$$

derived earlier can then be solved for the dynamics of the scaled uncertainty $\{\eta_t\}$ for all $t \geq 0$, starting from the initial condition $\eta_0 = K/(K+1)$.³² The dynamics of scaled uncertainty as

³¹See Appendix F.2 for details.

³²See Appendix F.2 for further discussion of the implied dynamics.

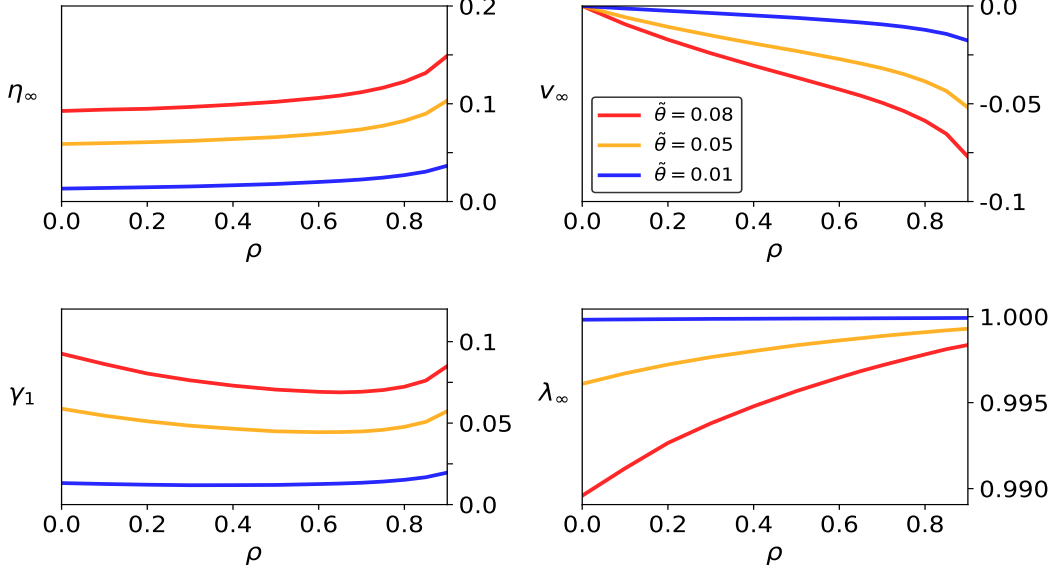


Figure 4: Coefficients describing the optimal memory structure in the long run, as a function of the degree of persistence ρ of the external state, in the case of a linear memory cost function, for alternative values of $\tilde{\theta}$. Respective panels show the long-run values for η , the direction vector v , the Kalman gain γ_1 , and the memory precision coefficient λ .

a function of the number of observations t are shown for progressively larger values of $\tilde{\theta}$ in Figure 3, using the same format as in Figure 1. Once again, we see that while uncertainty about μ eventually falls to zero as a result of when there is no cost of memory complexity, as long as the cost is positive, the value of η_t remains bounded away from zero, and converges asymptotically to a value η_∞ that is higher the higher the cost of memory complexity.

Associated with such an asymptotic degree of uncertainty is a particular long-run memory structure $(\lambda_\infty, v_\infty)$, which will imply a particular long-run value for the Kalman gain γ_1 . The way in which the long-run values of these different quantities vary with different assumptions about the values of ρ and $\tilde{\theta}$ is illustrated in Figure 4. (We use the same convention as in Figure 2 to indicate the direction of the vector v_∞ in the upper-right panel of the figure.) As we vary ρ for a given value of $\tilde{\theta}$, the associated value of λ_∞ changes; hence the fixed- $\tilde{\theta}$ curves shown in Figure 4 do not correspond exactly to any of the fixed- λ curves plotted in Figure 2, even though each of the long-run memory structures associated with a pair $(\rho, \tilde{\theta})$ is identical to the long-run memory structure associated with some pair $(\rho, \bar{\lambda})$. As shown in the lower-right panel of the figure, the optimal λ_∞ rises as ρ increases, for any value of the cost parameter $\tilde{\theta} > 0$; the more persistent the external state that must be forecasted, the more it becomes worthwhile to pay a larger information cost in order to retain a more precise memory of prior experience.

Not surprisingly, we observe that for any value of ρ , increasing the memory cost $\tilde{\theta}$ makes it optimal for the long-run precision of memory λ_∞ to be smaller, and consequently for the long-run degree of uncertainty about μ to be larger. In the case of a sufficiently high value of θ , it will be optimal for memory to be completely uninformative. In fact, this happens for a finite value of $\tilde{\theta}$, and it occurs abruptly, rather than through a gradual increase in the long-run degree of uncertainty η_∞ toward the limiting value of $\eta_0 = K/(K + 1)$ as $\tilde{\theta}$ is

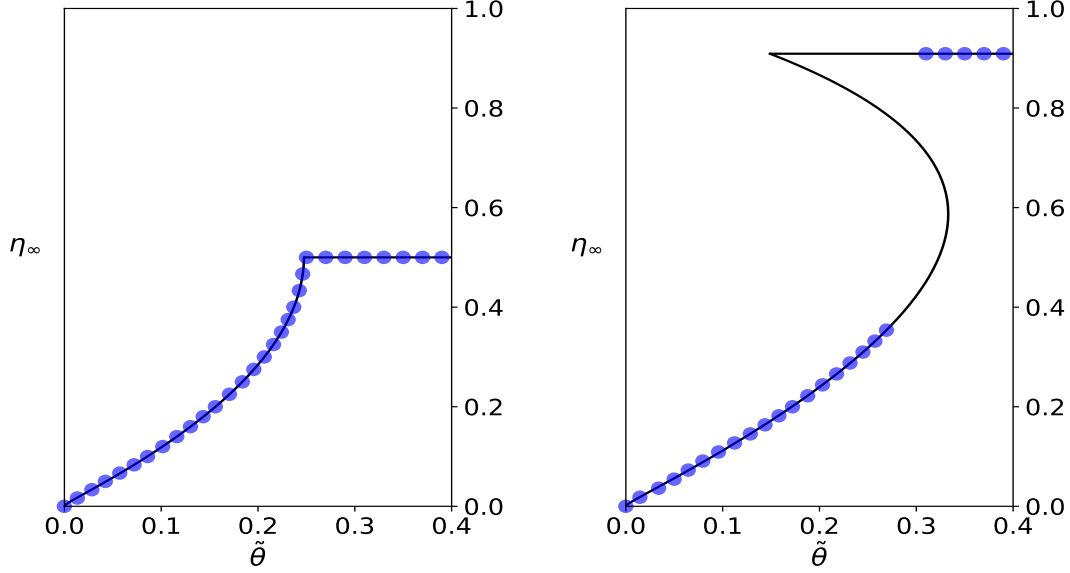


Figure 5: Long-run value of the scaled uncertainty measure η_∞ (blue dots) as a function of the cost parameter $\tilde{\theta}$, in the case of a linear memory cost function. Left panel: $K = 1, \rho = 0$. Right panel: $K = 10, \rho = 0$.

increased. A graph of the relationship between η_∞ and the value of $\tilde{\theta}$ is shown in Figure 5, for the case $\rho = 0$, and two different possible values of K : $K = 1$ and $K = 10$. For each value of $\tilde{\theta}$, the value of η_∞ associated with the optimal memory structure is shown by a large blue dot.

In each panel of this figure, the continuous black curve is the correspondence consisting of all points $(\tilde{\theta}, \eta_\infty)$ such that η_∞ is a stationary solution of the Euler equation associated with the optimization problem on the right-hand side of (3.3).³³ The Euler equation represents a first-order condition for the optimal choice of the degree of precision of memory; satisfaction of this condition is necessary but not sufficient for memory precision leading to $\eta_{t+1} = \eta$ to be optimal starting from a situation in which $\eta_t = \eta$. Because the objective function on the right-hand side of (3.3) is not a convex function, it can have multiple local minima (as well as a local maximum located between two local minima). Which of the local minima represents the global minimum (and hence the optimal memory structure) can jump abruptly as a result of a small change in parameters;³⁴ this is what happens when the value of η_∞ changes abruptly in the right panel of Figure 5, for a value of $\tilde{\theta}$ slightly above 0.28.

In the $K = 10$ case, we see that there need not be a unique value of η_∞ for a given value of $\tilde{\theta}$ that represents a stationary solution to the Euler equation. For any value of $\tilde{\theta}$ greater than a critical value around 0.15, if one starts from $\eta_t = \eta_0$ (a completely uninformative memory), the choice of $\eta_{t+1} = \eta_0$ again represents a local minimum of the objective; hence $\eta = \eta_0$ is a stationary solution of the Euler equation for all of these values of $\tilde{\theta}$, as shown in the figure. However, for values of $\tilde{\theta}$ only moderately larger than the critical value (such as $\tilde{\theta} = 0.20$), this is not the only local minimum, and the global minimum is instead at an

³³See Appendix F.4 for derivation of this equation.

³⁴See Appendix F.3 for a numerical example.

interior choice for λ_t ; this value results in a path $\{\eta_t\}$ that converges to a different stationary value for η_∞ , on the lower branch of the correspondence (as shown for example by the blue dot for $\tilde{\theta} = 0.20$). Yet for values of $\tilde{\theta}$ that exceed a second critical value just above 0.28, the global minimum shifts from the interior minimum to the local minimum at $\eta_{t+1} = \eta_0$. For all values beyond this point, the optimal memory structure involves $\lambda_t = 0$ for all t , so that $\eta_\infty = \eta_0$ (as shown by the blue dots on the upper branch of the correspondence).

Thus while the locus of fixed points $\eta_\infty(\lambda)$ is the same in Figures 1 and 3, all points on this locus represent possible long-run memory structures (attainable through an appropriate choice of $\bar{\lambda}$) in the case of a fixed upper bound on mutual information, but not all of them are always attainable in the case of a linear memory cost function. In the case $K = 1$, the two sets of long-run solutions are identical; but in the case $K = 10$, there is a range of values for η_∞ that are associated with particular (relatively low) values of $\bar{\lambda}$ but do not correspond to any possible value of $\tilde{\theta}$.³⁵

3.3 Stationary fluctuations in the long run

Because our model implies that a DM does not learn the true value of μ with certainty even in the long run, despite an arbitrarily long sequence of observations of the external state, over which time the coefficients of the data-generating process (1.1) are assumed not to change, it follows that the DM’s forecasts can be quite different from rational-expectations forecasts — that is, the forecasts of an ideal statistician who knows the true coefficient values. From the standpoint of an observer who is able to determine the true process, the forecasts of the DM with limited memory will appear to be systematically biased. The biases in the DM’s forecasts will furthermore fluctuate over time, in response both to variations in the external state (to which the DM reacts differently than someone with rational expectations would) and to noise in the evolution of the memory state.

We obtain a particularly simple characterization of the systematic pattern of forecast biases if we consider the long run — the predictions of the equations in the previous two sections in the case of very large values of t , so that η_t has converged to the constant value η_∞ , λ_t has converged to λ_∞ , and so on. In this case, our model, like the model of “natural expectations” of Fuster *et al.* (2010, 2011), predicts a stationary pattern of forecast biases that do not reflect incomplete adjustment to a new environment.

In the long run, equations (1.1), (2.1), and (2.7) become a system of linear equations with constant coefficients and Gaussian innovation terms, describing the evolution of the DM’s cognitive state. This system of equations can be reduced to a VAR(1) system

$$\tilde{s}_{t+1} = f\mu + F\tilde{s}_t + u_{t+1}, \quad u_{t+1} \sim N(0, \Sigma_u) \quad (3.5)$$

where

$$\tilde{s}_t \equiv \begin{bmatrix} \tilde{m}_t \\ y_t \end{bmatrix}, \quad u_{t+1} \equiv \begin{bmatrix} \tilde{\omega}_{t+1} \\ \epsilon_{y,t+1} \end{bmatrix},$$

³⁵We can show analytically that the continuous relationship shown in the left panel of Figure 5 occurs for all $K \leq 1$ when $\rho = 0$, while the backward-bending correspondence and consequent discontinuous relationship between $\tilde{\theta}$ and η_∞ occurs for all $K > 1$. See Appendix F.4 for further explanation.

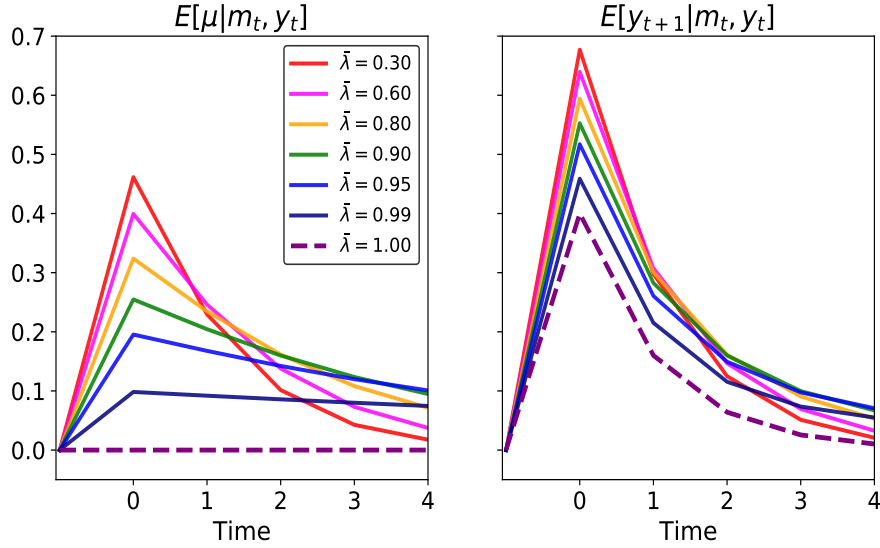


Figure 6: Impulse responses of the DM’s estimate of μ (left panel) and one-period-ahead forecast of the state (right panel) to a unit positive innovation in the observed value of y_t at the time marked as “time = 0” on the horizontal axis. Responses are plotted for alternative values of the information bound $\bar{\lambda}$, in the case that $K = 1, \rho = 0.4$.

and f, F and Σ_u are a 2-vector and two 2×2 matrices of constant coefficients respectively. In this vector system, the first equation is obtained by substituting (2.1) into (2.19), while the second equation is given by (1.1).

The matrix F furthermore has an upper-triangular form, while Σ_u is diagonal. We show in the appendix that the eigenvalues of the matrix F are ρ and ρ_m .³⁶ We further show that $0 < \rho_m < 1$, so that both y_t and $\tilde{m}_t$ exhibit stationary fluctuations around well-defined long-run average values which depend linearly on μ . The two independent exogenous sources of variation in this system are the innovations $\epsilon_{y,t+1}$ in the external state and the memory noise innovations $\tilde{\omega}_{t+1}$.

The DM’s optimal estimate of μ at each point in time, $\hat{\mu}_t$, as well as her optimal forecast of the external state at any horizon $\tau > t$,

$$\hat{y}_{\tau|t} = E[y_{\tau} | \tilde{m}_t, y_t] = (1 - \rho^{\tau-t})\hat{\mu}_t + \rho^{\tau-t}y_t, \quad (3.6)$$

will then be linear functions of the elements of $\tilde{s}_t$, with coefficients that are also time-invariant. We thus obtain a stationary multivariate Gaussian distribution for any number of leads and lags of the external state, the DM’s memory state, and the DM’s estimates and forecasts. This allows us to analyze not only the extent to which the DM’s forecasts should differ from rational-expectations forecasts, but the correlation that one should observe between the bias in the DM’s forecasts and other observable variables.

In particular, the biases in the DM’s forecasts will be correlated with the evolution of the external state. An unexpectedly high observed value for y_t will be interpreted (because of the DM’s uncertainty about μ) as implying a higher optimal estimate of μ , and this increase

³⁶See Appendix G.1 for the derivation.

in the DM's estimate of μ will furthermore persist, decaying only gradually in subsequent periods. This is illustrated in the left panel of Figure 6, which shows the impulse response function for $\hat{\mu}_\tau$ to a unit positive innovation in the value of y_t . The response is plotted for a variety of alternative values for the information bound $\bar{\lambda}$, in the case that $K = 1$ and $\rho = 0.4$

³⁷

In the case that $\bar{\lambda} = 1$ (perfect memory), the value of μ is learned with perfect precision, and as a consequence there is no effect (in the long run, depicted here) of fluctuations in y_t on the DM's estimate of μ . (The Kalman gain γ_1 has a long-run value of zero in this case.) Instead, for values of $\bar{\lambda} < 1$, a higher observed value of y_t leads the DM to increase her estimate $\hat{\mu}_t$ (the Kalman gain is positive). The estimate $\hat{\mu}_\tau$ remains higher (on average) in subsequent periods as well. The memory state $\tilde{m}_{t+1}$ carried into the period following the innovation is a noisy record of $\hat{\mu}_t$, and hence is higher because of the increase in y_t ; this increases the average value of the estimate $\hat{\mu}_{t+1}$, which increases the average value of the memory state $\tilde{m}_{t+2}$, and so on. The tighter the memory constraint (the lower the value of $\bar{\lambda}$), the greater the effect of the innovation in y_t on $\hat{\mu}_t$, because the DM is more uncertain about the value of μ before observing y_t ; however, the effect on the DM's estimate of μ is also more transient the lower the value of $\bar{\lambda}$, because less information is retained from one period to the next about past cognitive states.

These effects on the DM's optimal estimate of μ then feed into her optimal forecast of the external state at any future horizon τ , because of (3.6). As an illustration, the right panel of Figure 6 shows the impulse response of the one-quarter-ahead forecast $\hat{y}_{\tau+1|t}$ to a unit positive innovation in y_t , using the same conventions as in the left panel.³⁸ When $\rho > 0$, the rational-expectations forecast (corresponding to $\bar{\lambda} = 1$ in the figure) is itself increased by a positive innovation in y_t (by an amount equal to fraction ρ of the innovation), and the increase in the forecast is furthermore persistent (decaying back to its original level at a rate proportional to $\rho^{\tau-t}$). But when $\bar{\lambda} < 1$, the forecast is increased by even more, owing to the fact that the higher observation of y_t increases the DM's estimate of μ as well. This additional effect on the forecast is initially larger the smaller is $\bar{\lambda}$; but a smaller $\bar{\lambda}$ (tighter memory constraint) also causes the additional effect to die out more rapidly, since its propagation can only be through the DM's memory of her previous judgment about the value of μ .

Thus our model predicts that forecasts of the future value of a variable will over-react to news about the current value of that variable (assuming, as is often the case with economic time series, that the variable in question exhibits positive serial correlation). Positive serial correlation means that a higher current observation should increase somewhat one's forecast of the variable's future value, even under rational expectations; but imperfect memory results in a larger increase in the forecast than is consistent with rational expectations. The model also predicts that biases of this kind will persist for some time. Once a situation occurs that leads the DM to over-estimate the future level of some time series, the DM will as a consequence continue (on average) to over-estimate the future level of that variable for several more quarters.

³⁷See Appendix G.1 for illustration of how this figure would change under alternative assumptions about the degree of persistence of the fluctuations in the external state.

³⁸The corresponding impulse responses for alternative values of ρ are again shown in Appendix G.1.

3.4 “Recency bias” in expectation formation

One type of systematic difference between observed expectations and those of a perfect Bayesian decision maker that has often been reported is “recency bias” (e.g., Malmendier *et al.*, 2017) — a tendency for expectations to be influenced more by more recent observations, even when in principle, observations of a given time series at earlier dates should be equally relevant as a basis for inference. Our model predicts that such a bias should exist, as a consequence of optimal adaptation to limited memory precision (or to the cost of maintaining a more precise memory). Observations of the external state farther in the past are recalled with more noise, and as a consequence are given less weight in estimating parameters of the data generating process than would be optimal in the case of a perfect memory of past data.

The system (3.5) implies that, in the case that data have been generated in accordance with this law of motion for a sufficiently long time, we can express the value of the memory state $\tilde{m}_{t+1}$ as a function of the sequence of external states $\{y_\tau\}$ for $\tau \leq t$ and the sequence of memory noise realizations $\{\tilde{\omega}_{\tau+1}\}$ for $\tau \leq t$:

$$\tilde{m}_{t+1} = F_{12} \cdot \sum_{j=0}^{\infty} (\rho_m)^j y_{t-j} + \tilde{\omega}_{t+1}^{sum}, \quad (3.7)$$

where F_{12} is the $(1, 2)$ element of the matrix F in (3.5) and

$$\tilde{\omega}_{t+1}^{sum} \equiv \sum_{j=0}^{\infty} (\rho_m)^j \tilde{\omega}_{t+1-j} \quad (3.8)$$

is a serially correlated Gaussian noise term.³⁹

Equation (2.1) implies that a DM’s estimate of the unknown mean μ of the external state is given by a linear relation of the form

$$\hat{\mu}_t = \xi \tilde{m}_t + \gamma_1 y_t, \quad (3.9)$$

where the coefficient $\xi > 0$ is defined in the appendix. Using (3.7) to substitute for the memory state in this expression, we see that we can write the estimate in the form

$$\hat{\mu}_t = \sum_{j=0}^{\infty} \alpha_j y_{t-j} + \xi \tilde{\omega}_t^{sum}, \quad (3.10)$$

where the weights $\{\alpha_j\}$ are all positive, and the weights for $j \geq 1$ decrease exponentially: $\alpha_j = \alpha_1 (\rho_m)^{j-1}$.

The forecasts specified by (3.6) using this value for $\hat{\mu}_t$ are similar to those implied by a model of least-squares learning (Evans and Honkapohja, 2001) in which the DM is assumed to know that the variable’s law of motion is of the form (1.1); the value of the coefficient ρ is assumed to be known while μ must be estimated; and the unknown coefficient is estimated using a “constant-gain” estimator.⁴⁰ The biases in forecasts predicted by our model will

³⁹This is a stationary random process with a finite unconditional variance, since $0 < \rho_m < 1$ as shown in Appendix G.1.

⁴⁰The differences between (3.10) and a standard constant-gain estimate of the mean of a series are the fact that the coefficient α_0 is differently specified, and the presence of the Gaussian error term.

therefore have important similarities to those of a model of constant-gain learning, of the kind included in estimated macroeconomic models by authors such as Milani (2007, 2014) and Slobodyan and Wouters (2012).

We provide, however, a justification for the declining weight on observations farther in the past, as a consequence of optimal forecasting based on an imperfect memory, and furthermore endogenize the nature of that memory. The fact that our model predicts decreasing weights on observations made farther in the past is a notable difference between our model and the one proposed by Afrouzi *et al.* (2020).

4 Predictable Forecast Errors

An important consequence of optimal Bayesian inference with perfect memory (as assumed under the hypothesis of rational expectations) is that the error in a forecast should not itself be forecastable on the basis of any information available to the forecaster, at or before the time of the forecast in question. Thus if we let $\hat{y}_{t+h|t}$ denote a DM’s forecast at time t of the value of the external state at time $t+h$, the forecast error⁴¹ $FE_t \equiv y_{t+h} - \hat{y}_{t+h|t}$ should be uncorrelated with any variable z_t the value of which is known to the DM at time t (or earlier), either because it has been publicly observable or because it is part of the DM’s own cognitive state. Many econometric investigations of the consistency of observed forecasts with the hypothesis of rational expectations have accordingly been based on regressions of FE_t on other variables that ought to be known to the forecaster, testing the null hypothesis that all such regression coefficients should equal zero. Here we discuss the extent to which our model can account for some widely discussed examples of evidence against this null hypothesis; we particularly discuss evidence indicating over-reaction of subjective expectations to news about the series that is to be forecasted.

4.1 Evidence of over-reaction: the response of forecasts to fluctuations in the state

As noted in the introduction, Afrouzi *et al.* (2020) conduct a laboratory experiment in which forecasts of a stationary AR(1) process are elicited from subjects. They find that subjects’ expectations over-react to innovations in this process, as predicted by our model (as well as the related model of noisy memory that they discuss). They give particular emphasis to a measure of over-reaction in which a subject’s forecast $\hat{y}_{t+h|t}$ (where h is the number of realizations in advance for which the forecast is solicited in trial t) is regressed on the realization of the variable just before the forecast is solicited:

$$\hat{y}_{t+h|t} = \alpha_h^{subj} + \rho_h^{subj} y_t + v_t. \quad (4.1)$$

A separate regression (with coefficients α_h, ρ_h^{subj}) can be estimated for each of several horizons h . Afrouzi *et al.* are interested in the difference between the “subjective degree of

⁴¹Note that we give this variable a time subscript t to indicate that it is the error in the forecast made at time t ; the value of the random variable FE_t is not revealed however until date $t+h$.

persistence” measured by the estimated coefficient ρ_h^{subj} and the corresponding coefficient ρ_h in a regression using actual outcomes:

$$y_{t+h} = \alpha_h + \rho_h y_t + u_{t+h}. \quad (4.2)$$

The authors measure the degree of over-reaction of expectations to news by the extent to which ρ_h^{subj} is larger than ρ^h . Note that this is an example of a test of the predictability of forecast errors, since the coefficient of a regression of FE_t on y_t will equal $\rho^h - \rho_h^{subj}$.

We can investigate what our model of expectation formation on the basis of an imperfect memory implies about the relationship between ρ_h^{subj} and ρ_h in the case of a stationary AR(1) process. Here we consider the predicted values of the regression coefficients in the long run, as the length of the time series used to estimate them goes to infinity. The law of motion (1.1) implies that for any horizon $h \geq 1$, the joint distribution of y_t and y_{t+h} (conditional on the value of μ) will be bivariate Gaussian, with

$$E[y_{t+h} | \mu, y_t] = (1 - \rho^h)\mu + \rho^h y_t.$$

Hence with a sufficiently long series of observations, the coefficients in a regression of the form (4.2) should approach the asymptotic values

$$\alpha_h = (1 - \rho^h)\mu, \quad \rho_h = \rho^h.$$

(Here we assume that the regression uses an arbitrarily long sequence of realizations of a process for which there is a single, unchanging value of μ .)

Equation (3.6) implies that subjective forecasts should be given by

$$\hat{y}_{t+h|t} = (1 - \rho^h)\hat{\mu}_t + \rho^h y_t,$$

so that the predicted coefficient ρ_h^{subj} in regression (4.1) will equal

$$\rho_h^{subj} = (1 - \rho^h)\beta_{\hat{\mu}|y} + \rho^h = (1 - \rho_h)\beta_{\hat{\mu}|y} + \rho_h, \quad (4.3)$$

where $\beta_{\hat{\mu}|y}$ is the coefficient in a regression of $\hat{\mu}_t$ on y_t ,

$$\beta_{\hat{\mu}|y} = \frac{\text{cov}[\hat{\mu}_t, y_t | \mu]}{\text{var}[y_t | \mu]} = \frac{\text{cov}[\hat{\mu}_t, y_t | \mu]}{\sigma_y^2}.$$

We show in the appendix how to calculate this coefficient as a function of the model parameters.⁴²

Importantly, our numerical solutions indicate that $\hat{\mu}_t$ and y_t are always positively correlated (conditional on μ). This is because a positive innovation in the external state y_t raises (or at least never lowers) the expected value of y_τ for all $\tau \geq t$, and at the same time also raises the expected value of $\hat{\mu}_\tau$ for all $\tau \geq t$ (as illustrated in Figure 6 and similar figures in the appendix). Since the memory noise has no effect on the evolution of the external state, there are no shocks that move $\hat{\mu}_t$ and y_t in opposite directions, while some (at least the innovation ϵ_{yt}) move both of them in the same direction. But given that $\beta_{\hat{\mu}|y} > 0$, equation

⁴²See Appendix G.3 for details.

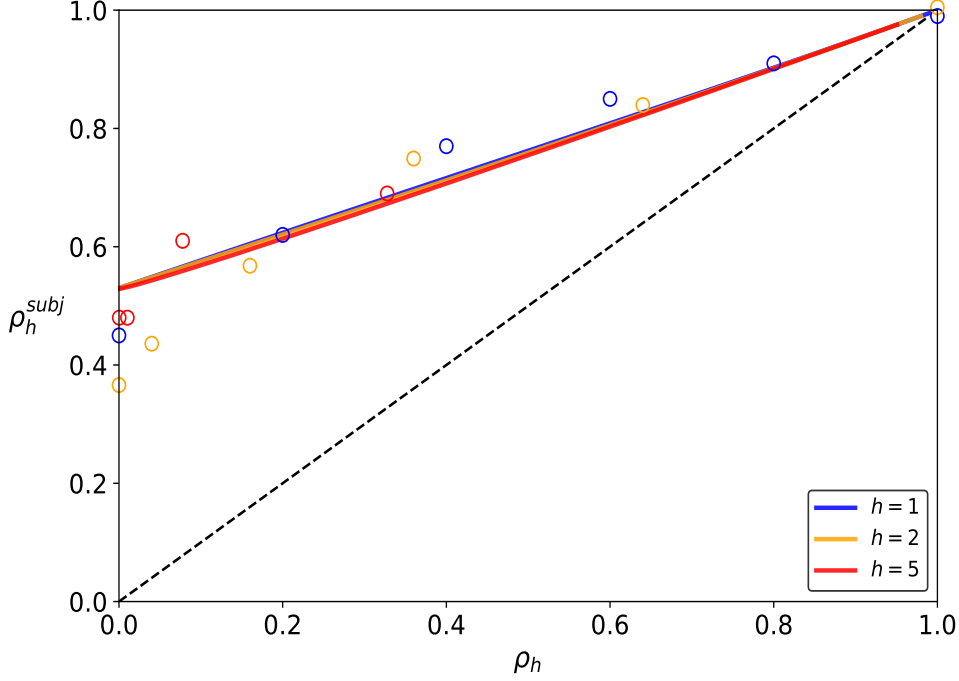


Figure 7: Comparison of the values for the regression coefficients ρ_h and ρ_h^{subj} for different values of ρ and h . (The figure is shown for the case $K = 1, \bar{\lambda} = 0.3$.) The diagonal line indicates the prediction of the rational-expectations hypothesis.

(4.3) implies that $\rho_h^{subj} > \rho_h$; that is, our model implies over-reaction of the kind exhibited by the forecasts of the subjects of Afrouzi *et al.*

Equation (4.3) also implies that for fixed values of the model parameters other than ρ , the over-reaction measure $\rho_h^{subj} - \rho_h$ converges to zero as $\rho \rightarrow 1$, for any forecast horizon h .⁴³ This is also approximately true of the regression coefficients reported by Afrouzi *et al.* (see their Figures 2B, 5A, and 5B). Indeed, these authors stress the finding that in their data, the discrepancy $\rho_h^{subj} - \rho_h$ is much larger when ρ_h is relatively small (either because ρ is small, or because ρ is well below one and h is large). This is also true in numerical solutions of our model as indicated in Figure 7.

One of the more striking features of the regressions reported by Afrouzi *et al.* is that ρ_h^{subj} is well approximated by an increasing function of ρ_h , with approximately the same functional relationship regardless of whether the variation in ρ_h occurs as a result of variation in ρ or variation in h .⁴⁴ The relationship $\rho^{subj}(\rho)$ is furthermore an upward-sloping one, with a slope much less than one, starting well above the diagonal for low values of ρ and approaching the

⁴³This prediction depends on $\beta_{\hat{\mu}|y}$ remaining bounded as ρ approaches 1. This is the case in our numerical solutions, both when $\bar{\lambda}$ is held constant as ρ is varied (as in Figure 2) and when $\bar{\theta}$ is held constant as ρ is varied (as in Figure 4).

⁴⁴This was shown in an earlier version of the paper now circulated as Afrouzi *et al.* (2020), though this figure is omitted from their most recent draft.

diagonal as $\rho \rightarrow 1$. (See the plot of their regression coefficients in Figure 7.⁴⁵) While our model does not imply that a functional relationship of that kind should hold precisely, it is worth noting that to the extent that the value of $\beta_{\hat{\mu}|y}$ remains approximately the same as one varies ρ , (4.3) implies that the value of ρ_h^{subj} should be nearly the same for all pairs (ρ, h) that imply the same value of ρ_h . Perhaps more to the point, our model can be parameterized so that it simultaneously fits the experimental evidence for each of the three different horizons for which forecasts are solicited in the experiment of Afrouzi *et al.*

Figure 7 plots the predicted value of ρ_h^{subj} against the value of ρ_h , for each of several different horizons h , each represented by a distinct curve; the curves are shown for the case in which $K = 1$ and $\bar{\lambda} = 0.3$. Along each curve, the variation in ρ_h is due purely to variation in ρ . (The fact that $\bar{\lambda}$ is fixed despite variation in ρ means that we assume a fixed upper bound on the mutual information, as in section 3.1, rather than a convex cost function.) The horizons used are $h = 1, 2$ and 5 , as these are the horizons for which Afrouzi *et al.* elicit forecasts from their subjects; the regression coefficients that they estimate for various combinations of ρ and h are indicated by the circles in the figure (with colors indicating the horizon h).

The three curves are not exactly the same, since in our model $\beta_{\hat{\mu}|y}$ is a function of ρ (but the same for all values of h), rather than being a function only of ρ_h . Nonetheless, for the parameterization chosen here, $\beta_{\hat{\mu}|y}$ is nearly constant as ρ is varied; as a consequence, the relationship between ρ_h and ρ_h^{subj} predicted by (4.3) is close to a linear one, and is nearly the same for all values of h . Our model therefore provides quite a good account of the effects of variation in either ρ or h on the value of ρ_h^{subj} , as indicated by the fact that none of the circles in Figure 7 are far from the corresponding curve.

4.2 Evidence of over-reaction: forecast revisions and forecast error

A comparison of the coefficient ρ_h^{subj} with ρ_h provides a fairly straightforward test of over-reaction of forecasts to news about the variable being forecasted; but the null hypothesis that ρ_h^{subj} should equal ρ_h only follows from rational expectations in the case that one is sure that forecasters at time t have observed the external state y_t , and calculation of the predicted value ρ_h requires knowledge of the true data-generating process. Both of these assumptions make sense in the case of the laboratory experiment of Afrouzi *et al.* (2020); but they are more debatable in the case of forecasts of economic time series outside laboratory settings.

Bordalo *et al.* (2020a) use a different approach to document systematic departures from rational expectations in surveys of professional forecasters. Following a proposal by Coibion and Gorodnichenko (2015), they regress the error (that eventually becomes known) in a given forecaster's forecast of a future data release on the revision that the forecast represents, relative to the same forecaster's forecast of the same future variable at an earlier time. That is, they test whether the coefficient b is different from zero in a regression specification of the form

$$FE_t = a + b \cdot FR_t + u_t, \quad (4.4)$$

⁴⁵The data plotted here are based on Figures 2B, 5A, and 5B of Afrouzi *et al.* (2020).

where FE_t is the error (as defined above) in the forecast at time t for some horizon $h > 0$, and

$$FR_t \equiv \hat{y}_{t+h|t} - \hat{y}_{t+h|t-1}$$

is the revision of this forecast between time $t - 1$ and time t .

If the forecast $\hat{y}_{t+h|t}$ represents the correctly calculated expectation of y_{t+h} conditional on the forecaster's information set at time t , then the forecast error FE_t should not be forecastable on the basis of any information available to the forecaster at time t , as discussed above. Bordalo *et al.* point out that even if one is agnostic about which external developments are observed by forecasters (or how accurately they observe them), as long as one supposes that the forecaster's *own past cognitive states* are known with complete precision, then the size of the forecast revision FR_t should be part of the information set; hence the coefficient b in (4.4) should equal zero under a hypothesis of correct Bayesian inference from the forecaster's information set. A coefficient $b \neq 0$ allows one to reject not just the full-information rational-expectations hypothesis, but also models that assume that people's choices are optimal Bayesian responses conditional upon a noisy cognitive state (but with perfect memory), such as the models of Sims (2003) or Woodford (2003).

Bordalo *et al.* (2020a) find instead that b is significantly negative in the case of many macroeconomic and financial time series. (Afrouzi *et al.*, 2020, find that the same is true of the forecasts elicited in their experiment.⁴⁶) Bordalo *et al.* interpret this negative sign as evidence of over-reaction of forecasts to economic news arriving between the dates of the two successive forecasts: news that implies that one's previous forecast was too low results in an upward revision that is too large, so that the occurrence of an upward revision is correlated with the second forecast turning out to be too high.

As discussed above, our model implies that there will be over-reaction to new realizations of the external state, and indeed our model predicts that one can easily have a negative coefficient b in a regression of the form (4.4). The theoretically predicted asymptotic value for the coefficient b in the case of a long enough series of observations from an environment with unchanging statistics (including a fixed value of μ) is given by

$$b = \frac{\text{cov}[FE_t, FR_t | \mu]}{\text{var}[FR_t | \mu]},$$

where the forecast error and forecast revision variables are defined above. Since the denominator is necessarily positive (in any case in which the forecast is not constant at all times), the coefficient b should be negative if and only the covariance between FE_t and FR_t is negative.

That this can easily be the case can be illustrated by considering the simple case of an i.i.d. process ($\rho = 0$) for the external state. In this case, (3.6) implies that $\hat{y}_{t+h|t} = \hat{\mu}_t$, for any forecast horizon $h \geq 1$. In this case we have

$$\text{var}[FR_t | \mu] = \text{var}[\hat{\mu}_t - \hat{\mu}_{t-1} | \mu] = 2(1 - \rho_{\hat{\mu}})\text{var}[\hat{\mu}_t | \mu],$$

⁴⁶See Figure 2A in their paper. Indeed, the evidence for $b < 0$ as a general regularity is even stronger in the laboratory data of Afrouzi *et al.* than in the field data of Bordalo *et al.*

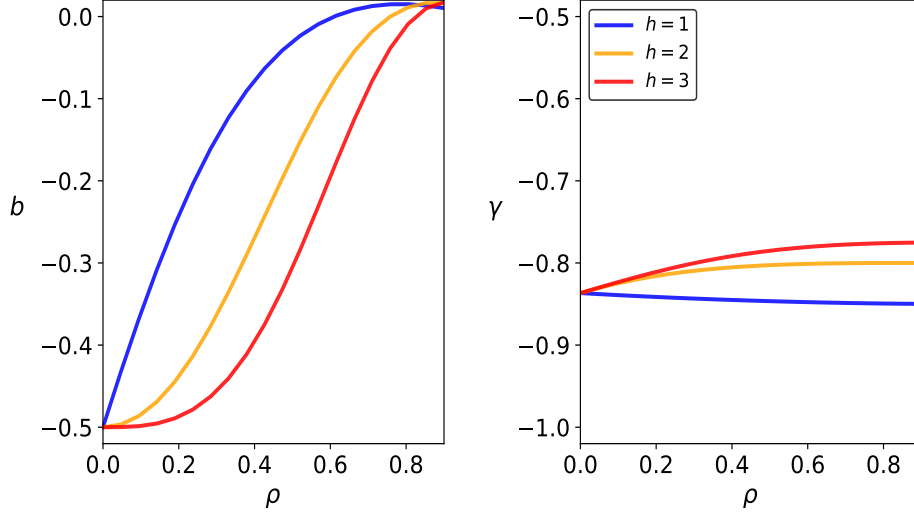


Figure 8: Predicted value of the coefficient b from a regression of forecast errors on the size of the revision of the forecast (left panel), and the coefficient γ from a regression of forecast revisions on the deviation of an individual's forecast from the consensus forecast (right panel). In each case, results are shown for external state processes of different degrees of serial correlation ρ . As in Figure 7, we assume that $K = 1$, $\bar{\lambda} = 0.3$.

where $\rho_{\hat{\mu}}$ is the coefficient of serial correlation of the stationary fluctuations in $\hat{\mu}_t$, and

$$\begin{aligned} \text{cov}[FE_t, FR_t | \mu] &= \text{cov}[y_{t+h} - \hat{\mu}_t, \hat{\mu}_t - \hat{\mu}_{t-1} | \mu] \\ &= \text{cov}[-\hat{\mu}_t, \hat{\mu}_t - \hat{\mu}_{t-1} | \mu] = -(1 - \rho_{\hat{\mu}})\text{var}[\hat{\mu}_t | \mu]. \end{aligned}$$

(Here the first expression on the second line follows from the fact that y_{t+h} is completely uncorrelated with any variables observable at time t or earlier, conditional on the value of μ , when $\rho = 0$.) Hence in this case we obtain the prediction $b = -1/2$, simply as a consequence of the fact that our model implies stationary fluctuations in $\hat{\mu}_t$ around some long-run average estimate, for any parameter values with $\lambda_{\infty} < 1$.

Our numerical solutions indicate that b is often negative in the case of positive serial correlation in the external state as well, as illustrated in the left panel of Figure 8.⁴⁷ The figure shows the predicted value of the coefficient b as the coefficient of serial correlation ρ varies between 0 and 1. The figure is computed under the assumption that $K = 1$ and $\bar{\lambda} = 0.3$, as in Figure 7. We see that the predicted degree of over-reaction (as measured by the degree to which $b < 0$) is greatest when $\rho = 0$; the coefficient equals -0.5 when $\rho = 0$ (as explained in the previous paragraph), but is less negative when $\rho > 0$, and near zero for large values of ρ . This is also what Afrouzi *et al.* find to be true of the forecasts elicited in their laboratory experiment.⁴⁸

A similar regularity is observed in the case of professional forecasts of the economic time series considered by Bordalo *et al.* (2020a). The estimates that they obtain for b mainly

⁴⁷Formulas that can be used to calculate b are given in Appendix G.3.

⁴⁸Again, see their Figure 2A.

fall in or near the interval $[-0.5, 0]$. Moreover, in the case of the highly persistent series that they consider, the value of b is around zero on average (sometimes slightly negative, but sometimes slightly positive); in the case of the series that they consider with a coefficient of serial correlation less than 0.1, the value of b is nearly as negative as -0.5; and for the series that they consider with intermediate degrees of serial correlation, b is negative but much less negative than -0.5.⁴⁹ Our model is not only able to explain why negative values of b are often obtained, but also why these are almost always between small positive values and -0.5, and why the coefficient is more negative for the least persistent time series.

4.3 Idiosyncratic noise in individual forecasts

Another kind of evidence of systematic bias in the forecasts announced by individual forecasters is provided by Fuhrer (2018). Fuhrer shows that subsequent revisions of the forecasts of an individual forecaster are partially forecastable on the basis of information available (at least to the community of forecasters in general) at the time of the original forecast, a result inconsistent with the hypothesis of full-information rational expectations (under which not only should all forecasters be ideal Bayesian statisticians, but all should share a common information set).

Suppose now that we let $\hat{y}_{t+h|t}^i$ be the forecast of y_{t+h} at time t by forecaster i , while $\hat{y}_{t+h|t}^{cons}$ is the “consensus forecast,” the median of the forecasts at time t made by the different forecasters in a given survey. Fuhrer reports regressions of the form

$$\hat{y}_{t+h|t}^i - \hat{y}_{t+h|t-1}^i = a + \gamma \cdot (\hat{y}_{t+h|t-1}^i - \hat{y}_{t+h|t-1}^{cons}) + e'Z_t^i + u_t, \quad (4.5)$$

where Z_t^i is a vector of other forecaster-specific control variables (that vary across specifications). His main finding, obtained for forecasts of several different aggregate variables, and robust both to different choices for the horizon h and the control variables included, is that the coefficient γ is found to be significantly negative (for example, between -0.5 and -0.6 in the case of revisions of inflation forecasts by members of the Survey of Professional Forecasters, where the forecasts are collected at a quarterly frequency). This indicates a tendency of forecasters to subsequently revise their forecasts so as to partially eliminate the previous gap between their forecast and the consensus forecast.

Many models of biased expectation formation proposed in the previous literature predict systematic departures from rational expectations, and hence allow forecast revisions to be predictable by information that was publicly available at the earlier date; but they nonetheless do not predict that the variable considered by Fuhrer should predict subsequent forecast revisions, insofar as they do not explain why the forecasts of different forecasters should respond in different ways to the same publicly available information. Our model instead requires that there should be idiosyncratic noise in individual forecasts; for it is only because of the noise term in the law of motion for the memory state (2.19) that the DM is unable to learn the value of μ with perfect precision, and there is no reason for the noise term ω_{t+1}^i to be correlated across decision makers.

It is clear that not only the forecasts of different households, but even the forecasts of professional forecasters exhibit substantial dispersion at a given point in time. A defender

⁴⁹See Figure 1 of Afrouzi *et al.*, who stress this feature of the results of Bordalo *et al.*

of the hypothesis of Bayesian rationality might argue that this simply reflects the fact that different forecasters have access to different private sources of information. Yet experiments in which forecasts are elicited from different subjects who are shown identical sequences of observations indicate that dispersion of forecasts exists even the experimenter can be certain that each subject was exposed to precisely the same information; see in particular Khaw *et al.* (2017) and Afrouzi *et al.* (2020). This indicates noisy cognitive processing of the information presented to the subjects; our model provides at least one possible example of the nature of such idiosyncratic noise in the process by which individuals' expectations are formed.

Our model not only explains why there should exist non-trivial dispersion in the discrepancy between individual forecasts and the consensus forecast, but why this discrepancy should predict subsequent forecast revisions in the way that it does. In our model, all forecasters observe the same external states, but their memories of past states are subject to idiosyncratic noise. In the case of a large enough sample of forecasters, the mean realizations of the memory noise term $\tilde{\omega}_{t+1}^i$ across forecasters i should be close to zero each period, so that (3.7) implies that the mean memory state should be essentially a deterministic function of the past external states,

$$\tilde{m}_{t+1}^{avg} = F_{12} \cdot \sum_{j=0}^{\infty} (\rho_m)^j y_{t-j}.$$

If we let $\tilde{m}_{t+1}^{diff,i} \equiv \tilde{m}_{t+1}^i - \tilde{m}_{t+1}^{avg}$ be the difference between the memory state of DM i and the average memory state, (3.7) implies that

$$\tilde{m}_{t+1}^{diff,i} = \tilde{\omega}_{t+1}^{sum,i}.$$

If we similarly let $\hat{\mu}_t^{diff,i}$ be the difference between DM i 's estimate of μ and the average estimate, it follows from (3.9) that this difference must entirely be due to the difference between i 's memory state and the average memory state, so that we can write

$$\hat{\mu}_t^{diff,i} = \xi \tilde{m}_t^{diff,i}.$$

Finally, it follows from (3.6) that for any forecast horizon h , the difference between i 's forecast and the average forecast will be due entirely to the difference in i 's estimate of μ , so that

$$\hat{y}_{t+h|t}^i - \hat{y}_{t+h|t}^{avg} = (1 - \rho^h) \hat{\mu}_t^{diff,i} = (1 - \rho^h) \xi \tilde{\omega}_t^{sum,i}.$$

In the large-sample limit, the population distribution of realizations of the variable $\tilde{\omega}_{t+1}^{sum,i}$ across forecasters i should be essentially be identical to the distribution of the random variable $\tilde{\omega}_{t+1}^{sum}$. It follows that the distribution of individual forecasts $\{\hat{y}_{t+h|t}^i\}$ should be approximately a Gaussian distribution, with a median very close to its mean, $\hat{y}_{t+h|t}^{avg}$. Thus the consensus forecast should be essentially the same as $\hat{y}_{t+h|t}^{avg}$, and we obtain the prediction

$$\Delta_t^i \equiv \hat{y}_{t+h|t}^i - \hat{y}_{t+h|t}^{cons} = (1 - \rho^h) \xi \tilde{\omega}_t^{sum,i}. \quad (4.6)$$

In a regression of the form (4.5), but where, for purposes of our theoretical derivation, we assume there are no control variables Z_t^i included, the asymptotic value of the coefficient

γ (with a long enough sample) should equal

$$\gamma = \frac{\text{cov}[FR_t^i, \Delta_{t-1}^i]}{\text{var}[\Delta_{t-1}^i]}.$$

The precise predicted value depends on both the degree of persistence of the variable being forecasted and the length of the forecast horizon h , but we can show in general that our model predicts that $-1 < \gamma < 0$.⁵⁰ If we use parameter values $K = 1$, $\bar{\lambda} = 0.3$, as in Figure 7, the model predicts a regression coefficient on the order of -0.8, as shown in the right panel of Figure 8. And indeed, Fuhrer (2018) finds that the predictable forecast revision over the next quarter should be more than half of the discrepancy between i 's forecast and the consensus forecast (though less than a complete correction of the discrepancy), for the most of the variables and alternative regression specifications that he considers.

5 Conclusion

We have shown that it is possible to characterize the optimal structure of memory, for a class of linear-quadratic-Gaussian forecasting problems, when the cost of a more precise memory is proportional to Shannon's mutual information, and when we assume that the joint distribution of past cognitive states and the memory state is of a multivariate Gaussian form, but with no *a priori* restriction on the dimension of the memory state or the dimensions of past experience that may be more or less precisely recalled. Strikingly, we find that for the class of problems that we consider, the optimal memory structure is necessarily at most one-dimensional. This means that what can be recalled at any time about past observations is simply a noisy recollection of a single summary statistic for past experience. We show how the model parameters determine the law of motion for that summary statistic, and hence what single dimension of past experience will be (imprecisely) available as an input to the DM's forecasts.

Among the implications of our model, two seem of particularly general interest. First, while our formalism allows for the possibility of an independent noisy record of each past observation (as assumed for example in the model of Neligh, 2019), this is not optimal; instead, the optimal memory structure is one in which only a particular weighted average of past observations can be recalled with noise. And second, this weighted average places much larger weights on recent observations than on ones at earlier dates, even though observations at all dates are equally relevant to inference about the value of the parameter μ , which matters for the DM's decisions. Thus our model provides an explanation for "recency bias" in the influence of past observations on current decisions, unlike the model of endogenous memory precision proposed by Afrouzi *et al.* (2020).

We have shown that our model predicts "over-reaction" of forecasts of an autoregressive process to current realizations of the process, of a kind similar to that observed in the forecasts of experimental subjects (Afrouzi *et al.*, 2020), and in survey forecasts of macroeconomic and financial time series (Bordalo *et al.*, 2020a). Moreover, it predicts that the degree of over-reaction should be greater in the case of less persistent time series. This prediction

⁵⁰See Appendix G.4 for an explicit solution for this coefficient.

is confirmed both by laboratory experiments and survey forecasts, as discussed by Afrouzi *et al.* (2020), but cannot be explained by competing explanations for over-reaction, such as the theory of “diagnostic expectations” proposed by Bordalo *et al.* (2018).⁵¹ Our model also predicts that over-reaction should be observed even in the case of time series with a very simple autocorrelation structure — even a white noise time series — unlike the model of “natural expectations” proposed by Fuster *et al.* (2010, 2011), and again in conformity with experimental evidence. And unlike either diagnostic or natural expectations (at least in their most basic forms), our model provides an explanation for heterogeneity in forecasts on the part of forecasters with access to the same information; indeed, the noise in people’s memories is at the heart of our explanation for forecast biases.

In the applications sketched above, we have focused on biases that have been observed in people’s stated expectations. But we suspect that the expectational biases implied by our model can help to explain puzzling aspects of market outcomes as well. For example, Bordalo *et al.* (2020b) argue that a number of well-known puzzles about the behavior of the aggregate stock market are in fact all consistent with a simple dividend discount model of stock prices, under the hypothesis that market expectations regarding firms’ future earnings differ systematically from rational expectations in a particular way, that is furthermore consistent with the biases observed in survey expectations of earnings. They further show that a particular sort of bias in market expectations is needed in order to explain both the biases in survey expectations and the asset pricing anomalies, one very much like the kind of forecast bias predicted by our model.

Briefly, Bordalo *et al.* propose a model in which asset prices at time t are based on market expectations of dividend growth g_{t+h} at various future horizons h . Dividend growth is assumed to be a stationary autoregressive process; market expectations of g_{t+h} differ from rational expectations by an expectational error term $\epsilon_{h,t}$. For any horizon h , $\epsilon_{h,t}$ is assumed to be a stationary, mean-zero autoregressive process, with a substantial degree of persistence; and the innovations in $\epsilon_{h,t}$ are positively correlated with the innovations in g_t , though fluctuations in $\epsilon_{h,t}$ also occur that are uncorrelated with fundamentals. Finally, the fluctuations in $\epsilon_{h,t}$ for different horizons h are perfectly correlated, and $\epsilon_{h,t}$ remains different from zero as $h \rightarrow \infty$, so that innovations in the error process bias expectations about dividend growth in the far future and not only in the near term.

These assumptions are all features of subjective forecasts of the future evolution of the state y_{t+h} in our model (if we identify our y_t with dividend growth). We have shown (in the right panel of Figure 6) that in our model, innovations in y_t cause subjective expectations of the future state to rise more than the RE forecast would, and the effect persists for several periods, though the bias caused by the innovation in any single period t eventually converges to zero. For each horizon h , (3.6) implies that the bias term is equal to $(1 - \rho^h)\hat{\mu}_t$; thus the biases for different forecast horizons are all perfectly correlated. Moreover, as the horizon is increased, the bias term becomes simply $\hat{\mu}_t$ for all large enough h ; thus the forecast errors predicted by the model are above all errors in long-term forecasts.

Our model also implies that there will be random fluctuations in forecast bias that are uncorrelated with any underlying fundamentals; these innovations are indicated by the $\tilde{\omega}_{t+1}$

⁵¹See Afrouzi *et al.* (2020) for a thorough demonstration of the inability of a variety of familiar models to explain this pattern.

shock in (3.5). The most important difference with the reduced-form specification of expectational bias proposed by Bordalo *et al.* is that in their model, there are arbitrary random variations in the “market expectations” that determine the value of the stock market; our model instead implies the existence of idiosyncratic random variation in the beliefs of an individual forecaster, but one might expect that these idiosyncratic variations should cancel out in their effects on the market price. It is possible that a satisfactory model of asset pricing will require us to suppose that some individual traders are large enough for their idiosyncratic beliefs to have a non-negligible effect on aggregate outcomes, as in the model of Gabaix *et al.* (2006). We leave the development of a complete model of asset prices for future work. But it seems likely that imperfect memory of the kind modeled here will be a necessary element in such a model.

References

- [1] Afrouzi, Hassan, Spencer Youngwook Kwon, Augustin Landier, Yueran Ma, and David Thesmar, “Overreaction and Working Memory,” NBER Working Paper no. 27947, October 2020.
- [2] Bordalo, Pedro, Nicola Gennaioli, and Andrei Shleifer, “Diagnostic Expectations and Credit Cycles,” *Journal of Finance* 73: 199-227 (2018).
- [3] Bordalo, Pedro, Nicola Gennaioli, Yueran Ma, and Andrei Shleifer, “Over-Reaction in Macroeconomic Expectations,” *American Economic Review* 110: 2748-2782 (2020a).
- [4] Bordalo, Pedro, Nicola Gennaioli, Rafael LaPorta, and Andrei Shleifer, “Expectations of Fundamentals and Stock Market Puzzles,” NBER Working Paper no. 27283, May 2020b.
- [5] Caplin, Andrew, Mark Dean, and John Leahy, “Rationally Inattentive Behavior: Characterizing and Generalizing Shannon Entropy,” working paper, New York University, February 2019.
- [6] Coibion, Olivier, and Yuriy Gorodnichenko, “Information Rigidity and the Expectations Formation Process: A Simple Framework and New Facts,” *American Economic Review* 105: 2644-78 (2015).
- [7] Cover, Thomas M., and Joy A. Thomas, *Elements of Information Theory*, New York: Wiley, 2d ed., 2006.
- [8] Evans, George W., and Seppo Honkapohja, *Learning and Expectations in Macroeconomics*, Princeton: Princeton University Press, 2001.
- [9] Fuhrer, Jeffrey, “Expectations as a Source of Macroeconomic Persistence: Evidence from Survey Expectations in Dynamic Macro Models,” *Journal of Monetary Economics* 86: 22-35 (2017).
- [10] Fuhrer, Jeffrey, “Intrinsic Expectations Persistence: Evidence from Professional and Household Survey Expectations,” Federal Reserve Bank of Boston Working Paper no. 18-9, May 2018.
- [11] Fuster, Andreas, Ben Hébert, and David Laibson, “Natural Expectations, Macroeconomic Dynamics, and Asset Pricing,” *NBER Macroeconomics Annual* 26: 1-48 (2011).
- [12] Fuster, Andreas, David I. Laibson, and Brock Mendel, “Natural Expectations and Macroeconomic Fluctuations,” *Journal of Economic Perspectives* 24(4): 67-84 (2010).
- [13] Gabaix, Xavier, Parameswaran Gopikrishnan, Vasiliki Plerou, and H. Eugene Stanley, “Institutional Investors and Stock-Market Volatility,” *Quarterly Journal of Economics* 121: 461-504 (2006).

- [14] Khaw, Mel Win, Luminita L. Stevens, and Michael Woodford, “Discrete Adjustment to a Changing Environment: Experimental Evidence,” *Journal of Monetary Economics* 91: 88-103 (2017).
- [15] Malmendier, Ulrike, and Stefan Nagel, “Learning from Inflation Experiences,” *Quarterly Journal of Economics* 131: 53–87 (2016).
- [16] Malmendier, Ulrike, Demian Pouzo, and Victoria Vanasco, “A Theory of Experience Effects,” working paper, UC Berkeley, January 2017.
- [17] Metzler, Lloyd A., “The Nature and Stability of Inventory Cycles,” *Review of Economics and Statistics* 23: 113-129 (1941).
- [18] Milani, Fabio, “Expectations, Learning and Macroeconomic Persistence,” *Journal of Monetary Economics* 54: 2065–2082 (2007).
- [19] Milani, Fabio, “Learning and Time-varying Macroeconomic Volatility,” *Journal of Economic Dynamics and Control* 47(C): 94–11 (2014).
- [20] Muth, John F., “Rational Expectations and the Theory of Price Movements,” *Econometrica* 29: 315-335 (1961).
- [21] Neligh, Nathaniel, “Rational Memory with Decay,” working paper, Chapman University, November 2019.
- [22] Sims, Christopher A., “Implications of Rational Inattention,” *Journal of Monetary Economics* 50: 665-690 (2003).
- [23] Slobodyan, Sergey, and Raf Wouters, “Learning in an Estimated Medium-scale DSGE Model,” *Journal of Economic Dynamics and Control* 36: 26–46 (2012).

For Online Publication

APPENDIX

Azeredo da Silveira, Sung, and Woodford,
“Optimally Imprecise Memory and Biased Forecasts”

A Reduction of the General Forecasting Problem to Estimation of μ

Consider the problem of choosing the vector of forecasts z_t each period so as to minimize (1.2). The elements of z_t must be chosen as a function of the DM’s cognitive state at time t (after observing the external state y_t). As explained in the text, the DM’s cognitive state at time t is assumed to consist of the value of the current external state y_t (observed with perfect precision), along with whatever additional information is reflected in the DM’s period t memory state m_t . (In this section, it is not yet necessary to specify the nature of the vector m_t .)

If we use the notation $E_t[\cdot]$ for the expectation of a random variable conditional on a complete description of the state at date t (including knowledge of the true value of μ), then

$$E[(z_t - E_t \tilde{z}_t)' W (\tilde{z}_t - E_t \tilde{z}_t)] = 0,$$

since $\tilde{z}_t - E_t \tilde{z}_t$ is a function of innovations in the external state subsequent to date t , that must be distributed independently of all of the determinants of both z_t and $E_t \tilde{z}_t$. It follows that the term in (1.2) involving z_t can be equivalently expressed as⁵²

$$\begin{aligned} E[(z_t - \tilde{z}_t)' W (z_t - \tilde{z}_t)] &= E[(z_t - E_t \tilde{z}_t)' W (z_t - E_t \tilde{z}_t)] \\ &\quad + E[(\tilde{z}_t - E_t \tilde{z}_t)' W (\tilde{z}_t - E_t \tilde{z}_t)] \\ &\equiv L_{1t} + L_{2t}. \end{aligned}$$

Moreover, L_{2t} is independent of the decisions of the DM, and thus irrelevant to a determination of the optimal decision rule. The loss function (1.2) can thus equivalently be written as the discounted sum of the L_{1t} terms, which involve squared differences between z_t and $E_t \tilde{z}_t$.

It further follows from the law of motion (1.1) that

$$E_t \tilde{z}_t = \sum_{j=0}^{\infty} A_j [\mu + \rho^j (y_t - \mu)].$$

Since the precise value of y_t is presumed to be part of the cognitive state on the basis of which z_t can be chosen, one can write any decision rule in the form

$$z_t = \hat{z}_t + \left(\sum_{j=0}^{\infty} \rho^j A_j \right) \cdot y_t,$$

⁵²Here we omit the factor β^t that multiplies this term in (1.2).

where $\hat{z}_t$ must be some function of the cognitive state at date t . In terms of this notation, the relevant part of the loss function (1.2) can then be written as

$$L_{1t} = E[(\hat{z}_t - \mu a)'W(\hat{z}_t - \mu a)],$$

where we define $a \equiv \sum_{j=0}^{\infty} (1 - \rho^j) A_j$ and make use of the fact that $E_t[\mu] = \mu$.

The term L_{1t} that we wish to minimize can further be expressed as the expected value (integrating over all possible realizations of the cognitive state s_t in period t) of the quantity

$$\begin{aligned} \tilde{L}_1(s_t) &\equiv E[(\hat{z}_t - \mu a)'W(\hat{z}_t - \mu a) | s_t] \\ &= E[\hat{z}_t | s_t]'W E[\hat{z}_t | s_t] + E[\check{z}_t' W \check{z}_t | s_t] \\ &\quad - 2a'W E[\hat{z}_t | s_t] \cdot E[\mu | s_t] + a'W a \cdot E[\mu^2 | s_t], \end{aligned}$$

where we define $\check{z}_t \equiv \hat{z}_t - E[\hat{z}_t | s_t]$. (In expanding the right-hand side in this way, we use the fact that $E[\check{z}_t | s_t] = 0$, and that $\check{z}_t$ must be independent of the deviation of μ from $E[\mu | s_t]$, since the DM has no way to condition her action on μ except through the information about μ revealed by the cognitive state.) The expression $\tilde{L}_1(s_t)$ can then be separately minimized for each possible cognitive state s_t , by choosing a distribution for $\hat{z}_t$ conditional on that state. We further note that the random component $\check{z}_t$ of the action affects only the second term on the right-hand side, and so should be chosen to minimize that term; since W is positive definite, this is achieved by setting $\check{z}_t = 0$ with certainty, so that $\hat{z}_t$ must be a deterministic function of s_t .

We can then simply write $E[\hat{z}_t | s_t]$ as $\hat{z}_t$, and observe that

$$\tilde{L}_1(s_t) = (\hat{z}_t - aE[\mu | s_t])'W(\hat{z}_t - aE[\mu | s_t]) + a'W a \cdot \text{var}[\mu | s_t], \quad (\text{A.1})$$

where the final term on the right-hand side is independent of the choice of $\hat{z}_t$. Thus in each cognitive state s_t , $\hat{z}_t$ must be chosen to minimize the first term on the right-hand side; since W is positive definite, this is achieved by setting $\hat{z}_t = a \cdot \hat{\mu}_t$, where $\hat{\mu}_t = E[\mu | s_t]$.

Thus there is no loss of generality in restricting the DM to response rules of the form $\hat{z}_t = a \cdot \hat{\mu}_t$, where $\hat{\mu}_t$ is a scalar choice that depends on the cognitive state in period t , and that can be interpreted as the DM's estimate of μ given the cognitive state. Substituting this expression for $\hat{z}_t$ into (A.1), we have

$$\begin{aligned} \tilde{L}_1(s_t) &= a'W a \cdot \{(\hat{\mu}_t - E[\mu | s_t])^2 + \text{var}[\mu | s_t]\} \\ &= a'W a \cdot E[(\hat{\mu}_t - \mu)^2 | s_t]. \end{aligned}$$

Then taking the unconditional expectation of this expression, we obtain

$$L_{1t} = \alpha \cdot MSE_t,$$

where $\alpha \equiv a'W a > 0$ and MSE_t is defined as in the text.

Under any forecasting rule of the kind assumed here, then, the value of the loss function (1.2) will equal (1.4), plus an additional term

$$\sum_{t=0}^{\infty} \beta^t L_{2t}$$

that is independent of the DM's forecasting rule. Hence within this class of forecasting rules, the rule that minimizes (1.2) must be the one that minimizes (1.4); and since any other kind of forecasting rule can only lead to a higher value of (1.2), we can replace the problem of choosing a rule for determining z_t that minimizes (1.2) by the problem of choosing a rule for determining $\hat{\mu}_t$ that minimizes (1.4).

B Bayesian Updating After the External State is Observed: A Kalman Filter

Let the elements of the memory state be partitioned as

$$m_t = \begin{bmatrix} \underline{m}_t \\ \bar{m}_t \end{bmatrix}, \quad (\text{B.1})$$

where the lower block consists of the elements of the reduced memory state

$$\bar{m}_t \equiv E[x_t | m_t], \quad \text{where } x_t \equiv \begin{bmatrix} \mu \\ y_{t-1} \end{bmatrix},$$

while the upper block consists of the conditional expectations $E[y_{t-j} | m_t]$ for $2 \leq j \leq t$. (This simply requires an appropriate ordering of the elements of m_t , using the notation for this vector introduced in the main text.)

We assume a posterior distribution of the form

$$x_t | m_t \sim N(\bar{m}_t, \Sigma_t)$$

conditional on the memory state m_t , where $\bar{m}_t$ is a 2-vector and Σ_t is a 2×2 symmetric, p.s.d. matrix. Under our assumption of linear-Gaussian dynamics for the memory state, the vector $\bar{m}_t$ will also be drawn from a multivariate Gaussian distribution. Since the prior for the hidden state vector is specified to be

$$x_t \sim N(0, \Sigma_0), \quad \Sigma_0 \equiv \begin{bmatrix} \Omega & \Omega \\ \Omega & \Omega + \sigma_y^2 \end{bmatrix}, \quad (\text{B.2})$$

it follows that the unconditional distribution for the reduced memory state $\bar{m}_t$ must be of the form

$$\bar{m}_t \sim N(0, \Sigma_0 - \Sigma_t).$$

The complete set of variables (x_t, m_t) also have a multivariate Gaussian distribution. Moreover, since (by assumption) the expectation of x_t conditional on the realization of m_t depends only on the elements of $\bar{m}_t$, it follows that the entire distribution of x_t conditional on m_t depends only on $\bar{m}_t$, so that

$$x_t | m_t = x_t | \bar{m}_t.$$

Hence the joint distribution of the variables (x_t, m_t) can be factored as

$$p(x_t, \underline{m}_t, \bar{m}_t) = p(x_t, \bar{m}_t) \cdot p(\underline{m}_t | \bar{m}_t).$$

The DM then observes the external state y_t , which is assumed to depend on the hidden state vector x_t through an “observation equation” of the form

$$y_t = c'x_t + \epsilon_{yt}, \quad \epsilon_{yt} \sim N(0, \sigma_\epsilon^2)$$

as a consequence of (1.1), where we further assume that ϵ_{yt} is distributed independently of both m_t and x_t . It follows that the variables (x_t, m_t, y_t) will have a joint distribution that is multivariate Gaussian; and that this distribution can be factored as

$$\begin{aligned} p(x_t, m_t, y_t) &= p(x_t, m_t) \cdot p(y_t | x_t) \\ &= p(\underline{m}_t | \bar{m}_t) \cdot p(x_t, \bar{m}_t) \cdot p(y_t | x_t) \\ &= p(\underline{m}_t | \bar{m}_t) \cdot p(x_t, \bar{m}_t, y_t). \end{aligned}$$

From this it follows that

$$x_t | m_t, y_t = x_t | \bar{m}_t, y_t.$$

Thus both the expectation of x_t conditional on the cognitive state $s_t \equiv (m_t, y_t)$, and the variance-covariance matrix of the errors in the estimation of x_t based on the cognitive state, will depend only on the joint distribution of the variables $(x_t, \bar{m}_t, y_t)$. Moreover, the distribution for x_t conditional on the realizations of the elements of the cognitive state will be multivariate Gaussian,

$$x_t | \bar{m}_t, y_t \sim N(\bar{\mu}_t, \bar{\Sigma}_t), \quad (\text{B.3})$$

where $\bar{\mu}_t$ is a linear function of $\bar{m}_t$ and y_t , while $\bar{\Sigma}_t$ is independent of the realizations of either $\bar{m}_t$ or y_t .

We can further decompose the vector of means $\bar{\mu}_t$ as

$$\begin{aligned} \bar{\mu}_t &= E[x_t | \bar{m}_t, y_t] \\ &= E[x_t | \bar{m}_t] + \{E[x_t | \bar{m}_t, y_t] - E[x_t | \bar{m}_t]\} \\ &= \bar{m}_t + \gamma_t \cdot (y_t - E[y_t | \bar{m}_t]) \\ &= \bar{m}_t + \gamma_t \cdot (y_t - c'E[x_t | \bar{m}_t]) \\ &= \bar{m}_t + \gamma_t \cdot (y_t - c'\bar{m}_t), \end{aligned}$$

where γ_t is the vector of *Kalman gains*. (The first element of this vector equation is then just equation (2.1) in the main text.)

The vector of Kalman gains must be chosen so that the estimation errors $x_t - \bar{\mu}_t$ are orthogonal to the surprise in the observation of the external state, $y_t - c'\bar{m}_t$. This requires that

$$\begin{aligned} 0 &= \text{cov}(x_t - \bar{\mu}_t, y_t - c'\bar{m}_t) \\ &= \text{cov}((x_t - \bar{m}_t) - \gamma_t(y_t - c'\bar{m}_t), y_t - c'\bar{m}_t) \\ &= \text{var}[x_t - \bar{m}_t]c - \text{var}[c'(x_t - \bar{m}_t) + \epsilon_{yt}] \cdot \gamma_t \\ &= \Sigma_t c - [c'\Sigma_t c + \sigma_\epsilon^2] \cdot \gamma_t. \end{aligned}$$

Hence

$$\gamma_t = \frac{\Sigma_t c}{c'\Sigma_t c + \sigma_\epsilon^2}. \quad (\text{B.4})$$

The gain coefficient γ_{1t} in equation (2.1) is just the first element of this vector, $\gamma_{1t} \equiv e'_1 \gamma_t$. This together with (B.4) yields the formula (2.3) given in the main text.

The variance-covariance matrix in the conditional distribution (B.3) will be given by

$$\begin{aligned}
\bar{\Sigma}_t &= \text{var}[x_t - \bar{\mu}_t] = \text{var}[(x_t - \bar{m}_t) - \gamma_t(y_t - c' \bar{m}_t)] \\
&= \text{var}[(I - \gamma_t c')(x_t - \bar{m}_t) - \gamma_t \epsilon_{yt}] \\
&= (I - \gamma_t c') \Sigma_t (I - \gamma_t c')' + \sigma_\epsilon^2 \gamma_t \gamma_t' \\
&= \Sigma_t - 2[c' \Sigma_t c + \sigma_\epsilon^2] \gamma_t \gamma_t' + [c' \Sigma_t c] \gamma_t \gamma_t' + \sigma_\epsilon^2 \gamma_t \gamma_t' \\
&= \Sigma_t - [c' \Sigma_t c + \sigma_\epsilon^2] \gamma_t \gamma_t'.
\end{aligned}$$

The remaining uncertainty about the value of μ given the cognitive state, $\hat{\sigma}_t^2$, is then equal to $\bar{\Sigma}_{11,t}$, so that

$$\hat{\sigma}_t^2 = e'_1 \bar{\Sigma}_t e_1 = e'_1 \Sigma_t e_1 - (c' \Sigma_t c + \sigma_\epsilon^2) (\gamma_{1t})^2,$$

which is just expression (2.2) in the main text.

Substituting expression (B.2) for Σ_0 into this solution, we obtain

$$\begin{aligned}
\hat{\sigma}_0^2 &= \Omega - (\Omega + \sigma_y^2) \cdot \left[\frac{\Omega}{\Omega + \sigma_y^2} \right]^2 \\
&= \frac{\Omega \sigma_y^2}{\Omega + \sigma_y^2},
\end{aligned}$$

which is the formula given in (2.4). It remains to be shown that this is an upper bound for $\hat{\sigma}_t^2$. To show this, we observe that

$$\begin{aligned}
\hat{\sigma}_t^2 &= \min_{\beta, \gamma_1} \text{var}[\mu - \beta' \bar{m}_t - \gamma_1 y_t] \\
&\leq \min_{\gamma_1} \text{var}[\mu - \gamma_1 y_t] \\
&\leq \text{var}[\mu - (\Omega/(\Omega + \sigma_y^2)) \cdot y_t] \\
&= \text{var}[(\sigma_y^2/(\Omega + \sigma_y^2))\mu - (\Omega/(\Omega + \sigma_y^2))(y_t - \mu)] \\
&= \left(\frac{\sigma_y^2}{\Omega + \sigma_y^2} \right)^2 \text{var}[\mu] + \left(\frac{\Omega}{\Omega + \sigma_y^2} \right)^2 \text{var}[y_t | \mu] \\
&= \left(\frac{\sigma_y^2}{\Omega + \sigma_y^2} \right)^2 \Omega + \left(\frac{\Omega}{\Omega + \sigma_y^2} \right)^2 \sigma_y^2 \\
&= \frac{\Omega \sigma_y^2}{\Omega + \sigma_y^2} = \sigma_0^2.
\end{aligned}$$

This establishes the upper bound (2.4) stated in the main text.

C Demonstration that an Optimal Memory Structure Records Information Only about the Reduced Cognitive State

Let (1.5) be written in the partitioned form

$$\begin{bmatrix} \underline{m}_{t+1} \\ \bar{m}_{t+1} \end{bmatrix} = \begin{bmatrix} \Lambda_{a,t} & \Lambda_{b,t} \\ \Lambda_{c,t} & \Lambda_{d,t} \end{bmatrix} \begin{bmatrix} \underline{s}_t \\ \bar{s}_t \end{bmatrix} + \begin{bmatrix} \underline{\omega}_{t+1} \\ \bar{\omega}_{t+1} \end{bmatrix}. \quad (\text{C.1})$$

Here m_{t+1} is again partitioned as in (B.1). The lower block of s_t consists of the elements of the reduced cognitive state

$$\bar{s}_t \equiv \begin{bmatrix} \hat{\mu}_t \\ y_t \end{bmatrix},$$

both elements of which are linear functions of s_t , as a consequence of equation (2.1). We choose a representation for the vector s_t such that the lower block consists of the elements of $\bar{s}_t$, the elements of $\underline{s}_t$ are all uncorrelated with the elements of $\bar{s}_t$, and the elements of the vectors $\bar{s}_t$ and $\underline{s}_t$ together span the same linear space of random variables as the elements of s_t . (We can necessarily write any memory structure of the form (1.5) in this way; it amounts simply to a choice of the basis vectors in terms of which the vectors m_{t+1} and s_t are each decomposed.)

Let us suppose furthermore that a representation for m_{t+1} is chosen consistent with the normalization $E[\bar{s}_t | m_{t+1}] = \bar{m}_{t+1}$. This holds if and only if both elements of the vector $\bar{s}_t - \bar{m}_{t+1}$ are uncorrelated with each of the elements of m_{t+1} . These consistency conditions can be reduced to two requirements: (i) the requirement that

$$\text{var}[\Lambda_{c,t}\underline{s}_t + \bar{\omega}_{t+1}] = (I - \Lambda_{d,t})X_t\Lambda'_{d,t}, \quad (\text{C.2})$$

where the matrix $X_t \equiv \text{var}[\bar{s}_t]$ is independent of the memory structure chosen for period t ; and (ii) the requirement that $\bar{s}_t - \bar{m}_{t+1}$ be uncorrelated with all elements of $\underline{m}_{t+1}$. (Note that $\bar{s}_t - \bar{m}_{t+1}$ is uncorrelated with $\bar{m}_{t+1}$ if and only if (C.2) holds.)

C.1 Forecast accuracy depends only on the matrices $\{\Lambda_{d,t}\}$

Suppose that in any period t , we take the memory structure in periods $\tau < t$ as given. This means that the DM's uncertainty about x_t given the memory state m_t (specified by the posterior variance-covariance matrix Σ_t) will be given. (If $t = 0$, Σ_0 is simply given by the prior.) Hence the value of $\hat{\mu}_t$ as a function of $\bar{m}_t$ and y_t will be given, and consequently the value of MSE_t will be given, following the discussion in the main text (and the previous section of this appendix). The elements of the matrix X_t will similarly be given.

We next consider how $\Lambda_{d,t}$ must be chosen, in order for it to be possible to choose matrices $\Lambda_{c,t}$ and $\text{var}[\bar{\omega}_{t+1}]$ such that (C.2) is satisfied. Equation (C.2) requires that $(I - \Lambda_{d,t})X_t\Lambda'_{d,t}$ be a symmetric matrix; this will hold if and only if the simpler requirement is satisfied that $\Lambda_{d,t}X_t = X_t\Lambda'_{d,t}$ be a symmetric matrix. In addition, it is necessary that $(I - \Lambda_{d,t})X_t\Lambda'_{d,t}$ be a p.s.d. matrix. The set of matrices $\Lambda_{d,t}$ with these properties is a non-empty set ($\Lambda_{d,t} = 0$ is

a trivial example), and depends only on the matrix X_t . Let this set of matrices be denoted $\mathcal{L}(X_t)$.

Now let $\Lambda_{d,t}$ be any matrix that belongs to $\mathcal{L}(X_t)$. Then it is possible to choose the matrices $\Lambda_{c,t}$ and $\text{var}[\bar{\omega}_{t+1}]$ so that (C.2) is satisfied; and given any such choice of these two matrices, it is further possible to choose the specification of the equation for $\underline{m}_{t+1}$ so that all elements of $\underline{m}_{t+1}$ are uncorrelated with the elements of $\bar{s}_t - \bar{m}_{t+1}$. Given any such specifications, both conditions (i) and (ii) above will be satisfied. Thus the matrix $\Lambda_{d,t}$ is admissible as part of the specification of a memory structure; and any possible memory structure consistent with the matrix $\Lambda_{d,t}$ will be one of those with the properties just assumed.

Given a matrix $\Lambda_{d,t}$ of this sort, we next observe that the equations determining $\bar{m}_{t+1}$ can be written in the form

$$\bar{m}_{t+1} = \Lambda_{d,t} \bar{s}_t + \nu_{t+1},$$

where $\nu_{t+1} \sim N(0, \Lambda_{d,t} X_t)$ is distributed independently of $\bar{s}_t$. Thus the joint distribution of $(\bar{s}_t, \bar{m}_{t+1})$ will be a multivariate Gaussian distribution, the parameters of which are completely determined by X_t and $\Lambda_{d,t}$. It then follows that the conditional distribution $\bar{s}_t | \bar{m}_{t+1}$ will be a bivariate Gaussian distribution, with a mean $\bar{m}_{t+1}$ and a variance independent of the realization of $\bar{m}_{t+1}$, which also depends only on X_t and $\Lambda_{d,t}$. Moreover, since the elements of $\underline{m}_{t+1}$ are all Gaussian random variables distributed independently of $\bar{s}_t - \bar{m}_{t+1}$, knowledge of $\underline{m}_{t+1}$ cannot further improve one's estimate of $\bar{s}_t$, and so the conditional distribution $\bar{s}_t | m_{t+1} = \bar{s}_t | \bar{m}_{t+1}$. Finally, since we can write

$$x_{t+1} = \bar{s}_t + \begin{bmatrix} u_t \\ 0 \end{bmatrix},$$

where $u_t \sim N(0, \hat{\sigma}_t^2)$ must be uncorrelated with any of the elements of s_t (and hence uncorrelated with any of the elements of m_{t+1}), we must further have

$$x_{t+1} | m_{t+1} \sim N(\bar{m}_{t+1}, \Sigma_{t+1})$$

where

$$\Sigma_{t+1} = \text{var}[\bar{s}_t | \bar{m}_{t+1}] + \hat{\sigma}_t^2 e_1 e_1'.$$

Since $\hat{\sigma}_t^2$ also depends only on Σ_t (see equation (2.2)), it follows that the elements of Σ_{t+1} depend only on Σ_t and $\Lambda_{d,t}$.

This argument can then be used recursively (starting from period $t = 0$) to show that given the initial uncertainty matrix Σ_0 implied by the prior (B.2), we can completely determine the entire sequence of matrices $\{\Sigma_t\}$, given a sequence of matrices $\{\Lambda_{d,t}\}$ for all $t \geq 0$ with the property that for each t , $\Lambda_{d,t} \in \mathcal{L}(X_t)$, where X_t is the matrix implied by Σ_t . Moreover, given such a sequence of matrices $\{\Lambda_{d,t}\}$, the value of MSE_t for each period t will be uniquely determined as well. Hence the terms in the loss function (1.6) that depend on the accuracy of forecasts that are possible using a given memory structure will depend only on the sequence of matrices $\{\Lambda_{d,t}\}$. (These matrices must be chosen to satisfy a set of consistency conditions, stated above, but these conditions can also be expressed purely in terms of the sequence of matrices $\{\Lambda_{d,t}\}$.) Thus the other elements of the specification (C.1) of the memory structure matter only to the extent that they have consequences for the information cost terms in (1.6).

C.2 Mutual information: a useful lemma

Information costs in period t are assumed to be an increasing function of $I_t = I(M; S)$, the Shannon mutual information between random variables M (the realizations of which are denoted m_{t+1}) and S (the realizations of which are denoted s_t).⁵³ Each of the random vectors M and S can further be partitioned as $M = (\underline{M}, \bar{M})$, $S = (\underline{S}, \bar{S})$.

Now for any random variables $X_1, X_2, \dots$, let $H(X_1, X_2, \dots, X_k)$ be the entropy of the joint distribution for variables $(X_1, X_2, \dots, X_k)$, and $H(X_1, \dots, X_k | X_{k+1}, \dots, X_{k+m})$ be the entropy of the joint distribution of the variables $(X_1, \dots, X_k)$ conditional on the values of the variables $(X_{k+1}, \dots, X_{k+m})$. The chain rule for entropy implies that

$$H(X_1, X_2, \dots, X_k) = H(X_1) + H(X_2 | X_1) + \dots + H(X_k | X_1, \dots, X_{k-1}).$$

We can then define the mutual information between the variables $(X_1, \dots, X_k)$ and the variables $(X_{k+1}, \dots, X_{k+m})$ as

$$I(X_1, \dots, X_k; X_{k+1}, \dots, X_{k+m}) \equiv H(X_1, \dots, X_k) - H(X_1, \dots, X_k | X_{k+1}, \dots, X_{k+m}).$$

(The information about the first set of variables that is revealed by learning the values of the second set of variables is measured by the average amount by which the entropy of the conditional distribution is smaller than the entropy of the unconditional distribution of the first set of variables.) Similarly, we can define the mutual information between the first set of variables and the second set of variables, conditioning on the values of some third set of variables as

$$\begin{aligned} I(X_1, \dots, X_k; X_{k+1}, \dots, X_{k+m} | X_{k+m+1}, \dots, X_{k+m+n}) \\ \equiv H(X_1, X_2, \dots, X_k | X_{k+m+1}, \dots, X_{k+m+n}) - H(X_1, \dots, X_k | X_{k+1}, \dots, X_{k+m+n}). \end{aligned}$$

Thus for any set of four random variables $\underline{M}, \bar{M}, \underline{S}, \bar{S}$, we must have

$$\begin{aligned} I(\underline{S}, \bar{S}; \underline{M}, \bar{M}) &= H(\underline{S}, \bar{S}) - H(\underline{S}, \bar{S} | \underline{M}, \bar{M}) \\ &= [H(\bar{S}) + H(\underline{S} | \bar{S})] - [H(\bar{S} | \underline{M}, \bar{M}) + H(\underline{S} | \bar{S}, \underline{M}, \bar{M})] \\ &= [H(\bar{S}) + H(\underline{S} | \bar{S})] - [H(\bar{S}, \underline{M}, \bar{M}) - H(\underline{M} | \bar{M}) - H(\bar{M})] - H(\underline{S} | \bar{S}, \underline{M}, \bar{M}) \\ &= [H(\bar{S}) + H(\underline{S} | \bar{S})] - [(H(\bar{M}) + H(\bar{S} | \bar{M}) + H(\underline{M} | \bar{M}, \bar{S})) - H(\underline{M} | \bar{M}) - H(\bar{M})] \\ &\quad - H(\underline{S} | \bar{S}, \underline{M}, \bar{M}) \\ &= [H(\bar{S}) + H(\underline{S} | \bar{S})] - [H(\bar{S} | \bar{M}) + H(\underline{M} | \bar{M}, \bar{S}) - H(\underline{M} | \bar{M})] - H(\underline{S} | \bar{S}, \underline{M}, \bar{M}) \\ &= [H(\bar{S}) - H(\bar{S} | \bar{M})] + [H(\underline{S} | \bar{S}) - H(\underline{S} | \bar{S}, \underline{M}, \bar{M})] + [H(\underline{M} | \bar{M}) - H(\underline{M} | \bar{M}, \bar{S})] \\ &= I(\bar{S}; \bar{M}) + I(\underline{S}; \underline{M}, \bar{M} | \bar{S}) + I(\underline{M}; \bar{S} | \bar{M}). \end{aligned}$$

Then, since mutual information is necessarily non-negative, we can establish the lower bound

$$I_t = I(\underline{S}, \bar{S}; \underline{M}, \bar{M}) \geq I(\bar{S}; \bar{M}). \quad (\text{C.3})$$

⁵³Here we adopt the notation used in Cover (2006), with different symbols for the random variables M and S and their realizations. This is to make it clear that I_t is not a function of the values taken by m_{t+1} and s_t along a particular history, but instead a function of the complete joint distribution of the two random variables; I_t is itself not a random variable, but a single number for each date t .

Furthermore, this lower bound is achieved if and only if

$$I(\underline{S}; \underline{M}, \bar{M} | \bar{S}) = I(\underline{M}; \bar{S} | \bar{M}) = 0.$$

For any three random variables X, Y, Z , the conditional mutual information $I(X; Y | Z) = 0$ if and only if the variables X and Y are distributed independently one another, conditional on the value of Z . Hence the lower bound (C.3) is achieved if and only if (a) conditional on the value of $\bar{m}_{t+1}$, the variables $\bar{s}_t$ and $\underline{m}_{t+1}$ are independent of one another; and (b) conditional on the value of $\bar{s}_t$, the variables $\underline{s}_t$ and m_{t+1} are independent of one another.

C.3 Optimality of Setting $\Lambda_{a,t} = \Lambda_{b,t} = \Lambda_{c,t} = 0$

We return now to the consideration of possible memory structures. Let the sequence of matrices $\{\Lambda_{d,t}\}$ be chosen to satisfy the consistency conditions discussed above, and for a given such sequence, consider an optimal choice of the remaining elements of the specification (C.1), from among those specifications that are consistent with the sequence $\{\Lambda_{d,t}\}$ (that is, that will satisfy both conditions (i) and (ii) stated above).

We have shown above that the sequence of values $\{MSE_t\}$ is completely determined by the specification of $\{\Lambda_{d,t}\}$. Hence other aspects of the specification of the memory structure can matter only to the extent that they affect the sequence of values $\{I_t\}$. Moreover, we have shown that the joint distribution of $(\bar{s}_t, \bar{m}_{t+1})$ each period is completely determined by X_t and $\Lambda_{d,t}$, which means that the lower bound for I_t given in (C.3) is completely determined by the choice of $\{\Lambda_{d,\tau}\}$ for $\tau \leq t$. It thus remains only to consider whether this lower bound can be achieved, and under what conditions.

We first observe that the lower bound is achievable. For any sequence of matrices $\{\Lambda_{d,t}\}$ satisfying the specified conditions, a memory structure specification with $\Lambda_{a,t} = \Lambda_{b,t} = \Lambda_{c,t} = 0$, together with a stipulation that $\underline{\omega}_{t+1}$ be distributed independently of $\bar{\omega}_{t+1}$ and that $\text{var}[\bar{\omega}_{t+1}] = \Lambda_{d,t}X_t$, will satisfy both conditions (i) and (ii) stated in the introduction to this appendix, and thus this represents a feasible memory structure. One can also show that such a specification satisfies both of conditions (a) and (b) stated at the end of section C.2, so that the lower bound (C.3) is achieved in each period. Thus such a specification achieves the lowest possible value for the combined objective function (1.6), and will be optimal, given our choice of the sequence $\{\Lambda_{d,t}\}$.

Not only will this specification be sufficient for achieving the lowest possible value of (1.6), but it will be essentially necessary. We have shown above that achieving the lower bound for I_t in period t requires that conditional on the value of $\bar{s}_t$, the variables $\underline{s}_t$ and m_{t+1} are independent of one another. This means that the values of the variables in the vector $\underline{s}_t$ cannot help at all in predicting any elements of m_{t+1} , once one is already using the reduced cognitive state $\bar{s}_t$ to forecast the next period's memory state; thus one must be able to write law of motion (C.1) for the memory state with $\Lambda_{a,t} = \Lambda_{c,t} = 0$.⁵⁴ Thus it is necessarily the

⁵⁴It might be possible to satisfy the condition required for the lower bound with non-zero elements in one of these matrices; but this will occur only because of collinearity in the fluctuations in the elements of the vector $\underline{s}_t$, so that it is possible to have a law of motion in which $\underline{s}_t$ has no effect on m_{t+1} , despite non-zero matrices $\Lambda_{a,t}$ and $\Lambda_{c,t}$. In such a case, the representation of the cognitive state by the vector s_t would involve redundancy; and in any event, there would be no loss of generality in setting $\Lambda_{a,t} = \Lambda_{c,t} = 0$, since the implied fluctuations in the memory state would be the same.

case that the elements of m_{t+1} convey information only about the reduced cognitive state $\bar{s}_t$, and not about any other aspects of the cognitive state s_t .

In addition, we have shown above that achieving the lower bound for I_t in period t requires that conditional on the value of $\bar{m}_{t+1}$, the variables $\bar{s}_t$ and $\underline{m}_{t+1}$ are independent of one another. Thus all of the information about $\bar{s}_t$ that is contained in the memory state m_{t+1} is contained in the elements $\bar{m}_{t+1}$. This means either that $\Lambda_{b,t} = 0$ as well, or, to the extent that some element of $\underline{m}_{t+1}$ corresponds to a row of $\Lambda_{b,t}$ with non-zero elements, that element of $\underline{m}_{t+1}$ must be a linear combination of the elements of $\bar{m}_{t+1}$, so that conditioning upon its value conveys no new information about $\bar{s}_t$. Thus any specification of the memory structure in which $\Lambda_{b,t} \neq 0$ in any period represents a redundant representation of the contents of memory available in period $t + 1$; we can equivalently describe the contents of memory by eliminating all such rows from m_{t+1} .

Thus there is no loss of generality in assuming that the lower bound is achieved by specifying $\Lambda_{a,t} = \Lambda_{b,t} = \Lambda_{c,t} = 0$ in each period. Finally, satisfaction of consistency condition (ii) in this case requires that the elements of $\underline{\omega}_{t+1}$ be distributed independently of the elements of $\bar{\omega}_{t+1}$. We might still allow $\text{var}[\underline{\omega}_{t+1}]$ to be non-zero; this would mean that $\underline{m}_{t+1}$ contains elements that fluctuate randomly, but are completely uncorrelated with the previous period's cognitive state s_t . Such an information structure is equally optimal, in the sense that (1.6) is made no larger by the existence of such components of the memory state, given our assumption that only mutual information is costly. But the additional components $\underline{m}_{t+1}$ of the memory structure will have no consequences for cognitive processing, and our inclusion of them as part of the representation of the memory state violates our assumption in the text that we label memory states by their implied posteriors for the values of μ and the past realizations of the external state; using labels $(\underline{m}_{t+1}, \bar{m}_{t+1})$ in which $\underline{m}_{t+1}$ is non-null will mean having separate labels for memory states that imply the same posterior (since the value of $\underline{m}_{t+1}$ would be completely uninformative about either μ or any past external states).

Hence in the case of any optimal memory structure, the memory state can be described more compactly by identifying it with the reduced memory state $\bar{m}_{t+1}$, which evolves according to

$$\bar{m}_{t+1} = \bar{\Lambda}_t \bar{s}_t + \bar{\omega}_{t+1}, \quad (\text{C.4})$$

where $\bar{\Lambda}_t$ is the matrix called $\Lambda_{d,t}$ in (C.1). (This corresponds to equation (2.7) in the main text.) We need only consider (at most) a two-dimensional memory state, and the optimal memory state conveys information only about the reduced cognitive state $\bar{s}_t$, not about any other aspects of the cognitive state s_t .

C.4 An alternative representation for the reduced cognitive state

We have shown in the main text (equation (2.10)) that the variance matrix of the reduced cognitive state $\bar{s}_t$ can be written as a function of the single parameter $\hat{\sigma}_t^2$:

$$X_t = X(\hat{\sigma}_t^2) \equiv \begin{bmatrix} \Omega - \hat{\sigma}_t^2 & \Omega \\ \Omega & \Omega + \sigma_y^2 \end{bmatrix}.$$

There is another way of writing this function that will be useful below.

We can orthogonalize the reduced cognitive state using the transformation $\bar{s}_t = \Gamma \check{s}_t$, where

$$\Gamma \equiv \begin{bmatrix} 1 & \frac{\Omega}{\Omega + \sigma_y^2} \\ 0 & 1 \end{bmatrix}. \quad (\text{C.5})$$

The elements of the orthogonalized cognitive state have the interpretation

$$\check{s}_t \equiv \begin{bmatrix} \hat{\mu}_t - \text{E}[\mu|y_t] \\ y_t \end{bmatrix},$$

from which it is obvious that the first element must be uncorrelated with the second.

The variance matrix of $\check{s}_t$ is therefore diagonal:

$$\text{var}[\check{s}_t] = \check{X}(\hat{\sigma}_t^2) \equiv \begin{bmatrix} \hat{\sigma}_0^2 - \hat{\sigma}_t^2 & 0 \\ 0 & \Omega + \sigma_y^2 \end{bmatrix}. \quad (\text{C.6})$$

We can then alternatively write

$$X(\hat{\sigma}_t^2) = \Gamma \check{X}(\hat{\sigma}_t^2) \Gamma'. \quad (\text{C.7})$$

D The Law of Motion for the Memory State and the Information Content of Memory

We now consider how the parameterization of the law of motion (C.4) for the memory state determines the degree of uncertainty about the external state vector that will exist when beliefs are conditioned on the memory state, and how the same parameters determine the mutual information between the memory state and the prior cognitive state, and hence the size of the information cost term $c(I_t)$.

We begin by recapitulating the conditions that the sequence of matrices $\{\bar{\Lambda}_t\}$ and $\{\Sigma_{\bar{\omega},t+1}\}$ must satisfy, in order for (C.4) to represent a memory structure consistent with the normalization according to which $\text{E}[x_{t+1} | \bar{m}_{t+1}] = \bar{m}_{t+1}$. Condition (C.2) will be satisfied if and only if

$$\Sigma_{\bar{\omega},t+1} = (I - \bar{\Lambda}_t) X_t \bar{\Lambda}_t'. \quad (\text{D.1})$$

In order for there to be a symmetric, p.s.d. matrix $\Sigma_{\bar{\omega},t+1}$ that satisfies (D.1), it must be the case that $\bar{\Lambda}_t \in \mathcal{L}(X_t)$. As explained above, this means that $\bar{\Lambda}_t X_t = X_t \bar{\Lambda}_t'$ must be a symmetric matrix, and in addition that $(I - \bar{\Lambda}_t) X_t \bar{\Lambda}_t'$ is p.s.d. Note that since

$$X_t \bar{\Lambda}_t' = (I - \bar{\Lambda}_t) X_t \bar{\Lambda}_t' + \bar{\Lambda}_t X_t \bar{\Lambda}_t',$$

and X_t is necessarily a p.s.d. matrix, it follows from the assumption that $(I - \bar{\Lambda}_t) X_t \bar{\Lambda}_t'$ is p.s.d. that $\bar{\Lambda}_t X_t = X_t \bar{\Lambda}_t'$ will also be a p.s.d. matrix; but this latter condition is weaker than the one assumed in our definition of the set $\mathcal{L}(X_t)$. This constitutes the complete set of conditions that must be satisfied for (C.4) to represent a memory structure consistent with our proposed normalization of the vector m_{t+1} .

We can further specialize these conditions in the case that $\bar{\Lambda}_t$ is a singular matrix. (Here we assume that X_t is of full rank.) If $\bar{\Lambda}_t$ is of rank one (or less), it can be written in the

form $\bar{\Lambda}_t = u_t v_t'$, where we are furthermore free to normalize the vector v_t' so that $v_t' X_t v_t = 1$. Then the condition that $\bar{\Lambda}_t X_t = X_t \bar{\Lambda}_t'$ will hold only if $u_t (v_t' X_t) = (X_t v_t) u_t'$. This means that u_t must be collinear with $X_t v_t$, so that we must be able to write $u_t = \lambda_t X_t v_t$, for some scalar λ_t . Thus in the singular case, we must be able to write

$$\bar{\Lambda}_t = \lambda_t X_t v_t v_t', \quad (\text{D.2})$$

where λ_t is a scalar and v_t is a vector such that $v_t' X_t v_t = 1$. Then

$$(I - \bar{\Lambda}_t) X_t \bar{\Lambda}_t' = \lambda_t (1 - \lambda_t) (X_t v_t) (X_t v_t)'$$

will be a p.s.d. matrix if and only if in addition $0 \leq \lambda_t \leq 1$. Thus a singular matrix $\bar{\Lambda}_t$ is an element of $\mathcal{L}(X_t)$ if and only if it is of the form (D.2) with $0 \leq \lambda_t \leq 1$ and v_t a vector such that $v_t' X_t v_t = 1$.

Consistency with the proposed normalization of m_{t+1} then further requires that

$$\Sigma_{\bar{\omega}, t+1} = \lambda_t (1 - \lambda_t) X_t v_t v_t' X_t. \quad (\text{D.3})$$

This implies that $\Sigma_{\bar{\omega}, t+1}$ is a singular matrix; the random vector $\bar{\omega}_{t+1}$ can be written as $\bar{\omega}_{t+1} = X_t v_t \tilde{\omega}_{t+1}$, where $\tilde{\omega}_{t+1}$ is a scalar random variable, with distribution $N(0, \lambda_t (1 - \lambda_t))$. It follows that in such a case, the memory state can be given a one-dimensional representation, writing $\bar{m}_{t+1} = X_t v_t \tilde{m}_{t+1}$, where the scalar memory state $\tilde{m}_{t+1}$ has a law of motion

$$\tilde{m}_{t+1} = \lambda_t v_t' \bar{s}_t + \tilde{\omega}_{t+1}, \quad \tilde{\omega}_{t+1} \sim N(0, \lambda_t (1 - \lambda_t)). \quad (\text{D.4})$$

In the case that $X_t = X_0$ (the only case in which it is possible for $X_t = X(\hat{\sigma}_t^2)$ to be singular), we have defined $\mathcal{L}(X_0)$ to include only matrices of the special form (2.11) with $0 \leq \lambda_t \leq 1$. In this case, $\bar{\Lambda}_t$ is necessarily of the form (D.2), with the vector v_t given by (2.20). Hence our comments above about the case in which $\bar{\Lambda}_t$ is singular apply also in the case in which X_t is singular, except that in this latter case we have the further restriction that v_t must be given by (2.20). In this special case, (D.3) reduces to

$$\Sigma_{\bar{\omega}, t+1} = \lambda_t (1 - \lambda_t) [\Omega + \sigma_y^2] w w'.$$

D.1 The degree of uncertainty implied by a given memory structure

We turn now to the question of how the posterior uncertainty Σ_{t+1} in the following period is determined by the law of motion for the memory state $\bar{m}_{t+1}$ that can be accessed at that time. Note that the variance of the marginal distribution for x_{t+1} can be decomposed as

$$\text{var}[x_{t+1}] = \text{E}[\text{var}[x_{t+1} | m_{t+1}]] + \text{var}[\text{E}[x_{t+1} | m_{t+1}]],$$

where in the first term on the right-hand side, the variance refers to the distribution of values for x_{t+1} conditional on the realization of m_{t+1} , and the expectation is over realizations of m_{t+1} , while in the second term the variance refers to the distribution of values for m_{t+1} , and the expectation is over values of x_{t+1} conditional on the realization of m_{t+1} . Since the

marginal distribution for x_{t+1} is the same for all t , and coincides with the prior distribution for x_0 specified in (B.2), the left-hand side must equal the matrix Σ_0 defined there. Hence the variance decomposition can be written as

$$\Sigma_0 = \Sigma_{t+1} + \text{var}[\bar{m}_{t+1}],$$

which implies that in any period,

$$\Sigma_{t+1} = \Sigma_0 - \text{var}[\bar{m}_{t+1}].$$

Thus in order to understand how the choice of $\bar{\Lambda}_t$ determines Σ_{t+1} , it suffices that we determine the implications for the degree of variation in $\bar{m}_{t+1}$.

A law of motion of the form (C.4) implies that

$$\begin{aligned} \text{var}[\bar{m}_{t+1}] &= \bar{\Lambda}_t X_t \bar{\Lambda}_t' + \Sigma_{\bar{\omega}, t+1} \\ &= \bar{\Lambda}_t X_t \bar{\Lambda}_t' + (I - \bar{\Lambda}_t) X_t \bar{\Lambda}_t' \\ &= X_t \bar{\Lambda}_t', \end{aligned}$$

where the second line uses (D.1). Hence we obtain the prediction that

$$\Sigma_{t+1} = \Sigma_0 - X_t \bar{\Lambda}_t'. \quad (\text{D.5})$$

Note that for any $\bar{\Lambda}_t \in \mathcal{L}(X_t)$, this must be a symmetric, p.s.d. matrix.

Hence for any value of $\hat{\sigma}_t^2$ satisfying $0 \leq \hat{\sigma}_t^2 \leq \hat{\sigma}_0^2$ and any transition matrix $\bar{\Lambda}_t \in \mathcal{L}(X(\hat{\sigma}_t^2))$, we can substitute $X_t = X(\hat{\sigma}_t^2)$ and the value of Σ_{t+1} given by (D.5) into (2.2) to obtain a solution for $\hat{\sigma}_{t+1}^2$ as a function of $\hat{\sigma}_t^2$ and $\bar{\Lambda}_t$. This defines the function $f(\hat{\sigma}_t^2, \bar{\Lambda}_t)$ referred to in the main text. We can then define $\mathcal{L}^{seq}$ as the set of sequences of transition matrices $\{\bar{\Lambda}_t\}$ for all $t \geq 0$ such that

$$\bar{\Lambda}_0 \in \mathcal{L}(X_0), \quad \bar{\Lambda}_1 \in \mathcal{L}(X(f(\hat{\sigma}_0^2, \bar{\Lambda}_0))), \quad \bar{\Lambda}_2 \in \mathcal{L}(X(f(f(\hat{\sigma}_0^2, \bar{\Lambda}_0), \bar{\Lambda}_1))),$$

and so on.

Then given any sequence of transition matrices $\{\bar{\Lambda}_t\} \in \mathcal{L}^{seq}$, there will be uniquely defined sequences $\{\hat{\sigma}_t^2, X_t\}$ for all $t \geq 0$. Equation (D.5), together with (B.2), can then be used to uniquely define the implied sequence of matrices $\{\Sigma_t\}$ for all $t \geq 0$. These matrices can in turn be used in (2.3) to define the Kalman gain γ_{1t} for each $t \geq 0$. Thus for any sequence of transition matrices $\{\bar{\Lambda}_t\} \in \mathcal{L}^{seq}$, there will be uniquely determined sequences $\{\Sigma_t, \gamma_{1t}, \hat{\sigma}_t^2, X_t\}$, as stated in the text. These in turn will imply a uniquely determined sequence of losses $\{MSE_t\}$ from forecast inaccuracy, using (2.5).

D.2 The mutual information implied by a given memory structure

Finally, we compute the mutual information I_t in the case that the memory state consists only of a reduced memory state $\bar{m}_{t+1}$, with law of motion (C.4). We first review the definition of mutual information in the case of continuously distributed random variables.

Let X and Y be two random variables, each parameterized using a finite system of coordinates (so that realizations x and y are each represented by finite-dimensional vectors),

and suppose that at least Y has a continuous distribution, with a density function $p(y|x)$ such that $p(y|x) > 0$ for all y in the support of Y and all x in the support of X . Suppose also that the marginal distribution for Y can be characterized by a density function $p(y) = \mathbb{E}[p(Y|x)]$, where the expectation is over possible realizations of x , and $p(y) > 0$ for all y in the support of Y . Then we can measure the degree to which knowing the realization of x changes the distribution that one can expect y to be drawn from by the Kullback-Liebler divergence (or relative entropy) of the conditional distribution $p(y|x)$ relative to the marginal distribution $p(y)$, defined as

$$D_{KL}(p(\cdot|x)||p(\cdot)) \equiv \mathbb{E} \left[\log \frac{p(y|x)}{p(y)} \right] \geq 0, \quad (\text{D.6})$$

where the expectation is over possible realizations of y , and this quantity is a function of the particular realization x .⁵⁵ The *mutual information* $I(X; Y)$ can then be defined as the mean value of this expression,

$$I(X; Y) \equiv \mathbb{E}[D_{KL}(p(\cdot|x)||p(\cdot))], \quad (\text{D.7})$$

where the expectation is now over possible realization of x , and the mutual information is also necessarily non-negative.⁵⁶

This definition of the mutual information has the attractive feature of being independent of the coordinates used to parameterize the realizations of the variable Y . Suppose that we write $y = \phi(z)$, where $\phi(\cdot)$ is an invertible smooth coordinate transformation between two Euclidean spaces of the same dimension. Then corresponding to the conditional density $p(y|x)$ for any x , there will be a corresponding density function $\tilde{p}(z|x)$ for the random variable Z (which is just the variable Y described using the alternative coordinate system), such that $\tilde{p}(z|x) = p(\phi(z)|x) \cdot D\phi(z)$ for each z , where $D\phi(z)$ is the Jacobian matrix of the coordinate transformation, evaluated at z . It follows that for any z in the support of Z and any x in the support of X ,

$$\frac{p(\phi(z)|x)}{p(\phi(z))} = \frac{\tilde{p}(z|x)}{\tilde{p}(z)},$$

so that

$$D_{KL}(p(\cdot|x)||p(\cdot)) = D_{KL}(\tilde{p}(\cdot|x)||\tilde{p}(\cdot))$$

for all x . We thus find that the mutual information $I(X; Y)$ will be the same as $I(X; Z)$: it is unaffected by a change in the coordinates used to parameterize Y .⁵⁷

We can similarly define the mutual information in a case in which the support of Y is not the entire Euclidean space, because of the existence of redundant coordinates in the parameterization of realizations y . Suppose that all vectors y in the support of Y are of the form $y = \phi(z)$, where $\phi(\cdot)$ is a smooth embedding of some lower-dimensional Euclidean space (the support of Z) into a higher-dimensional Euclidean space. Then the information about

⁵⁵The value of this quantity is necessarily non-negative because of Jensen's inequality, owing to the concavity of the logarithm.

⁵⁶Note that this definition — rather than the one often given in terms of the average reduction in the entropy of Y from observing X — has the advantage of remaining well-defined even when the random variable Y has a continuous distribution. See Cover and Thomas (2006) for further discussion.

⁵⁷It is equally unaffected by a change in the coordinates used to parameterize X , though we need not show this here.

the possible realizations of y contained in a realization of x is given by the information that x contains about the possible realizations of z . If the joint distribution of X and Z is such that we can define conditional density functions $\tilde{p}(z|x)$, with $\tilde{p}(z|x) > 0$ for all z and x , and a marginal density function $\tilde{p}(z) > 0$ for all z , then we can define the mutual information between X and Z using (D.7) as above. Since mutual information should be independent of the coordinates used to parameterize the variables, we can use the value of $I(X; Z)$ as our definition of $I(X; Y)$ in this case as well (even though expression (D.6) is not defined in this case).

In the case of interest in this paper, X and Y are variables with a joint distribution that is multivariate Gaussian. Let us consider first the generic case in which the conditional variance-covariance matrix $\text{var}[Y|x]$ is of full rank. (Note that this matrix will be independent of the realization of x , and so can be written $\text{var}[Y|X]$, to emphasize that only the parameters of the joint distribution matter.) In this case $\text{var}[Y]$ is of full rank as well, and for any x and y , the ratio of the density functions satisfies

$$\begin{aligned} \log \frac{p(y|x)}{p(y)} &= -\frac{1}{2} \log \frac{\det(\text{var}[Y|x])}{\det(\text{var}[Y])} \\ &\quad - \frac{1}{2} (y - \text{E}[y|x])' \text{var}[Y|x]^{-1} (y - \text{E}[y|x]) + \frac{1}{2} (y - \text{E}[y])' \text{var}[Y]^{-1} (y - \text{E}[y]). \end{aligned}$$

Hence for any x , we have

$$D_{KL}(x) = -\frac{1}{2} \log \frac{\det(\text{var}[Y|x])}{\det(\text{var}[Y])},$$

and since this will be independent of the realization of x , we similarly will have

$$I(X; Y) = -\frac{1}{2} \log \frac{\det(\text{var}[Y|X])}{\det(\text{var}[Y])}. \quad (\text{D.8})$$

One case in which $\text{var}[Y|x]$ will not be of full rank is if $y = Uz$ for some matrix U , where z is a random vector of lower dimension than that of y . (In this case, the rank of $\text{var}[Y|x]$ cannot be greater than the rank of $\text{var}[Z|x]$, which is at most the dimension of z .) Let us suppose that the rank of U is equal to the dimension of z , so that any vector $y = Uz$ is associated with exactly one vector z . In such a case we can, as discussed above, define the mutual information between X and Y to equal the mutual information between X and Z . If $\text{var}[Z|x]$ is of full rank, then we can use the calculations of the previous paragraph to show that

$$I(X; Y) = I(X; Z) = -\frac{1}{2} \log \frac{\det(\text{var}[Z|X])}{\det(\text{var}[Z])}. \quad (\text{D.9})$$

We turn now to the calculation of the mutual information between the reduced cognitive state $\bar{s}_t$ and the memory state $\bar{m}_{t+1}$, in the case of a law of motion of the form (C.4) for the memory state. We first consider the case in which X_t is of full rank (which, as noted in the text, will be true except when the memory state m_t is completely uninformative). If $\bar{\Lambda}_t$ and $I - \bar{\Lambda}_t$ are also both matrices of full rank, then

$$\text{var}[\bar{m}_{t+1} | \bar{s}_t] = \Sigma_{\bar{\omega}, t+1} = (I - \bar{\Lambda}_t) X_t \bar{\Lambda}_t'$$

will be of full rank, and

$$\text{var}[\tilde{m}_{t+1}] = \bar{\Lambda}_t X_t \bar{\Lambda}_t' + \Sigma_{\bar{\omega}, t+1} = X_t \bar{\Lambda}_t'$$

will be of full rank as well. We can then apply (D.8) to obtain

$$I_t = -\frac{1}{2} \log \frac{\det[(I - \bar{\Lambda}_t) X_t \bar{\Lambda}_t']}{\det[X_t \bar{\Lambda}_t']} = -\frac{1}{2} \log \det(I - \bar{\Lambda}_t), \quad (\text{D.10})$$

in conformity with equation (2.12) in the text.

In the case that X_t is of full rank, but $\bar{\Lambda}_t$ is varied so that one of its eigenvalues approaches 1 (meaning that $I - \bar{\Lambda}_t$ approaches a singular matrix, while the determinant of $\bar{\Lambda}_t$ remains bounded away from zero), the value of I_t implied by (D.10) grows without bound. It thus makes sense to assign a value of $+\infty$ to the mutual information in the case that $\bar{\Lambda}_t$ is of full rank but $I - \bar{\Lambda}_t$ is not. Note that in this case there is a linear combination of the elements of $\bar{s}_t$ that is revealed with perfect precision by the memory state (since $\Sigma_{\bar{\omega}, t+1}$ will be singular), while this linear combination is a continuous random variable with positive variance (since X_t is of full rank). This is not consistent with any finite value for the mutual information (and so cannot represent a feasible memory structure).

Suppose instead that while X_t is of full rank, $\bar{\Lambda}_t$ is only of rank one. In this case, we have shown above that $\bar{\Lambda}_t$ must be of the form (D.2), as a consequence of which $\Sigma_{\bar{\omega}, t+1}$ must be given by (D.3). In this case, the memory state can be represented in the form $\tilde{m}_{t+1} = X_t v_t \cdot \tilde{m}_{t+1}$, where $\tilde{m}_{t+1}$ is a scalar random variable with law of motion (D.4). This implies that $\text{var}[\tilde{m}_{t+1} | s_t] = \text{var}[\tilde{\omega}_{t+1}] = \lambda_t(1 - \lambda_t)$, while $\text{var}[\tilde{m}_{t+1}] = \lambda_t$. In the case that $0 < \lambda_t < 1$, we can then apply (D.9) to show that

$$I_t = -\frac{1}{2} \log \frac{\lambda_t(1 - \lambda_t)}{\lambda_t} = -\frac{1}{2} \log(1 - \lambda_t), \quad (\text{D.11})$$

Since in this case, $\det(I - \hat{\Lambda}_t) = \det(I - \lambda_t v_t v_t') = 1 - \lambda_t$, result (D.11) is again just what (D.10) would imply, so that (D.10) continues to be correct even though $\bar{\Lambda}_t$ is singular.

If we consider a sequence of matrices of this kind in which λ_t approaches 1, the mutual information (D.11) grows without bound. Thus we can assign the value $+\infty$ to I_t in the case that $\bar{\Lambda}_t$ is a matrix of rank one with $\lambda_t = 1$. Indeed, in this case, the memory state reveals with perfect precision the value of $v_t' \bar{s}_t$, a continuous random variable with positive variance (under the assumption that X_t is of full rank); but this is not possible in the case of any finite bound on mutual information. Hence (D.10) can be applied to this case as well.

Suppose instead that X_t is of full rank, but $\bar{\Lambda}_t = 0$. In this case, the distribution of $\tilde{m}_{t+1}$ is independent of the value of s_{t+1} , and the mutual information between these two variables must be zero. This is also what (D.10) would imply, so that (D.10) is correct in this case as well.

Finally, consider the case in which $X_t = X_0$, the only possible case in which X_t is not of full rank. In this case, we have defined $\mathcal{L}(X_0)$ to consist only of matrices of the form (D.2), with the vector v_t given by (2.20). If $\lambda_t = 0$, then the entire matrix $\bar{\Lambda}_t = 0$, and the argument in the previous paragraph again applies. Suppose instead that $\lambda_t > 0$. Just as in the discussion above of the case of a singular transition matrix, the memory state can be

represented by a scalar state variable $\tilde{m}_{t+1}$ with law of motion (D.4), and we can apply (D.9) to show that I_t will be given by (D.11). Again this is just what (D.10) would imply, so that (D.10) also yields the correct conclusion when X_t is a singular matrix.

Thus in all cases, (D.10) applies, and the value of I_t depends only on the choice of the transition matrix $\bar{\Lambda}_t$. It follows that for any sequence of transition matrices $\{\bar{\Lambda}_t\} \in \mathcal{L}^{seq}$, there will be uniquely defined sequences $\{MSE_t, I_t\}$, allowing the objective (1.6) to be evaluated.

E Recursive Determination of the Optimal Memory Structure

We have shown in the text how the optimal memory structure can be characterized if we can find the value function $V(\hat{\sigma}_t^2)$ that satisfies the Bellman equation

$$V(\hat{\sigma}_t^2) = \min_{\bar{\Lambda}_t \in \mathcal{L}(X(\hat{\sigma}_t^2))} [\alpha \hat{\sigma}_t^2 + c(I(\bar{\Lambda}_t)) + \beta V(f(\hat{\sigma}_t^2, \lambda_t, v_t))]. \quad (\text{E.1})$$

Here we establish some properties of the solution to the optimization problem on the right-hand side of (E.1) for an arbitrary function $V \in \mathcal{F}$, which we can then be used to establish properties of the value function $V(\hat{\sigma}_t^2)$ that solves this equation, and properties of the optimal memory structure.

E.1 Monotonicity of the value function

We first show that, for any function V that may be assumed in the problem on the right-hand side of (E.1), the minimum achievable value of the right-hand side is a monotonically increasing function of $\hat{\sigma}_t^2$. This in turn implies that the value function (which must satisfy (E.1)) must be a monotonically increasing function of its argument.

Fix any value function V to be used in the problem on the right-hand side of (E.1), and consider any two possible degrees of uncertainty $\hat{\sigma}_a^2, \hat{\sigma}_b^2$, satisfying

$$0 \leq \hat{\sigma}_a^2 < \hat{\sigma}_b^2 \leq \sigma_0^2. \quad (\text{E.2})$$

Let $\bar{\Lambda}_t = \bar{\Lambda}_b$ be some element of $\mathcal{L}(X(\hat{\sigma}_b^2))$, and thus a feasible memory structure when $\hat{\sigma}_t^2 = \hat{\sigma}_b^2$, and let us further suppose that $I(\bar{\Lambda}_b) < \infty$, as must be true of an optimal memory structure. We wish to show that we can choose a transition matrix $\bar{\Lambda}_a \in \mathcal{L}(X(\hat{\sigma}_a^2))$ such that

$$f(\hat{\sigma}_a^2, \bar{\Lambda}_a) = f(\hat{\sigma}_b^2, \bar{\Lambda}_b), \quad (\text{E.3})$$

and in addition

$$I(\bar{\Lambda}_a) \leq I(\bar{\Lambda}_b). \quad (\text{E.4})$$

That is, in the case of the smaller degree of uncertainty $\hat{\sigma}_a^2$ in the cognitive state in period t , it is possible to choose a memory structure that implies exactly the same degree of uncertainty in period $t+1$, and hence the same value for $V(\hat{\sigma}_{t+1}^2)$, at no greater an information cost, and thus it is possible to achieve a strictly lower value for the right-hand side of (E.1).

If we can show this for an arbitrary transition matrix $\bar{\Lambda}_b \in \mathcal{L}(X(\hat{\sigma}_b^2))$, then it is also true when $\bar{\Lambda}_b$ is the transition matrix associated with the optimal memory structure (the solution

to the problem on the right-hand side of (E.1)) when $\hat{\sigma}_t^2 = \hat{\sigma}_b^2$. This implies that it is possible to achieve a lower value for the right-hand side of (E.1) when $\hat{\sigma}_t^2 = \hat{\sigma}_a^2$ than it is possible to achieve when $\hat{\sigma}_t^2 = \hat{\sigma}_b^2$. Since this must be true for any values of $\hat{\sigma}_a^2, \hat{\sigma}_b^2$ consistent with (E.2), the right-hand side of (E.1) defines a monotonically increasing function of $\hat{\sigma}_t^2$.

To show that such a construction is always possible, let us first consider the case in which $\hat{\sigma}_b^2 = \hat{\sigma}_0^2$, so that the memory state m_t is completely uninformative in case b . In this case, the assumption that $\bar{\Lambda}_b \in \mathcal{L}(X(\hat{\sigma}_b^2)) = \mathcal{L}(X_0)$ requires that

$$\bar{\Lambda}_b = \lambda_b \frac{ww'}{w'w}$$

for some $0 \leq \lambda_b < 1$.⁵⁸ In this case, the memory structure for the following period is equivalent to one in which there is a univariate memory state

$$\tilde{m}_b = \frac{\lambda_b}{(\Omega + \sigma_y^2)^{1/2}} y_t + \tilde{\omega}_b, \quad \tilde{\omega}_b \sim N(0, \lambda_b(1 - \lambda_b)).$$

The implied uncertainty in the following period (given the memory state, but before y_{t+1} is observed) is then given by

$$\Sigma_{t+1} = \Sigma_0 - \lambda_b(\Omega + \sigma_y^2)ww'. \quad (\text{E.5})$$

Now let $\bar{s}_a$ be the reduced cognitive state in period t , in the case of a more informative memory structure that implies the lower degree of uncertainty $\hat{\sigma}_a^2$, and let $X_a \equiv X(\hat{\sigma}_a^2)$ be the variance of this random vector. In this case, we can choose a memory structure for the following period defined by the transition matrix

$$\bar{\Lambda}_a = \lambda_b X_a \frac{e_2 e_2'}{\Omega + \sigma_y^2}$$

where $e_2 \equiv [0 \ 1]'$. This is a matrix of the form (D.2), and hence an element of $\mathcal{L}(X_a)$. Because $\bar{\Lambda}_a$ is singular, the specified memory structure is equivalent to one in which there is a univariate memory state

$$\tilde{m}_a = \lambda_b \frac{e_2' \bar{s}_a}{(e_2' X_a e_2)^{1/2}} + \tilde{\omega}_a, \quad \tilde{\omega}_a \sim N(0, \lambda_b(1 - \lambda_b)).$$

But this means that

$$\tilde{m}_a = \frac{\lambda_b}{(\Omega + \sigma_y^2)^{1/2}} y_t + \tilde{\omega}_a, \quad \tilde{\omega}_a \sim N(0, \lambda_b(1 - \lambda_b)).$$

Hence the joint distribution of $(\tilde{m}_a, x_{t+1})$ is identical to the joint distribution of $(\tilde{m}_b, x_{t+1})$, and the implied uncertainty in the following period given this memory structure is again given by (E.5). Hence the value of $\hat{\sigma}_{t+1}^2$ implied by memory structure a is the same as that

⁵⁸The upper bound is required in order to satisfy the assumption that $I(\bar{\Lambda}_b) < \infty$.

implied by memory structure b . This establishes condition (E.3). Moreover, for both memory structures we have the same mutual information,

$$I(\bar{\Lambda}_a) = I(\bar{\Lambda}_b) = -\frac{1}{2} \log(1 - \lambda_b).$$

This establishes condition (E.4). Hence the value of the right-hand side of (E.1) must be lower when $\hat{\sigma}_t^2 = \hat{\sigma}_a^2$.

Let us next consider the less trivial case in which $0 < \hat{\sigma}_b^2 < \hat{\sigma}_0^2$. Let $\bar{s}_b$ be the reduced cognitive state in period t that implies a degree of uncertainty $\hat{\sigma}_b^2$, and let $X_b \equiv X(\hat{\sigma}_b^2)$ be the variance of this random vector. Let the optimal memory structure for the following period (the solution to the problem on the right-hand side of (E.1)) in this case be

$$\bar{m}_b = \bar{\Lambda}_b \bar{s}_b + \bar{\omega}_b, \quad (\text{E.6})$$

where

$$\bar{\Lambda}_b \in \mathcal{L}(X_b), \quad \bar{\omega}_b \sim N(0, (I - \bar{\Lambda}_b)X_b\bar{\Lambda}_b').$$

The implied uncertainty in the following period will then be given by

$$\Sigma_{t+1} = \Sigma_0 - X_b \bar{\Lambda}_b'. \quad (\text{E.7})$$

Let us consider the memory structure for cognitive state a defined by the transition matrix

$$\bar{\Lambda}_a = \bar{\Lambda}_b \Gamma \Psi \Gamma^{-1}, \quad (\text{E.8})$$

where Γ is the invertible matrix defined in (C.5), and

$$\Psi \equiv \begin{bmatrix} \psi & 0 \\ 0 & 1 \end{bmatrix},$$

where $0 < \psi < 1$ is the quantity

$$\psi \equiv \frac{\hat{\sigma}_0^2 - \hat{\sigma}_b^2}{\hat{\sigma}_0^2 - \hat{\sigma}_a^2}.$$

Note that Ψ is a diagonal matrix, with the property that

$$\Psi \check{X}_a = \check{X}_a \Psi = \check{X}_b,$$

using the notation $\check{X}_i \equiv \check{X}(\hat{\sigma}_i^2)$ for $i = a, b$, where $\check{X}(\hat{\sigma}_i^2)$ is the function defined in (C.6). It is first necessary to verify that $\bar{\Lambda}_a \in \mathcal{L}(X_a)$, so that this matrix defines a possible memory structure.

We first show that $\bar{\Lambda}_a X_a = X_a \bar{\Lambda}_a'$. Definition (E.8) implies that

$$\begin{aligned} \bar{\Lambda}_a X_a &= \bar{\Lambda}_b \Gamma \Psi \Gamma^{-1} X_a \\ &= \bar{\Lambda}_b \Gamma \Psi \check{X}_a \Gamma' \\ &= \bar{\Lambda}_b \Gamma \check{X}_b \Gamma' \\ &= \bar{\Lambda}_b X_b. \end{aligned}$$

The fact that $\bar{\Lambda}_b \in \mathcal{L}(X_b)$ implies that $\bar{\Lambda}_b X_b$ must be a symmetric matrix; hence $\bar{\Lambda}_a X_a$, which is the same matrix, must also be symmetric. Thus $\bar{\Lambda}_a X_a = X_a \bar{\Lambda}'_a$.

Next, we must also show that $(I - \bar{\Lambda}_a)X_a \bar{\Lambda}'_a$ is a p.s.d. matrix. We first note that $I - \Psi$ is a diagonal matrix with non-negative elements on the diagonal; it follows that $(I - \Psi)\check{X}_b$ is also a diagonal matrix with non-negative elements on the diagonal, and hence p.s.d. From this it follows that

$$\begin{aligned} \bar{\Lambda}_b \Gamma \cdot (I - \Psi) \check{X}_b \cdot \Gamma' \bar{\Lambda}'_b &= \bar{\Lambda}_b \Gamma (\check{X}_b - \Psi \check{X}_a \Psi) \Gamma' \bar{\Lambda}'_b \\ &= \bar{\Lambda}_b (\Gamma \check{X}_b \Gamma') \bar{\Lambda}'_b - (\bar{\Lambda}_b \Gamma \Psi \Gamma^{-1}) (\Gamma \check{X}_a \Gamma') (\bar{\Lambda}_b \Gamma \Psi \Gamma^{-1})' \\ &= \bar{\Lambda}_b X_b \bar{\Lambda}'_b - \bar{\Lambda}_a X_a \bar{\Lambda}'_a \\ &= (X_a \bar{\Lambda}'_a - \bar{\Lambda}_a X_a \bar{\Lambda}'_a) - (X_b \bar{\Lambda}'_b - \bar{\Lambda}_b X_b \bar{\Lambda}'_b) \\ &= (I - \bar{\Lambda}_a) X_a \bar{\Lambda}'_a - (I - \bar{\Lambda}_b) X_b \bar{\Lambda}'_b \end{aligned}$$

must be p.s.d. as well. But since the fact that $\bar{\Lambda}_b \in \mathcal{L}(X_b)$ implies that $(I - \bar{\Lambda}_b)X_b \bar{\Lambda}'_b$ must be p.s.d., it follows that $(I - \bar{\Lambda}_a)X_a \bar{\Lambda}'_a$ can be expressed as the sum of two p.s.d. matrices, and so must also be p.s.d. This verifies the second of the conditions required in order to show that $\bar{\Lambda}_a \in \mathcal{L}(X_a)$.

Thus if $\bar{s}_a$ is a reduced cognitive state for period t that implies a degree of uncertainty $\hat{\sigma}_a^2$, a possible memory structure for the following period is

$$\bar{m}_a = \bar{\Lambda}_a \bar{s}_a + \bar{\omega}_a, \quad (\text{E.9})$$

where the transition matrix $\bar{\Lambda}_a$ is defined in (E.8), and

$$\bar{\omega}_a \sim N(0, (I - \bar{\Lambda}_a)X_a \bar{\Lambda}'_a).$$

The implied uncertainty in the following period will then be given by

$$\Sigma_{t+1} = \Sigma_0 - X_a \bar{\Lambda}'_a.$$

This latter matrix is the same as the one in (E.7); it follows that the implied value of $\hat{\sigma}_{t+1}^2$ is also the same as for the memory structure (E.6). Thus we have shown that in the case of the smaller degree of uncertainty $\hat{\sigma}_a^2$, it is possible to choose a memory structure that implies exactly the same degree of uncertainty in period $t + 1$ as when the degree of uncertainty in period t is given by the larger quantity $\hat{\sigma}_b^2$.

It remains to be shown that memory structure (E.9) involves no greater information cost than memory structure (E.6). Consider first the case in which the memory state $\bar{m}_b$ is non-degenerate, in the sense that $\text{var}[\bar{m}_b] = X_b \bar{\Lambda}'_b$ is non-singular. It follows that the same must be true of memory state $\bar{m}_a$. Then for either of the two memory structures $i = a, b$ just discussed, (D.10) implies that the mutual information will be given by

$$I_t = -\frac{1}{2} \log \frac{\det[(I - \bar{\Lambda}_i)X_i \bar{\Lambda}'_i]}{\det[X_i \bar{\Lambda}'_i]}.$$

We have shown above that the value of the denominator in this expression is the same for $i = a, b$ (and under the assumption that $X_b \bar{\Lambda}'_b$ is non-singular, it must be positive). Hence

the relative size of the two mutual informations depends on the relative size of the numerator in the two cases. But we have shown above that $(I - \bar{\Lambda}_a)X_a\bar{\Lambda}'_a$ can be expressed as the sum of $(I - \bar{\Lambda}_b)X_b\bar{\Lambda}'_b$ plus a p.s.d. matrix. Since both of these matrices are also p.s.d., their determinants satisfy

$$\det[(I - \bar{\Lambda}_a)X_a\bar{\Lambda}'_a] \geq \det[(I - \bar{\Lambda}_b)X_b\bar{\Lambda}'_b] > 0,$$

where the final inequality is necessary in order for memory structure b to have a finite information cost. It follows that condition (E.4) must hold in this case.

Now suppose instead that $\text{var}[\bar{m}_b]$ is a singular matrix. In the case that the matrix is zero in all elements, $\bar{\Lambda}_b = 0$, and so (E.8) implies that $\bar{\Lambda}_a = 0$ as well. In this case, $\det(I - \bar{\Lambda}_a) = \det(I - \bar{\Lambda}_b) = 1$, so that $I(\bar{\Lambda}_a) = I(\bar{\Lambda}_b) = 0$, and (E.4) is satisfied in this case as well. Thus we need only consider further the case in which $\text{var}[\bar{m}_b]$ is of rank one, which requires that $\bar{\Lambda}_b$ be of rank one as well.

In this case, we can write

$$\bar{\Lambda}_b = \lambda_b X_b v_b v'_b,$$

where $0 < \lambda_b < 1$ ⁵⁹ and v_b is a vector such that $v'_b X_b v_b = 1$. All columns of $\bar{\Lambda}_b$ are multiples of the vector $X_b v_b$, and as a consequence the unique non-null right eigenvector of $\bar{\Lambda}_b$ is given by $X_b v_b$, with the associated eigenvalue λ_b . Alternatively, using the orthogonalized representation of the cognitive state introduced in section C.4, we can write

$$\Gamma^{-1} \bar{\Lambda}_b \Gamma = \lambda_b \check{X}_b \check{v}_b \check{v}'_b,$$

where we define $\check{v}_b \equiv \Gamma' v_b$, and note that $\check{v}'_b \check{X}_b \check{v}_b = 1$.

Then (E.8) implies that the columns of $\bar{\Lambda}_a$ must also all be multiples of the vector $X_b v_b$. It follows that $\bar{\Lambda}_a$ must also be singular, and that its unique non-null eigenvector must be $X_b v_b$, with an associated eigenvalue

$$\begin{aligned} \lambda_a &= \lambda_b v'_b \Gamma \Psi \Gamma^{-1} (X_b v_b) \\ &= \lambda_b \check{v}'_b \Psi \check{X}_b \check{v}_b \\ &= \lambda_b (\check{v}'_b \Psi^{1/2}) \check{X}_b (\Psi^{1/2} \check{v}_b) \\ &\leq \lambda_b \check{v}'_b \check{X}_b \check{v}_b = \lambda_b. \end{aligned}$$

Thus we must have

$$\det(I - \bar{\Lambda}_a) = (1 - \lambda_a) \geq (1 - \lambda_b) = \det(I - \bar{\Lambda}_b),$$

from which it follows that (E.4) must hold in this case as well.

Thus we have shown that whenever $\hat{\sigma}_a^2, \hat{\sigma}_b^2$ satisfy (E.2), for any memory structure for case b with a finite information cost, it is possible to choose a memory structure for case a satisfying both (E.3) and (E.4). This means that it must be possible to achieve a lower value for the right-hand side of (E.1) when $\hat{\sigma}_t^2 = \hat{\sigma}_a^2$ than when $\hat{\sigma}_t^2 = \hat{\sigma}_b^2$. This in turn implies that the right-hand side of (E.1) defines a monotonically increasing function of $\hat{\sigma}_t^2$, regardless of the nature of the function $V(\hat{\sigma}_{t+1}^2)$ that is assumed in this optimization problem. Hence the value function $V(\hat{\sigma}_t^2)$ defined by (E.1) must be a monotonically increasing function of its argument.

⁵⁹Again, the upper bound is required in order for $I(\bar{\Lambda}_b)$ to be finite.

E.2 Optimality of a unidimensional memory state

Here we establish, as stated in the text, that the matrix $\bar{\Lambda}_t$ that solves the problem

$$\min_{\bar{\Lambda}_t \in \mathcal{L}(X(\hat{\sigma}_t^2))} I(\bar{\Lambda}_t) \quad \text{s.t.} \quad f(\hat{\sigma}_t^2, \bar{\Lambda}_t) \leq \hat{\sigma}_{t+1}^2, \quad (\text{E.10})$$

for given values of $(\hat{\sigma}_t^2, \hat{\sigma}_{t+1}^2)$ is necessarily at most of rank one. As explained in the text, we need only consider the case in which $\hat{\sigma}_t^2 < \hat{\sigma}_0^2$. Given a matrix $\bar{\Lambda}_t$ of rank two that satisfies the constraint in (E.10), we wish to show that we can choose an alternative transition matrix of at most rank one, that also satisfies the constraint, but which achieves a lower value of $I(\bar{\Lambda}_t)$.

We first note that when $\hat{\sigma}_t^2 < \hat{\sigma}_0^2$, $X(\hat{\sigma}_t^2)$ is non-singular. Under the hypothesis that $\bar{\Lambda}_t$ is non-singular, it follows that $X_t \bar{\Lambda}_t'$ is non-singular as well (where we now simply write X_t for $X(\hat{\sigma}_t^2)$), and hence positive definite. Similarly, $\bar{\Lambda}_t X_t \bar{\Lambda}_t'$ must be non-singular and hence positive definite.

Then let the alternative transition matrix be given by

$$\bar{\Lambda}_t^{1D} = \lambda_t X_t v_t v_t', \quad (\text{E.11})$$

with

$$\lambda_t = \frac{\delta'_{t+1} \bar{\Lambda}_t X_t \bar{\Lambda}_t' \delta_{t+1}}{\delta'_{t+1} X_t \bar{\Lambda}_t' \delta_{t+1}}, \quad v_t = \frac{\bar{\Lambda}_t' \delta_{t+1}}{(\delta'_{t+1} \bar{\Lambda}_t X_t \bar{\Lambda}_t' \delta_{t+1})^{1/2}},$$

where $\delta_{t+1} \equiv e_1 - \gamma_{1,t+1}c$ is the vector introduced in (2.16), and let the matrix $\Sigma_{\bar{\omega},t+1}$ be correspondingly modified in the way specified by (2.9). The fact that $X_t \bar{\Lambda}_t'$ is positive definite implies that the denominator of the expression for λ_t is necessarily positive, so that this quantity is well-defined. Similarly, the fact that $\bar{\Lambda}_t X_t \bar{\Lambda}_t'$ is positive definite implies that the denominator of the expression for v_t is necessarily positive, so that this vector is well-defined as well.

In addition, the fact that (by assumption) $\bar{\Lambda}_t \in \mathcal{L}(X_t)$ implies that $(I - \bar{\Lambda}_t)X_t \bar{\Lambda}_t'$ must be p.s.d. From this it follows that

$$\delta'_{t+1} (I - \bar{\Lambda}_t) X_t \bar{\Lambda}_t' \delta_{t+1} \geq 0,$$

and hence that

$$\delta'_{t+1} X_t \bar{\Lambda}_t' \delta_{t+1} \geq \delta'_{t+1} \bar{\Lambda}_t X_t \bar{\Lambda}_t' \delta_{t+1} > 0,$$

where the final inequality follows from the fact that $\bar{\Lambda}_t X_t \bar{\Lambda}_t'$ is positive definite. Thus the proposed definition of λ_t satisfies $0 < \lambda_t \leq 1$. One also observes from the definition of v_t that $v_t' X_t v_t = 1$. These conditions suffice to establish that the alternative transition matrix $\bar{\Lambda}_t^{1D}$ is also an element of $\mathcal{L}(X_t)$. That is, it represents a feasible memory structure for period t , given the value of $\hat{\sigma}_t^2$.

This alternative transition matrix corresponds to a memory structure in which $\bar{m}_{t+1} = X_t v_t \bar{m}_{t+1}$, where $\bar{m}_{t+1}$ is the unidimensional memory state with law of motion (2.19). From this it follows that

$$\delta'_{t+1} \bar{m}_{t+1} = \lambda_t \delta'_{t+1} X_t v_t v_t' \bar{s}_t + \delta'_{t+1} X_t v_t \tilde{\omega}_{t+1}$$

will be a normally distributed random variable, with conditional first and second moments given by

$$\begin{aligned}
\mathbb{E}[\delta'_{t+1} \bar{m}_{t+1} | s_t] &= \lambda_t \delta'_{t+1} X_t v_t v'_t \bar{s}_t \\
&= \frac{\delta'_{t+1} \bar{\Lambda}_t X_t \bar{\Lambda}'_t \delta_{t+1}}{\delta'_{t+1} X_t \bar{\Lambda}'_t \delta_{t+1}} \frac{\delta'_{t+1} X_t \bar{\Lambda}'_t \delta_{t+1} \cdot \delta'_{t+1} \bar{\Lambda}_t \bar{s}_t}{\delta'_{t+1} \bar{\Lambda}_t X_t \bar{\Lambda}'_t \delta_{t+1}} \\
&= \delta'_{t+1} \bar{\Lambda}_t \bar{s}_t
\end{aligned}$$

and

$$\begin{aligned}
\text{var}[\delta'_{t+1} \bar{m}_{t+1} | s_t] &= \lambda_t (1 - \lambda_t) (\delta'_{t+1} X_t v_t)^2 \\
&= (1 - \lambda_t) \frac{\delta'_{t+1} \bar{\Lambda}_t X_t \bar{\Lambda}'_t \delta_{t+1}}{\delta'_{t+1} X_t \bar{\Lambda}'_t \delta_{t+1}} \frac{(\delta'_{t+1} X_t \bar{\Lambda}'_t \delta_{t+1})^2}{\delta'_{t+1} \bar{\Lambda}_t X_t \bar{\Lambda}'_t \delta_{t+1}} \\
&= (1 - \lambda_t) \delta'_{t+1} X_t \bar{\Lambda}'_t \delta_{t+1} \\
&= \delta'_{t+1} X_t \bar{\Lambda}'_t \delta_{t+1} - \delta'_{t+1} \bar{\Lambda}_t X_t \bar{\Lambda}'_t \delta_{t+1} \\
&= \delta'_{t+1} [(I - \bar{\Lambda}_t) X_t \bar{\Lambda}'_t] \delta_{t+1} \\
&= \delta'_{t+1} \Sigma_{\bar{\omega}_{t+1}} \delta_{t+1}.
\end{aligned}$$

These are the same conditional mean and variance as in the case of the memory structure specified by the transition matrix $\bar{\Lambda}_t$. Since the optimal estimate $\hat{\mu}_{t+1}$ depends on m_{t+1} only through the value of $\delta'_{t+1} \bar{m}_{t+1}$ (from equation (2.16)), it follows that the conditional distribution $\hat{\mu}_{t+1} | s_t, y_{t+1}$ will be the same under the alternative memory structure. This in turn implies that the variance of $\hat{\mu}_{t+1}$ will be the same, and hence that

$$\hat{\sigma}_{t+1}^2 = \Omega - \text{var}[\hat{\mu}_{t+1}]$$

will be the same. Thus $\bar{\Lambda}_t^{1D}$ also satisfies the constraint in (E.10).

Next we show that $I(\bar{\Lambda}_t^{1D})$ must be lower than $I(\bar{\Lambda}_t)$. Let u'_1 and u'_2 be the two left eigenvectors of $\bar{\Lambda}_t$, with associated eigenvalues μ_1 and μ_2 respectively, and let the eigenvectors be normalized so that $u'_i X_t u_i = 1$ for $i = 1, 2$. The corresponding right eigenvectors must then be $X_t u_1$ and $X_t u_2$ respectively. Thus we have

$$\bar{\Lambda}_t X_t u_i = \mu_i X_t u_i, \quad u'_i \bar{\Lambda}_t = \mu_i u'_i,$$

for $i = 1, 2$, and

$$u'_1 X_t u_1 = u'_2 X_t u_2 = 1, \quad u'_1 X_t u_2 = 0.$$

The vector δ'_{t+1} introduced in (2.16) can be written as a linear combination of the two left eigenvectors,

$$\delta'_{t+1} = \alpha_1 u'_1 + \alpha_2 u'_2,$$

for some coefficients α_1, α_2 . This implies that

$$\delta'_{t+1} X_t \bar{\Lambda}'_t \delta_{t+1} = \alpha_1^2 \mu_1 + \alpha_2^2 \mu_2,$$

$$\delta'_{t+1} \bar{\Lambda}_t X_t \bar{\Lambda}'_t \delta_{t+1} = \alpha_1^2 \mu_1^2 + \alpha_2^2 \mu_2^2,$$

and hence that

$$\lambda_t = \frac{\alpha_1^2 \mu_1}{\alpha_1^2 \mu_1 + \alpha_2^2 \mu_2} \mu_1 + \frac{\alpha_2^2 \mu_2}{\alpha_1^2 \mu_1 + \alpha_2^2 \mu_2} \mu_2.$$

Thus we see that λ_t must be a convex combination of μ_1 and μ_2 .

The fact that $\bar{\Lambda}_t \in \mathcal{L}(X_t)$ requires that both eigenvalues satisfy $0 \leq \mu_i \leq 1$, and the assumption that $\bar{\Lambda}_t$ is non-singular further requires that $\mu_i > 0$ for both. Thus we must have

$$1 - \mu_i > (1 - \mu_1)(1 - \mu_2)$$

for both $i = 1, 2$. Since λ_t is a convex combination of μ_1 and μ_2 , it follows that

$$1 - \lambda_t > (1 - \mu_1)(1 - \mu_2).$$

Thus

$$\det(I - \bar{\Lambda}_t^{1D}) = 1 - \lambda_t > (1 - \mu_1)(1 - \mu_2) = \det(I - \bar{\Lambda}_t).$$

Results (D.10) and (D.11) then imply that $I(\bar{\Lambda}_t^{1D}) < I(\bar{\Lambda}_t)$.

Thus $\bar{\Lambda}_t$ cannot be the solution to the optimization problem (E.10). Since this argument can be made in the case of any matrix $\bar{\Lambda}_t \in \mathcal{L}(X_t)$ that is of full rank, we conclude that the optimal transition matrix can be at most of rank one.

E.3 The optimal univariate memory state

We turn now to the question of which linear combination of the elements of the reduced cognitive state constitutes the single variable for which it is optimal to retain a noisy record in memory — that is, we wish to characterize the optimal weight vector v_t in (D.4). Here we take as given the value of λ_t (or equivalently, the mutual information between the period t cognitive state and the memory carried into period $t + 1$), and solve for the optimal choice of v_t for any given value of λ_t . With this in hand, it will then be possible to characterize an optimal memory structure in terms of the single parameter λ_t .

Given the value of $\hat{\sigma}_t^2$ and the matrix $X_t \equiv \text{var}[\bar{s}_t]$, and taking as given the value of λ_t , we wish to choose v_t so as to minimize $\hat{\sigma}_{t+1}^2$. Note that

$$\hat{\sigma}_{t+1}^2 = \min_{\xi, \gamma_1} \text{var}[\mu - \xi \tilde{m}_{t+1} - \gamma_1 y_{t+1}].$$

Hence we can write our problem as the choice of ξ, γ_1 , and the vector v_t so as to minimize

$$\begin{aligned} f(\hat{\sigma}_t^2, \lambda_t, v_t; \xi, \gamma_1) &\equiv \text{var}[\mu - \xi(\lambda_t v_t' \bar{s}_t + \tilde{\omega}_{t+1}) - \gamma_1 y_{t+1}] \\ &= \text{var}[\mu - \xi \lambda_t v_t' \bar{s}_t - \gamma_1 y_{t+1}] + \xi^2 \lambda_t (1 - \lambda_t), \end{aligned}$$

subject to the constraint that $v_t' X_t v_t = 1$. Note that the solution to this problem will simultaneously determine the optimal choice of v_t (and hence the optimal memory structure, given a choice of λ_t) and the coefficients of the optimal estimate

$$\hat{\mu}_{t+1} = \xi \tilde{m}_{t+1} + \gamma_1 y_{t+1} \tag{E.12}$$

based on that memory structure.

We can alternatively define this problem as the choice of a weighting vector $\psi \equiv \xi \lambda_t v_t$ and a Kalman gain γ_1 . The values of these quantities suffice to determine the value of the objective (if we know the values of $\hat{\sigma}_t^2$ and λ_t), since we can reconstruct ξ and v_t from them:

$$v_t = \frac{\psi}{(\psi' X_t \psi)^{1/2}}, \quad \xi = (\psi' X_t \psi)^{1/2} \lambda_t.$$

Moreover, there is no theoretical restriction on the elements of the vector ψ , since the scale factor ξ can be of arbitrary size in the previous formulation of the optimization problem. Thus we can alternatively state our problem as the choice of a weighting vector ψ and a Kalman gain γ_1 to minimize

$$f(\hat{\sigma}_t^2, \lambda_t; \psi, \gamma_1) = \text{var}[\mu - \psi' \bar{s}_t - \gamma_1 y_{t+1}] + \frac{1 - \lambda_t}{\lambda_t} \psi' X_t \psi. \quad (\text{E.13})$$

We can write the first term in this objective as

$$\begin{aligned} \text{var}[\mu - \psi' \bar{s}_t - \gamma_1 y_{t+1}] &= \text{var}[(1 - (1 - \rho)\gamma_1)(\mu - \hat{\mu}_t) - \gamma_1(y_{t+1} - \mu) + (e'_1 - \gamma_1 c') \bar{s}_t - \psi' \bar{s}_t] \\ &= (e'_1 - \gamma_1 c' - \psi') X_t (e_1 - \gamma_1 c - \psi) + (1 - (1 - \rho)\gamma_1)^2 \hat{\sigma}_t^2 + \gamma_1^2 \sigma_\epsilon^2. \end{aligned}$$

Substituting this into (E.13), we see that the objective is a strictly convex quadratic function of ψ and γ_1 , for any values of $\hat{\sigma}_t^2$ and λ_t . It follows that the objective has an interior minimum, given by the unique solution to the first-order conditions.

The FOCs for the minimization of (E.13) are given by the linear equations

$$\psi = \lambda_t (e_1 - \gamma_1 c), \quad (\text{E.14})$$

$$c' X_t (e_1 - \gamma_1 c - \psi) + (1 - \rho)(1 - (1 - \rho)\gamma_1) \hat{\sigma}_t^2 - \gamma_1 \sigma_\epsilon^2 = 0. \quad (\text{E.15})$$

Equation (E.14) already allows one valuable insight: the optimal weight vector v_t is simply a normalized version of the vector δ_{t+1} defined in (2.16). However, this does not yet tell us how to choose v_t , since the vector δ_{t+1} depends on the Kalman gain $\gamma_{1,t+1}$, which depends on the memory structure chosen in period t .

But together equations (E.14)–(E.15) provide a linear system that can be solved for ψ and γ_1 , given the values of $\hat{\sigma}_t^2$ and λ_t . We obtain

$$\gamma_{1,t+1} = \frac{(1 - \lambda_t)\Omega + \lambda_t(1 - \rho)\hat{\sigma}_t^2}{(1 - \lambda_t)(\Omega + \rho^2 \sigma_y^2) + \lambda_t(1 - \rho)^2 \hat{\sigma}_t^2 + \sigma_\epsilon^2} \quad (\text{E.16})$$

as an explicit solution for the Kalman gain. It is worth noting that this implies that

$$0 < \gamma_{1,t+1} < \frac{1}{1 - \rho}. \quad (\text{E.17})$$

We can then use this solution to evaluate the elements of the vector δ . We obtain

$$\delta_{1,t+1} \equiv 1 - (1 - \rho)\gamma_{1,t+1} = \frac{(1 - \lambda_t)\rho(\Omega + \rho \sigma_y^2) + \sigma_\epsilon^2}{(1 - \lambda_t)(\Omega + \rho^2 \sigma_y^2) + \lambda_t(1 - \rho)^2 \hat{\sigma}_t^2 + \sigma_\epsilon^2} > 0,$$

$$\delta_{2,t+1} \equiv -\rho\gamma_{1,t+1} = -\frac{(1-\lambda_t)\rho\Omega + \lambda_t\rho(1-\rho)\hat{\sigma}_t^2}{(1-\lambda_t)(\Omega + \rho^2\sigma_y^2) + \lambda_t(1-\rho)^2\hat{\sigma}_t^2 + \sigma_\epsilon^2} \leq 0.$$

The weight vector v_t is then just a normalized version of δ_{t+1} .

We note that when $\rho = 0$, the optimal weight vector has $v_2 = 0$; that is, the memory state $\tilde{m}_{t+1}$ is just a noisy record of $\hat{\mu}_t$. (This is intuitive, since when the state is i.i.d., and given the estimate $\hat{\mu}_t$ of the mean, the value of y_t provides no information about anything that needs to be estimated or forecasted in period $t + 1$ or later.) Instead when $\rho > 0$, we see that the sign of v_2 is necessarily opposite to the sign of v_1 : the optimal memory state averages $\hat{\mu}_t$ and y_t with a negative relative weight on y_t .

Given this solution for γ_1 , the implied solution for the vector ψ is given by (E.14). Substituting the solutions for γ_1 and ψ into the quadratic objective, we obtain for the minimum possible value of the objective

$$\hat{\sigma}_{t+1}^2 = (1-\lambda_t)\delta'_{t+1}\Sigma_0\delta_{t+1} + \lambda_t(\delta_{1,t+1})^2\hat{\sigma}_t^2 + \gamma_{1,t+1}^2\sigma_\epsilon^2. \quad (\text{E.18})$$

This provides an equation for the evolution of the uncertainty measure $\hat{\sigma}_{t+1}^2$, given a choice each period of λ_t , and using the formulas above for the values of $\gamma_{1,t+1}$ and δ_{t+1} .

F Numerical Solutions

Here we provide further details of the numerical calculations reported in section 3 of the main text.

F.1 Dynamics of uncertainty given the path of $\{\lambda_t\}$

We begin by discussing our approach to numerical solution for the law of motion $\eta_{t+1} = \phi(\eta_t; \lambda_t)$ for the scaled uncertainty measure $\{\eta_t\}$, given a path for the memory-sensitivity coefficient $\{\lambda_t\}$. In terms of this rescaled state variable, the law of motion (E.18) becomes

$$\eta_{t+1} = (1-\lambda_t)(1-\gamma_{1,t+1})^2K + (1-\rho^2\lambda_t)\gamma_{1,t+1}^2 + \lambda_t(1-(1-\rho)\gamma_{1,t+1})^2\eta_t, \quad (\text{F.1})$$

and (E.16) becomes

$$\gamma_{1,t+1} = \frac{(1-\lambda_t)K + (1-\rho)\lambda_t\eta_t}{(1-\lambda_t)(K + \rho^2) + (1-\rho^2) + (1-\rho)^2\lambda_t\eta_t}. \quad (\text{F.2})$$

Substitution of (F.2) for $\gamma_{1,t+1}$ in the right-hand side of (F.1) yields an analytical expression for the function $\phi(\eta_t; \lambda_t)$.

This result suffices to allow us to compute the optimal dynamics of the uncertainty measure $\{\eta_t\}$ in the case that the only limit on the complexity of memory is an upper bound $\lambda_t \leq \bar{\lambda} < 1$ each period. We observe from (E.13) that the objective $f(\hat{\sigma}_t^2, \lambda_t; \psi, \gamma_1)$ is minimized, for given values of the other parameters, by making λ_t as large as possible. Hence the same is true for the function $f(\hat{\sigma}_t^2, \lambda_t, v_t)$ obtained by minimizing the objective over possible choices of ξ and γ_1 . It follows that it will be optimal to choose $\lambda_t = \bar{\lambda}$ each period in the case of this kind of constraint.

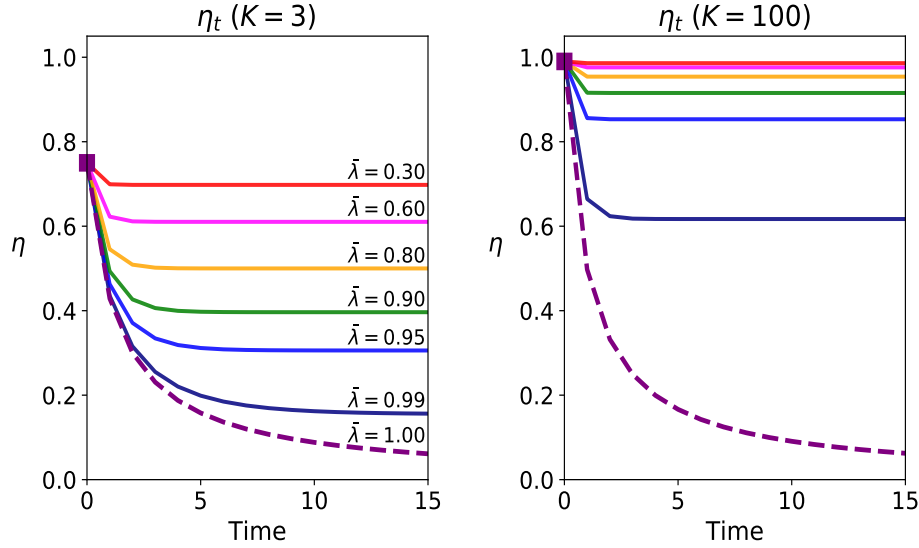


Figure 9: The evolution of scaled uncertainty about μ as the number t of previous (imperfectly remembered) observations grows. Each panel corresponds to a particular value of K (maintaining the assumption that $\rho = 0$, as in Figure 1). Each panel shows the evolution for several different possible values of $\bar{\lambda}$ (color code is the same in both panels).

We thus obtain a nonlinear difference equation

$$\eta_{t+1} = \phi(\eta_t; \bar{\lambda})$$

for the dynamics of the scaled uncertainty measure. We can iterate this mapping, starting from the initial condition $\eta_0 = K/(K+1)$, to obtain the complete sequence of values $\{\eta_t\}$ for all $t \geq 0$ implied by any given value of $\bar{\lambda}$. This is the method used to compute the dynamic paths shown in Figure 1 in the main text.

Figure 1 shows the dynamics for $\{\eta_t\}$ implied by this solution, for various possible values of $\bar{\lambda}$, in the case that $K = 1$ and $\rho = 0$. Figure 9 shows how this graph would be different in the case of two larger values for K (but again assuming $\rho = 0$). A higher value of K (greater prior uncertainty) implies a higher value for the initial value η_0 of our normalized measure of uncertainty (since $\eta_0 = K/(K+1)$). This means that the curves all start higher, the larger the value of K . But the value of K also affects the long-run level of uncertainty η_∞ , even though the initial condition becomes irrelevant in the long run. Except when $\bar{\lambda} = 1$ (perfect memory), a higher value of K implies greater long-run uncertainty; and when K is large (as illustrated in the right panel), η_∞ is large (not much below the degree of uncertainty implied by the prior) except in the case of quite high values of $\bar{\lambda}$.

Figure 10 similarly shows how Figure 1 would look in the case of two larger values of ρ , but again assuming $K = 1$. We see that for a given degree of prior uncertainty and a given bound on memory precision, the rate at which uncertainty is reduced is slower when the external state is more serially correlated. This is because there are effectively fewer independent observations over a given number of periods when the state is serially correlated. In the case of perfect memory ($\bar{\lambda} = 1$), this affects the speed of learning but not the long-run value

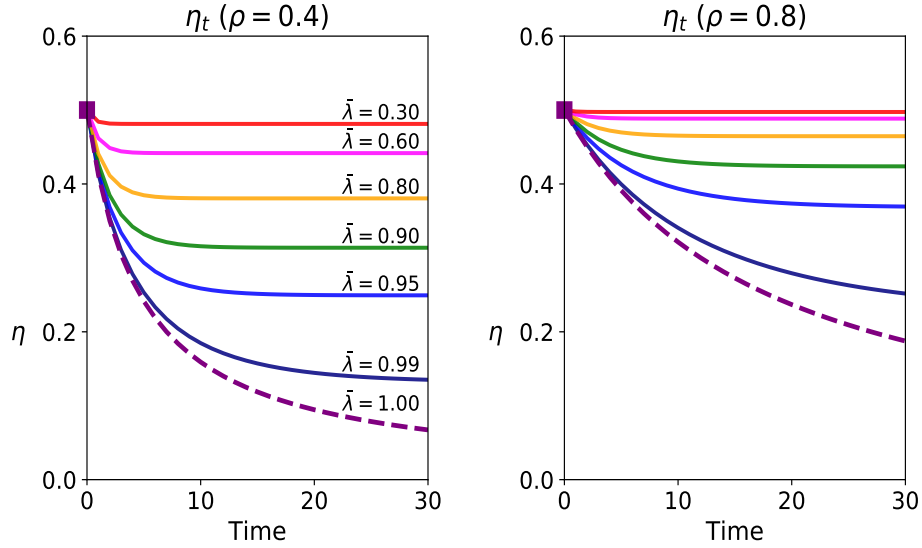


Figure 10: The evolution of scaled uncertainty about μ as the number t of previous (imperfectly remembered) observations grows. Each panel corresponds to a particular value of ρ (maintaining the assumption that $K = 1$, as in Figure 1). Each panel shows the evolution for several different possible values of $\bar{\lambda}$ (color code is the same in both panels).

$\eta_\infty = 0$ that is eventually reached. Instead, when memory is imperfect, the long-run value η_∞ is also higher when the state is more serially correlated; effectively, the limited number of recent observations of the state that can be retained in memory reveal less about the value of μ when the state is more serially correlated.

F.2 Solving for the value function $\tilde{V}(\eta)$ and policy function $\lambda^*(\eta)$ in the case of a linear information cost

In the case of a linear information cost (or any other cost function with a positive marginal cost of increasing I_t), it is necessary to solve the Bellman equation for the value function $\tilde{V}(\eta)$, in order to determine the optimal dynamics of $\{\lambda_t\}$. Here we explain the methods used to solve this problem in the case of a linear information cost (the results reported in section 3.2).

Once we have solved for the function $\phi(\eta_t; \lambda_t)$, as in the previous subsection, the Bellman equation for the case of a linear information cost can be written

$$\tilde{V}(\eta_t) = \min_{\lambda_t \in [0,1]} \left[\eta_t - \frac{\tilde{\theta}}{2} \log(1 - \lambda_t) + \beta \tilde{V}(\phi(\eta_t; \lambda_t)) \right]. \quad (\text{F.3})$$

We use the value function iteration algorithm to find the value function that is a fixed point of this mapping.

When iterating the mapping to update the value function, we use a grid search method to find the optimal policy function, because the right-hand side of the Bellman equation is in general a non-convex function of the policy variable λ_t (as we illustrate in Figure 13 below).

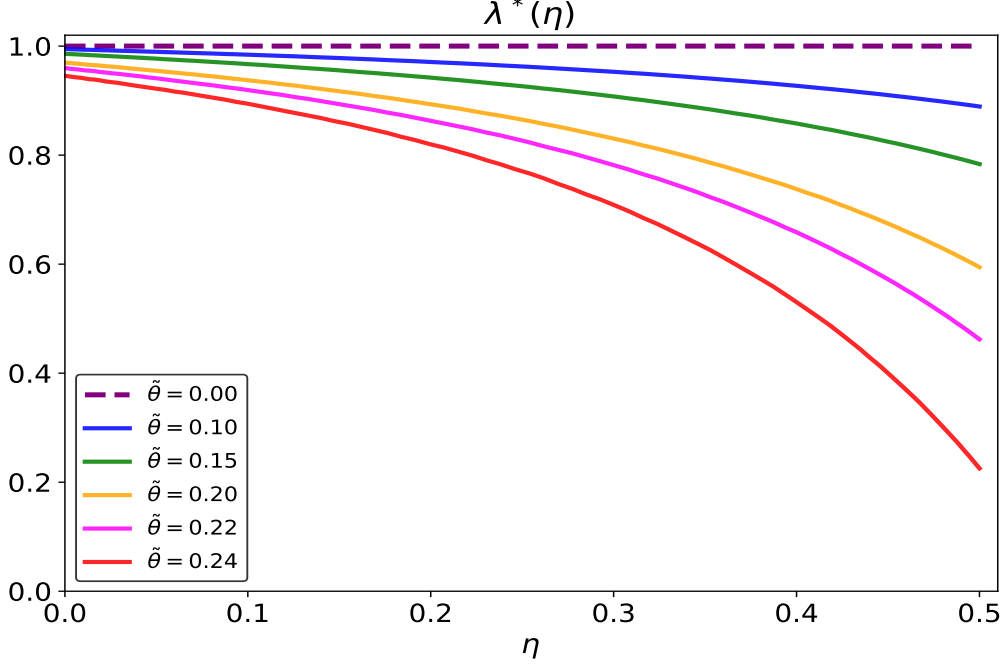


Figure 11: The optimal policy function $\lambda^*(\eta)$, in the case of progressively larger values for the information cost parameter $\tilde{\theta}$, under the assumption that $K = 1, \rho = 0$.

We approximate the value function with Chebyshev polynomials. Once the value function has converged, we can use our solution for $\tilde{V}(\eta)$ to solve numerically for the policy function $\lambda^*(\eta)$, the solution to the minimization problem on the right-hand side of (F.3).

This function is graphed for several values of $\tilde{\theta}$ in Figure 11, where we maintain the parameter values $K = 1, \rho = 0$ as in Figure 1. When $\tilde{\theta} = 0$ (no cost of memory precision), it is optimal to choose $\lambda_t = 1$ (perfect memory) in all cases. But for any value of η , the optimal $\lambda^*(\eta) < 1$ when $\tilde{\theta} > 0$ (since in this case, perfect memory becomes infinitely costly); furthermore it is lower (memory is more imperfect) the higher is $\tilde{\theta}$. We also see that for any cost parameter $\tilde{\theta} > 0$, the optimal $\lambda^*(\eta)$ is a decreasing function of η . This indicates that the less accurate the information contained in the cognitive state s_t (as indicated by the higher value of η_t), the less information about the cognitive state that it will be optimal to store in memory, when the memory cost can be reduced by storing a less informative record.

The policy function $\lambda_t = \lambda^*(\eta_t)$ together with the law of motion

$$\eta_{t+1} = \phi(\eta_t; \lambda_t) \quad (\text{F.4})$$

derived in section F.1 can then be solved for the dynamics of the scaled uncertainty $\{\eta_t\}$ for all $t \geq 0$, starting from the initial condition $\eta_0 = K/(K + 1)$. The dynamics implied by these equations can be graphed in a phase diagram, as illustrated in Figure 12. In the phase diagrams shown in each of the two panels, the value of η_t is indicated on the horizontal axis and the value of λ_t on the vertical axis. Equation (F.4), which holds regardless of the nature of the information cost function and the degree to which the future is discounted, determines a locus $\eta_\infty(\lambda)$, indicating for each value of λ the unique value of η that will be a fixed point of these dynamics if λ_t is held at the value λ . We can further show that whenever

Figure 4: The dynamics of scaled uncertainty and memory precision

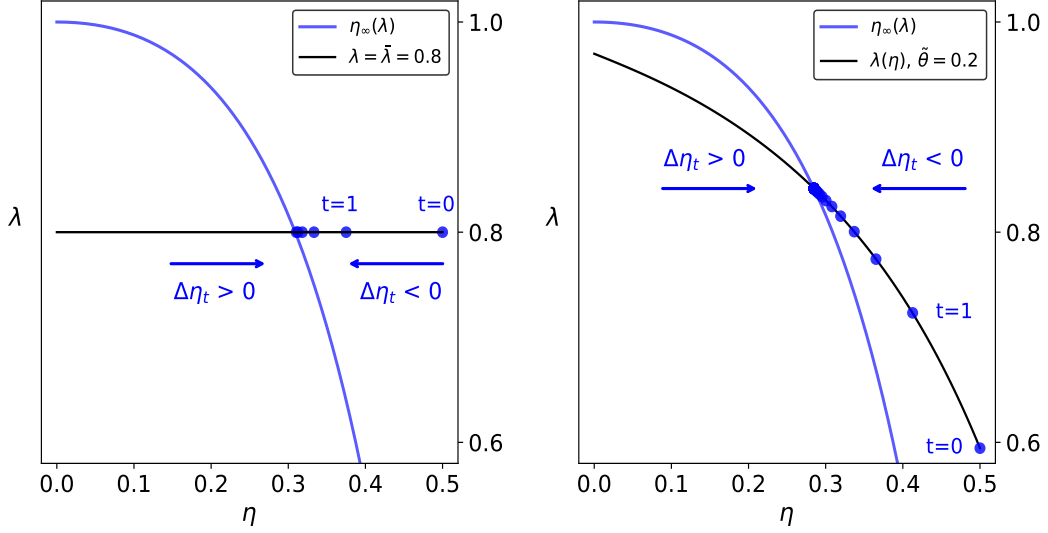


Figure 12: The dynamics of scaled uncertainty η_t and memory precision λ_t graphed in the phase plane. The left panel gives an alternative graphical presentation of the dynamics plotted in Figure 1 for the case of a fixed upper bound $\bar{\lambda}$ on memory precision. The right panel shows the corresponding dynamics in the case of a linear cost of precision parameterized by $\tilde{\theta}$.

$\eta_t < \eta_\infty(\lambda_t)$, the law of motion (F.4) implies that $\eta_{t+1} > \eta_t$, so that uncertainty will increase, while if $\eta_t > \eta_\infty(\lambda_t)$, it implies instead that $\eta_{t+1} < \eta_t$, so that uncertainty will decrease.

The choice of λ_t (and hence the degree to which uncertainty will increase or decrease) is given by the policy function, that depends on the specification of information costs. When there is a fixed upper bound on information (the case discussed in the previous subsection), the policy function is just a horizontal line at the vertical height $\bar{\lambda}$, as shown in the left panel of the figure.⁶⁰ In this case, the values of (η_t, λ_t) in successive periods start at the point $(\eta_0, \bar{\lambda})$, labeled “ $t = 0$ ” in the figure, and then move left along the graph of the policy function (since $\eta_0 > \eta_\infty(\bar{\lambda})$ as shown). They continue to move left along the policy function, with η_t converging asymptotically to $\eta_\infty(\bar{\lambda})$ from above; the stationary long-run values $(\eta_\infty, \lambda_\infty)$ correspond to the point at which the policy function $\lambda = \bar{\lambda}$ intersects the locus of fixed points $\eta_\infty(\lambda)$.

The right-hand panel of the figure shows the corresponding phase-plane dynamics in the less trivial case of a linear cost function for information. In this case, the policy function is instead a downward-sloping curve, as shown in Figure 11.⁶¹ Again the values of (η_t, λ_t) in successive periods must always lie on the graph of the policy function; the direction of

⁶⁰The figure plots the location of this line for the case $\bar{\lambda} = 0.8$. The figure is drawn for parameter values $K = 1, \rho = 0$. Thus the dynamics of uncertainty shown in the figure correspond to the curve labeled $\bar{\lambda} = 0.8$ in Figure 1.

⁶¹In the figure, the policy function and the implied dynamics are shown for the case in which $\tilde{\theta} = 0.2$, corresponding to one of the intermediate curves shown in Figure 11. Again the figure is for the case $K = 1, \rho = 0$, so that the location of the locus of fixed points $\eta_\infty(\lambda)$ and the law of motion (F.4) remain the same as in the left panel.

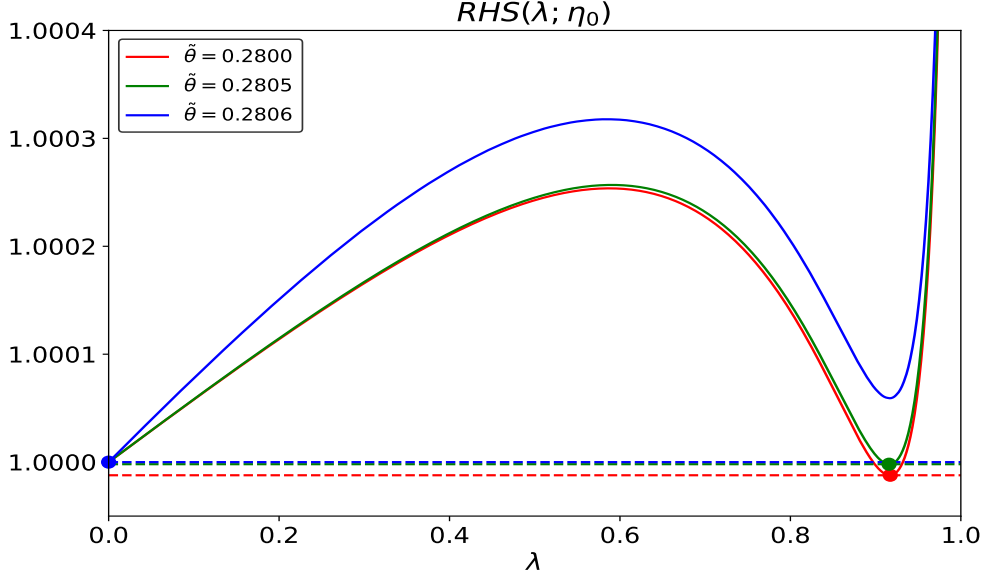


Figure 13: The objective function $RHS(\lambda_t, \eta_t)$ that is minimized in the Bellman equation, plotted as a function of λ_t for the initial level of uncertainty $\eta_t = \eta_0$. The function is normalized so that the value is 1.0 when $\lambda_t = 0$, and plotted for three nearby values of $\tilde{\theta}$, in the case that $K = 10$. The minimizing value of λ_t jumps discontinuously as $\tilde{\theta}$ passes a value between 0.2800 and 0.2805.

motion up or down this curve depends on whether the current position lies to the left or right of the locus of fixed points $\eta_\infty(\lambda)$. The initial point (labeled “ $t = 0$ ”) is determined as the point on the policy curve with horizontal coordinate given by the initial condition η_0 . Since this point lies to the right of the locus of fixed points, the points for successive periods move up and to the left on the policy curve, meaning that λ_t rises as η_t falls.

The scaled uncertainty continues to fall, and the precision of memory continues to rise, until the values (η_t, λ_t) converge to stationary long-run values $(\eta_\infty, \lambda_\infty)$, again corresponding to the point at which the policy function $\lambda^*(\eta)$ intersects the locus of fixed points $\eta_\infty(\lambda)$. Note that convergence is slower in the right panel of the figure than in the left, because in the early periods, when uncertainty is high, a less precise memory is chosen in the linear-cost case, resulting in slower learning from experience.

Different values of $\tilde{\theta}$ correspond to different locations for the policy function $\lambda^*(\eta)$, as shown in Figure 11, and hence to different dynamics in the phase plane, converging to different long-run levels of scaled uncertainty. The dynamics of scaled uncertainty as a function of the number of observations t are shown for progressively larger values of $\tilde{\theta}$ in Figure 3 in the main text, using the same format as in Figure 1.

F.3 The possibility of discontinuous solutions

Figure 13 illustrates our comment about the possible non-convexity of the optimization problem (F.3). Let $RHS(\lambda_t; \eta_t)$ be the function defined on the right-hand side of (F.3), i.e., the objective of the minimization problem. The figure plots the value of $RHS(\lambda; \eta_0)$, normalized by dividing by the positive constant $RHS(0; \eta_0)$ (so that a value of 1.0 on the

vertical axis means that $RHS(\lambda; \eta_0)$ is of exactly the same size as $RHS(0; \eta_0)$). This function is shown for each of three slightly different values of $\tilde{\theta}$, assuming in each case that $K = 10$, as in the right panel of Figure 5 in the text. In the case of each of these curves, a large dot (the same color as the curve) indicates the global minimum of the function. A horizontal dashed line (also the same color as the corresponding curve) indicates the minimum of $RHS(\lambda; \eta_0)$ — and thus the value of $\tilde{V}(\eta_0)$ — again normalized by dividing by $RHS(\eta_0)$.

The figure shows that for values of $\tilde{\theta}$ in this range, $RHS(\lambda)$ is not a convex function of λ . It is increasing for small enough values of λ , making the choice $\lambda_t = 0$ a local minimum in this case. (This is true for all values of $\tilde{\theta}$ greater than a critical value around 0.15, which explains the existence of the horizontal segment of the connected black curve in the right panel of Figure 5.) However, the function reaches a local maximum, and then decreases for larger values of λ , as the degree to which a larger value of λ_t reduces $\phi(\eta_0; \lambda_t)$ outweighs the increase in the information cost. (A large enough value of K is required for this to occur. A larger value of K increases the sensitivity of the value of $\phi(\eta_0; \lambda)$ to the value of λ ; see equation (F.5) below.) For even larger values of λ (values approaching 1), further increases in λ increase the information cost term so sharply that $RHS(\lambda; \eta_0)$ is again decreasing in λ . This means that there is a second local minimum of the objective function, at an interior value of λ . Which of the two local minima represents the global minimum of the function depends on parameter values.

In the case illustrated in the figure, the interior local minimum achieves a lower value of the objective than the choice $\lambda_t = 0$, for all values of $\tilde{\theta}$ less than a critical value that is slightly larger than 0.2805. (As shown in the figure, when $\tilde{\theta} = 0.2805$, the interior minimum achieves a value of the objective that is quite close to the value $RHS(0; \eta_0)$. However, the value achieved remains slightly smaller: there is a (barely visible) green dashed line, just below the blue dashed line at the normalized value 1.0.) But the normalized value of the objective at the interior minimum increases as $\tilde{\theta}$ is increased, and for a value of $\tilde{\theta}$ only slightly greater than 0.2805, the normalized value becomes greater than 1.0 (which is to say, the interior local minimum is no longer the global minimum of the objective). When this critical value of $\tilde{\theta}$ is passed, the optimal value $\lambda^*(\eta_0)$ jumps discontinuously from the interior local minimum (which is a continuously decreasing function of $\tilde{\theta}$) to the value zero. When this happens, the optimal long-run level for the normalized uncertainty measure η_∞ increases discontinuously, from a value on the lower branch of the correspondence shown in the right panel of Figure 5 to the value $\eta_0 = K/K + 1$. For all values of $\tilde{\theta}$ higher than this, it is optimal to choose a completely uninformative memory for all t , so that $\eta_t = \eta_0$ for all t , and hence $\eta_t \rightarrow \eta_\infty = \eta_0$.

For larger values of $\tilde{\theta}$ than those considered in Figure 11, the optimal policy function $\lambda^*(\eta)$ is equal to zero for all large enough (though still finite) values of η , as illustrated in Figure 14. Once $\tilde{\theta}$ is large enough for $\lambda^*(\eta_0)$ to equal zero, the optimal dynamics imply $\eta_t = \eta_0$ for all t , and hence $\eta_\infty = \eta_0 = K/K + 1$, as shown in Figure 5.

F.4 The case $\rho = 0$

Additional analytical results are possible in the case that $\rho = 0$ (the external state is an i.i.d. random variable). In this case, the law of motion for the scaled uncertainty measure

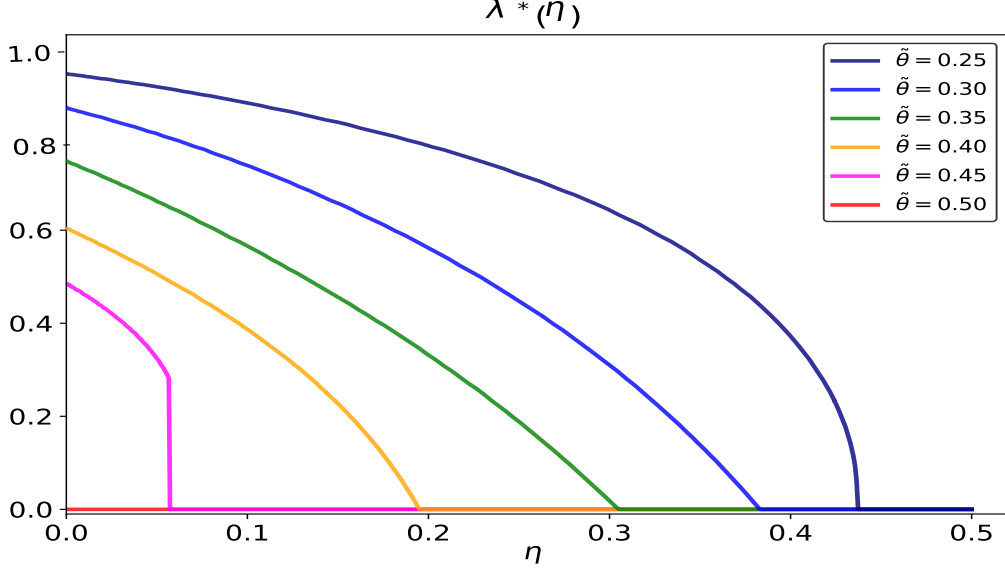


Figure 14: The optimal policy function $\lambda^*(\eta)$, in the case of progressively larger values for the information cost parameter $\tilde{\theta}$, under the assumption that $K = 1, \rho = 0$. Here we consider values of $\tilde{\theta}$ larger than those shown in Figure 11.

(derived in section F.1) simplifies to

$$\eta_{t+1} = 1 - \frac{1}{K + 1 - \lambda_t(K - \eta_t)} \equiv \phi(\eta_t; \lambda_t). \quad (\text{F.5})$$

In the case of an exogenous upper bound on mutual information, the nonlinear difference equation obtained by setting $\lambda_t = \bar{\lambda}$ in (F.5) is of an especially simple sort. The function on the right-hand side of this equation is a hyperbola, increasing and concave for all $\eta_t > 0$. We easily see that the right-hand side has a positive value when $\eta_t = 0$, and a value less than $K/(K + 1)$ when $\eta_t = K/(K + 1)$.

Thus for any $0 < \bar{\lambda} < 1$, the function $\phi(\eta_t; \bar{\lambda})$ is an increasing, concave function that is above the diagonal at $\eta_t = 0$ and below the diagonal at $\eta_t = K/(K + 1)$. It follows that the function must intersect the diagonal at exactly one point, $\eta_t = \eta_\infty$. We can furthermore give an explicit algebraic solution for this fixed point as the solution to a quadratic equation. Note in particular that it is necessarily strictly positive and strictly less than $K/(K + 1)$, and that it is a decreasing function of $\bar{\lambda}$, approaching $K/(K + 1)$ as $\bar{\lambda} \rightarrow 0$, and approaching 0 as $\bar{\lambda} \rightarrow 1$.

On the interval $\eta_\infty < \eta_t \leq K/(K + 1)$, the law of motion (F.5) implies that $\eta_\infty < \eta_{t+1} < \eta_t$. Hence when we start from the initial condition $\eta_0 = K/(K + 1)$, the implied dynamics must satisfy

$$\eta_0 > \eta_1 > \eta_2 > \eta_3 \dots,$$

a monotonically decreasing sequence. Because the sequence is bounded below by η_∞ , it must converge, and it is easily seen that it can only converge to the fixed point η_∞ that we have already calculated. Hence for each possible $\bar{\lambda}$, we obtain a monotonically decreasing, convergent sequence of the kind shown in Figure 1. We can also easily show that the curve must be lower for each value of t , the larger is $\bar{\lambda}$.

We can also obtain additional analytical results in the case of a linear information cost. The value function satisfies a Bellman equation of the form

$$\tilde{V}(\eta_t) = \min_{\lambda_t} \left[\beta^2 \eta_t - \frac{\tilde{\theta}}{2} \log(1 - \lambda) + \beta \tilde{V}(\phi(\eta_t; \lambda_t)) \right].$$

The first order condition with respect to λ_t is

$$\frac{\tilde{\theta}}{2} \frac{1}{1 - \lambda_t} + \beta \tilde{V}'(\eta_{t+1}) \frac{\partial \phi(\eta_t; \lambda_t)}{\partial \lambda_t} = 0. \quad (\text{F.6})$$

And the envelope condition is

$$\tilde{V}'(\eta_t) = \beta^2 + \beta \tilde{V}'(\eta_{t+1}) \frac{\partial \phi(\eta_t; \lambda_t)}{\partial \eta_t}.$$

We can use these two conditions to derive an Euler equation for the dynamics of the scaled uncertainty measure.

Substituting the solution (F.5) for $\phi(\eta_t; \lambda_t)$ and taking the derivative with respect to λ_t , we can rewrite (F.6) as

$$\begin{aligned} \tilde{V}'(\eta_{t+1}) &= -\frac{\tilde{\theta}}{2\beta} \frac{1}{1 - \lambda_t} \left(\frac{\partial \phi(\eta_t; \lambda_t)}{\partial \lambda_t} \right)^{-1} \\ &= -\frac{\tilde{\theta}}{2\beta} \frac{1}{1 - \lambda_t} \left(-\frac{(K - \eta_t)}{(K + 1 - \lambda_t(K - \eta_t))^2} \right)^{-1} \\ &= \frac{\tilde{\theta}}{2\beta} \frac{(K + 1 - \lambda_t(K - \eta_t))^2}{(1 - \lambda_t)(K - \eta_t)} \\ &= \frac{\tilde{\theta}}{2\beta} \frac{1}{(1 - \eta_{t+1})(1 - (1 - \eta_{t+1})(1 + \eta_t))}, \end{aligned}$$

where the last equality is derived by again substituting the law of motion (F.5). It follows that if $\eta_t \rightarrow \eta_\infty$ in the long run, the stationary solution η_∞ must satisfy

$$\tilde{V}'(\eta_\infty) = \frac{\tilde{\theta}}{2\beta} \frac{1}{(1 - \eta_\infty)\eta_\infty^2}. \quad (\text{F.7})$$

Next we rewrite (F.4), again taking the derivative of expression (F.5) for $\tilde{\eta}_t; \lambda_t$:

$$\begin{aligned} \tilde{V}'(\eta_t) &= \beta^2 + \beta \tilde{V}'(\eta_{t+1}) \frac{\partial \phi(\eta_t; \lambda_t)}{\partial \eta_t} \\ &= \beta^2 + \beta \tilde{V}'(\eta_{t+1}) \frac{\lambda_t}{(K + 1 - \lambda(K - \eta_t))^2} \\ &= \beta^2 + \beta \tilde{V}'(\eta_{t+1}) \frac{\lambda_t}{(1 - \eta_{t+1})^{-2}} \\ &= \beta^2 + \beta \tilde{V}'(\eta_{t+1}) (1 - \eta_{t+1})^2 \frac{(K + 1)(1 - \eta_{t+1}) - 1}{(K - \eta_t)(1 - \eta_{t+1})}. \end{aligned}$$

It follows that the stationary solution η_∞ must satisfy

$$\tilde{V}'(\eta_\infty) = \beta^2 + \beta \tilde{V}'(\eta_\infty) \frac{(1 - \eta_\infty)[(K + 1)(1 - \eta_\infty) - 1]}{K - \eta_\infty}. \quad (\text{F.8})$$

Moreover, in a stationary solution, the value $\tilde{V}'(\eta_\infty)$ given by (F.7) must also be the value of $\tilde{V}'(\eta_\infty)$ in (F.8). Using (F.7) to substitute for $\tilde{V}'(\eta_\infty)$ in (F.8), we obtain a condition that must be satisfied by η_∞ in any stationary solution with an interior optimum (i.e., a stationary solution in which $0 < \eta_\infty < K/(K + 1)$):

$$\tilde{\theta} = 2\beta^3(1 - \eta_\infty)\eta_\infty^2 \left[1 - \beta \frac{(K + 1)(1 - \eta_\infty)^2 - (1 - \eta_\infty)}{K - \eta_\infty} \right]^{-1}. \quad (\text{F.9})$$

This is the relationship between $\tilde{\theta}$ and η_∞ that is graphed as a connected black curve in Figure 5. Note that for any value $0 < \eta_\infty < K/(K + 1)$, there is a unique $\tilde{\theta} > 0$ consistent with this relationship; but (as shown in the right panel of Figure 5) there may be multiple solutions for η_∞ consistent with a given value of $\tilde{\theta}$.

G Predicted Values for the Quantitative Measures of Forecast Bias

Here we provide further explanation of the numerical results reported in section 4 of the main text.

G.1 Long-run stationary fluctuations

From the definition of the univariate memory state $\tilde{m}_{t+1} = \lambda_t v'_t \bar{s}_t + \omega_{t+1}$, we can derive a law of motion for the univariate memory state $\tilde{m}_t$. Using the subscript ∞ for the long-run stationary coefficients, we get

$$\begin{aligned} \tilde{m}_{t+1} &= \lambda_\infty v'_\infty \bar{s}_t + \tilde{\omega}_{t+1} \\ &= \lambda_\infty v'_\infty \begin{pmatrix} \hat{\mu}_t \\ y_t \end{pmatrix} + \tilde{\omega}_{t+1} \\ &= \lambda_\infty [e'_1 v_\infty \{(e'_1 - \gamma_1 c')m_t + \gamma_1 y_t\} + (e'_2 v_\infty)y_t] + \tilde{\omega}_{t+1} \\ &= \lambda_\infty [e'_1 v_\infty \{(e'_1 - \gamma_1 c')X_\infty v_\infty \tilde{m}_t + \gamma_1 y_t\} + (e'_2 v_\infty)y_t] + \tilde{\omega}_{t+1} \\ &= \rho_m \tilde{m}_t + \rho_{my} y_t + \tilde{\omega}_{t+1} \end{aligned}$$

where $\rho_m \equiv \lambda_\infty (e'_1 v_\infty) (e'_1 - \gamma_1 c') X_\infty v_\infty$ and $\rho_{my} \equiv \lambda_\infty (\gamma_1 + e'_2 v_\infty)$.

We can evaluate the numerical values of the coefficients defining the long-run dynamics as follows. Equations (F.1)–(F.2) imply that the long-run coefficients $\lambda_\infty, \eta_\infty, \gamma_{1,\infty}$ must satisfy the pair of nonlinear equations

$$\eta_\infty = \frac{(1 - \lambda_\infty)(1 - \gamma_{1,\infty})^2 K + (1 - \rho^2 \lambda_\infty) \gamma_{1,\infty}^2}{1 - \lambda_\infty (1 - (1 - \rho) \gamma_{1,\infty})^2},$$

$$\gamma_{1,\infty} = \frac{(1 - \lambda_\infty)K + (1 - \rho)\lambda_\infty\eta_\infty}{(1 - \lambda_\infty)(K + \rho^2) + (1 - \rho^2) + (1 - \rho)^2\lambda_\infty\eta_\infty}.$$

In the case of an exogenous bound on mutual information, we can set $\lambda_\infty = \bar{\lambda}$, in which case these provide two equations to solve for the values of η_∞ and $\gamma_{1,\infty}$. (Note that the relevant solution is the one that satisfies the bounds $0 < \eta_\infty < K/(K + 1)$, and that it necessarily also satisfies $0 < \gamma_{1,\infty} < 1/(1 - \rho)$.) This allows us to compute the long-run stationary values of the coefficients η and γ_1 plotted for alternative values of $\bar{\lambda}$ in Figure 2.

We have also shown in section E.3 that the optimal weight vector v_t is just a normalized version of the vector $\delta_{t+1} \equiv e_1 - \gamma_{1,t+1}c$. Hence in the long run, this vector must become

$$v_\infty = \frac{e_1 - \gamma_{1,\infty}c}{(e'_1 - \gamma_{1,\infty}c')X_\infty(e_1 - \gamma_{1,\infty}c)}.$$

In particular, the ratio $v_{2,\infty}/v_{1,\infty}$ (the quantity plotted as “ v_∞ ” in Figure 2) is given by

$$\frac{v_{2,\infty}}{v_{1,\infty}} = -\frac{\rho\gamma_{1,\infty}}{1 - (1 - \rho)\gamma_{1,\infty}} < 0.$$

Finally, we observe that the intrinsic persistence coefficient ρ_m defined above must satisfy

$$\begin{aligned} \rho_m &\equiv \lambda_\infty v_{1,\infty} \cdot (e'_1 - \gamma_{1,\infty}c')X_\infty v_\infty \\ &= \lambda_\infty v_{1,\infty} \\ &= \lambda_\infty(1 - (1 - \rho)\gamma_{1,\infty}). \end{aligned}$$

This allows us to calculate the other coefficient that is plotted in Figure 2. Note that because the Kalman gain necessarily satisfies the bounds $0 < \gamma_1 < 1/(1 - \rho)$, this solution for the intrinsic persistence coefficient implies that

$$0 < \rho_m < 1. \tag{G.1}$$

In the long run, we can describe the evolution of the DM’s cognitive state using the following system of equations:

$$\begin{aligned} \tilde{m}_{t+1} &= \rho_m \tilde{m}_t + \rho_{my} y_t + \tilde{\omega}_{t+1} \\ y_{t+1} &= (1 - \rho)\mu + \rho y_t + \epsilon_{y,t+1} \end{aligned}$$

Therefore, we can write it as a VAR(1) system with constant coefficients and Gaussian innovation terms:

$$\begin{pmatrix} \tilde{m}_{t+1} \\ y_{t+1} \end{pmatrix} = \begin{pmatrix} 0 \\ 1 - \rho \end{pmatrix} \mu + \begin{pmatrix} \rho_m & \rho_{my} \\ 0 & \rho \end{pmatrix} \begin{pmatrix} \tilde{m}_t \\ y_t \end{pmatrix} + \begin{pmatrix} \tilde{\omega}_{t+1} \\ \epsilon_{y,t+1} \end{pmatrix}$$

Because the two eigenvalues of this vector law of motion are ρ and ρ_m , (G.1) implies that this describes a stationary stochastic process. Hence we can compute stationary long-run values for the second moments of the variables, and use these to define the impulse response functions and predicted regression coefficients reported in the text.

For example, in the case of a fixed per-period bound on mutual information, we can compute the impulse responses for the DM’s estimate of μ and her one-quarter-ahead forecast

of the external state, as explained in section 3.3. Here we present additional figures, showing what the impulse responses shown in Figure 6 in the text would be like in the case of alternative values of ρ . In Figures 15 and 16 shown here, each panel corresponds to a different value of ρ , and shows the responses for several different possible values of $\bar{\lambda}$. (As with Figure 6 in the main text, we here assume that $K = 1$.)

Figure 15: Impulse responses of the DM's estimate of μ for alternative degrees of persistence ρ of the external state process.

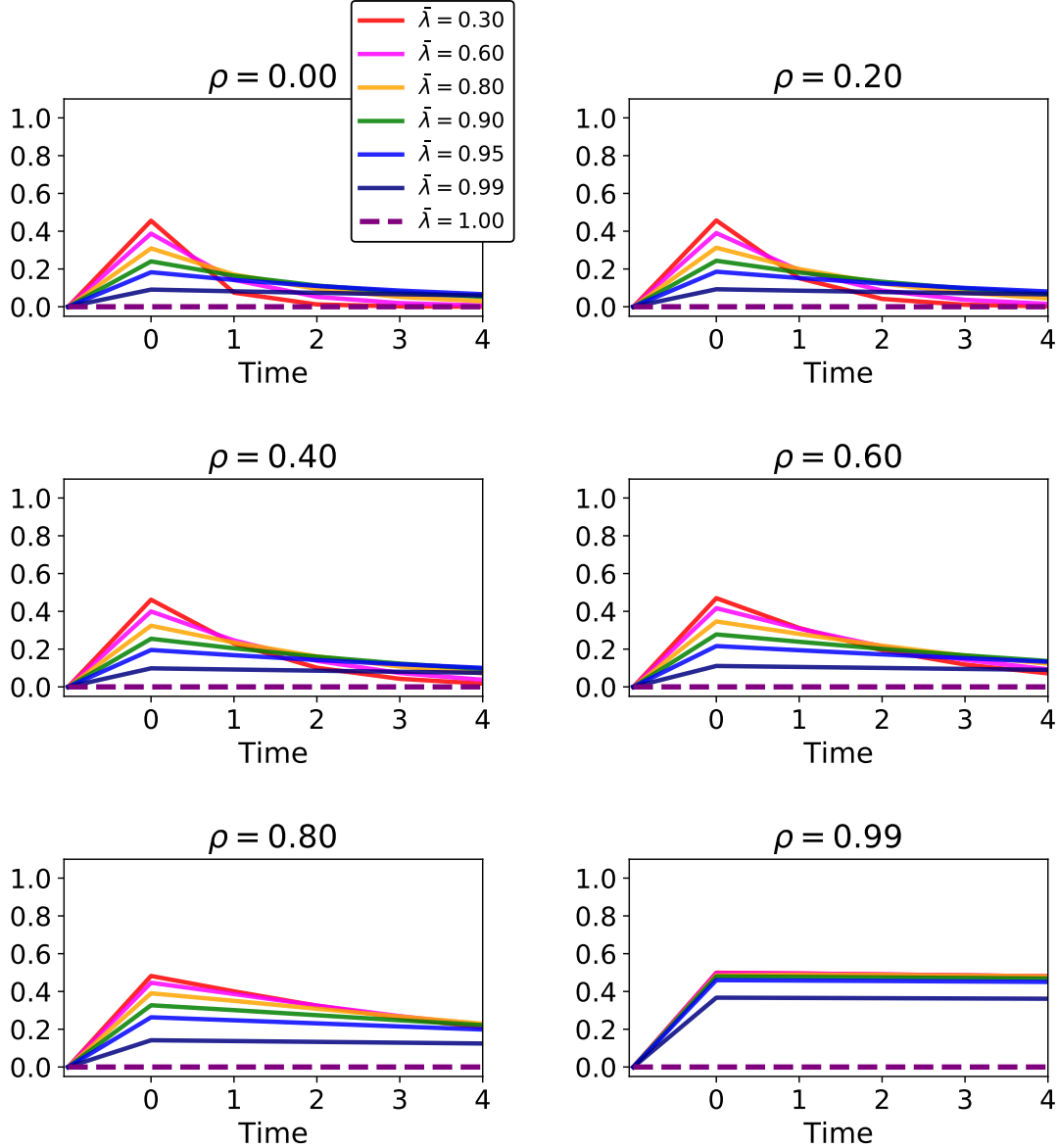
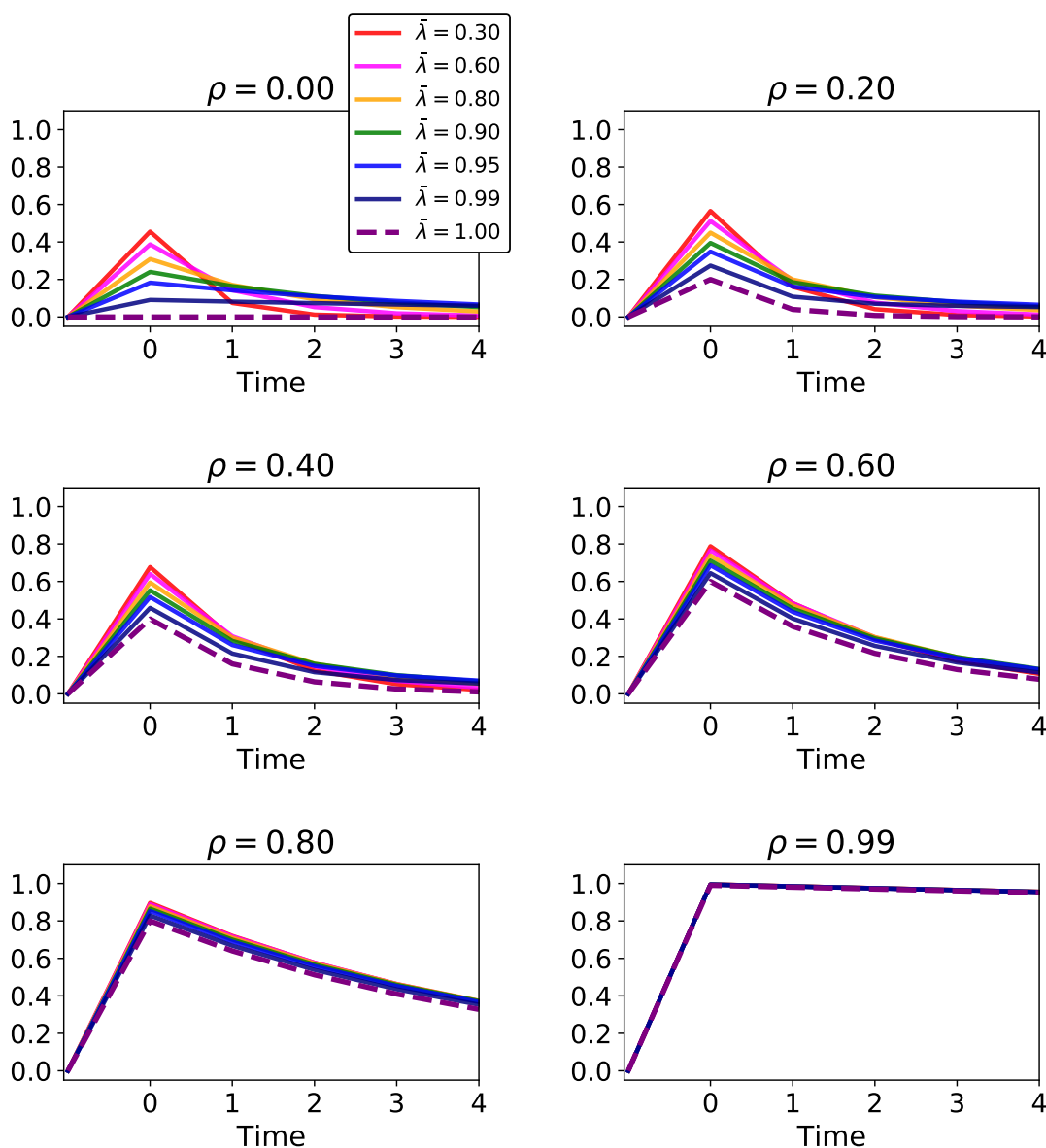


Figure 16: Impulse responses of the DM's one-quarter-ahead forecast of the external state for alternative degrees of persistence ρ of the external state process.



G.2 Predicted value of the regression coefficient ρ_h^{subj}

Given a long enough series of observations from an environment with a fixed μ , our model yields stationary values for the Kalman gain γ_1 and for the amplitude of fluctuations in the memory state $var[\bar{m}_t]$. We can then compute the values of the following long-run conditional second moments:

$$\begin{aligned} var[\bar{m}_t|\mu] &= var[\bar{m}_t] - cov[\bar{m}_t, \mu]var[\mu]^{-1}cov[\mu, \bar{m}_t] \\ &= var[\bar{m}_t] - cov[\bar{m}_t, x_t]e_1var[\mu]^{-1}e_1'cov[x_t, \bar{m}_t] \\ &= var[\bar{m}_t] - \frac{1}{var[\mu]}var[\bar{m}_t]e_1e_1'var[\mu] \end{aligned}$$

$$\begin{aligned} cov[\hat{\mu}_t, y_t|\mu] &= cov[(e_1' - \gamma_1 c')\bar{m}_t + \gamma_1 y_t, y_t|\mu] \\ &= (e_1' - \gamma_1 c')cov[\bar{m}_t, y_t|\mu] + \gamma_1 var[y_t|\mu] \\ &= (e_1' - \gamma_1 c')var[\bar{m}_t|\mu]c + \gamma_1 var[y_t|\mu] \end{aligned}$$

$$\begin{aligned} var[\hat{\mu}_t|\mu] &= var[(e_1' - \gamma_1 c')\bar{m}_t + \gamma_1 y_t|\mu] \\ &= (e_1' - \gamma_1 c')var[\bar{m}_t|\mu](e_1 - \gamma_1 c) + \gamma_1^2 var[y_t|\mu] + 2\gamma_1(e_1' - \gamma_1 c')cov[\bar{m}_t, y_t|\mu] \\ &= (e_1' - \gamma_1 c')var[\bar{m}_t|\mu](e_1 - \gamma_1 c) + \gamma_1^2 var[y_t|\mu] + 2\gamma_1(e_1' - \gamma_1 c')var[\bar{m}_t|\mu]c \end{aligned}$$

In order to write the dynamics of the model in terms of scale-invariant quantities, we divide each second moment by $var[y_t|\mu] = \sigma_y^2$. Thus we can write

$$\begin{aligned} \frac{var[\bar{m}_t|\mu]}{var[y_t|\mu]} &= \tilde{\Sigma}_{\bar{m}} - \frac{1}{K}\tilde{\Sigma}_{\bar{m}}e_1e_1'\tilde{\Sigma}_{\bar{m}} \\ \frac{cov[\hat{\mu}_t, y_t|\mu]}{var[y_t|\mu]} &= (e_1' - \gamma_1 c')\frac{var[\bar{m}_t|\mu]}{var[y_t|\mu]}c + \gamma_1 \\ \frac{var[\hat{\mu}_t|\mu]}{var[y_t|\mu]} &= (e_1' - \gamma_1 c')\frac{var[\bar{m}_t|\mu]}{var[y_t|\mu]}(e_1 - \gamma_1 c) + \gamma_1^2 + 2\gamma_1(e_1' - \gamma_1 c')\frac{var[\bar{m}_t|\mu]}{var[y_t|\mu]}c, \end{aligned}$$

using the notation $\tilde{\Sigma}_{\bar{m}} \equiv var[\bar{m}_t]/\sigma_y^2$.

We now wish to calculate the predicted asymptotic value of the regression coefficient

$$\rho_h^{subj} \equiv \frac{cov[\hat{y}_{t+h|t}, y_t|\mu]}{var[y_t|\mu]}$$

where $\hat{y}_{t+h|t} \equiv E[y_{t+h}|\bar{m}_t, y_t]$. From

$$\begin{aligned} cov[\hat{y}_{t+h|t}, y_t|\mu] &= cov[(1 - \rho^h)\hat{\mu}_t + \rho^h y_t, y_t|\mu] \\ &= (1 - \rho^h)cov[\hat{\mu}_t, y_t|\mu] + \rho^h var[y_t|\mu], \end{aligned}$$

where $\hat{\mu}_t \equiv E[\mu|\bar{m}_t, y_t]$, we can then compute

$$\begin{aligned} \rho_h^{subj} &= (1 - \rho^h)\frac{cov[\hat{\mu}_t, y_t|\mu]}{var[y_t|\mu]} + \rho^h \\ &= (1 - \rho^h) \left[(e_1' - \gamma_1 c') \left(\tilde{\Sigma}_{\bar{m}} - \frac{1}{K}\tilde{\Sigma}_{\bar{m}}e_1e_1'\tilde{\Sigma}_{\bar{m}} \right) c + \gamma_1 \right] + \rho^h. \end{aligned}$$

These are the coefficients whose values are plotted against the value of $\rho_h = \rho^h$ in Figure 7.

G.3 Predicted value of the Coibion-Gorodnichenko regression coefficient b

We wish to compute the predicted asymptotic value of the regression coefficient $b \equiv \frac{\text{cov}[FE_t, FR_t]}{\text{var}[FR_t]}$, where the forecast error $FE_t \equiv y_{t+1} - \hat{y}_{t+1|t}$ and the forecast revision $FR_t \equiv \hat{y}_{t+1|t} - \hat{y}_{t+1|t-1}$.

First, from $\hat{y}_{t+h|t} = (1 - \rho^h)\hat{\mu}_t + \rho^h y_t$ and $\hat{y}_{t+h|t-1} = (1 - \rho^{h+1})\hat{\mu}_{t-1} + \rho^{h+1}y_{t-1}$, we get

$$\begin{aligned} FE_t &= \left((1 - \rho^h)\mu + \rho^h y_t + \sum_{j=1}^h \rho^{h-j} \epsilon_{t+j} \right) - ((1 - \rho^h)\hat{\mu}_t + \rho^h y_t) \\ &= (1 - \rho^h)(\mu - \hat{\mu}_t) + \sum_{j=1}^h \rho^{h-j} \epsilon_{y,t+j} \\ FR_t &= ((1 - \rho^h)\hat{\mu}_t + \rho^h y_t) - ((1 - \rho^{h+1})\hat{\mu}_{t-1} + \rho^{h+1}y_{t-1}) \\ &= (1 - \rho^h)\hat{\mu}_t - (1 - \rho^{h+1})\hat{\mu}_{t-1} + \rho^h((1 - \rho)\mu + \epsilon_{y,t}) \end{aligned}$$

Then,

$$\begin{aligned} \text{cov}[FE_t, FR_t|\mu] &= \text{cov}[-(1 - \rho^h)\hat{\mu}_t, (1 - \rho^h)\hat{\mu}_t - (1 - \rho^{h+1})\hat{\mu}_{t-1}|\mu] \\ &= -(1 - \rho^h)^2 \text{var}[\hat{\mu}_t|\mu] + (1 - \rho^h)(1 - \rho^{h+1})\text{cov}[\hat{\mu}_t, \hat{\mu}_{t-1}|\mu] \\ \text{var}[FR_t|\mu] &= \text{var}[(1 - \rho^h)\hat{\mu}_t - (1 - \rho^{h+1})\hat{\mu}_{t-1} + \rho^h \epsilon_{y,t}|\mu] \\ &= \{(1 - \rho^h)^2 + (1 - \rho^{h+1})^2\} \text{var}[\hat{\mu}_t|\mu] - 2(1 - \rho^h)(1 - \rho^{h+1})\text{cov}[\hat{\mu}_t, \hat{\mu}_{t-1}|\mu] \\ &\quad + \rho^{2h} \text{var}[\epsilon_{y,t}|\mu] + 2\rho^h(1 - \rho^h)\text{cov}[\hat{\mu}_t, \epsilon_{y,t}|\mu] \\ &= \{(1 - \rho^h)^2 + (1 - \rho^{h+1})^2\} \text{var}[\hat{\mu}_t|\mu] - 2(1 - \rho^h)(1 - \rho^{h+1})\text{cov}[\hat{\mu}_t, \hat{\mu}_{t-1}|\mu] \\ &\quad + (\rho^{2h} + 2\rho^h(1 - \rho^h)\gamma_1) \sigma_\epsilon^2 \end{aligned}$$

where we have substituted $\text{cov}[\hat{\mu}_t, \epsilon_{y,t}|\mu] = \text{cov}[\gamma_1 \epsilon_{y,t}, \epsilon_{y,t}|\mu] = \gamma_1 \sigma_\epsilon^2$, to obtain the last equality.

It then follows that

$$\frac{\text{cov}[FE_t, FR_t|\mu]}{\text{var}[y_t|\mu]} = -(1 - \rho^h)^2 \frac{\text{var}[\hat{\mu}_t|\mu]}{\text{var}[y_t|\mu]} + (1 - \rho^h)(1 - \rho^{h+1}) \frac{\text{cov}[\hat{\mu}_t, \hat{\mu}_{t-1}|\mu]}{\text{var}[y_t|\mu]} \quad (\text{G.2})$$

$$\begin{aligned} \frac{\text{var}[FR_t|\mu]}{\text{var}[y_t|\mu]} &= \{(1 - \rho^h)^2 + (1 - \rho^{h+1})^2\} \frac{\text{var}[\hat{\mu}_t|\mu]}{\text{var}[y_t|\mu]} - 2(1 - \rho^h)(1 - \rho^{h+1}) \frac{\text{cov}[\hat{\mu}_t, \hat{\mu}_{t-1}|\mu]}{\text{var}[y_t|\mu]} \\ &\quad + (\rho^{2h} + 2\rho^h(1 - \rho^h)\gamma_1) (1 - \rho^2) \end{aligned} \quad (\text{G.3})$$

since $\text{var}[y_t|\mu] = \frac{1}{1-\rho^2} \sigma_\epsilon^2$. Using these expressions, we can compute the Coibion-Gorodnichenko regression coefficient

$$b = \frac{\text{cov}[FE_t, FR_t|\mu]}{\text{var}[y_t|\mu]} \left(\frac{\text{var}[FR_t|\mu]}{\text{var}[y_t|\mu]} \right)^{-1}. \quad (\text{G.4})$$

Note that when $\rho = 0$, this implies that $b = -\frac{1}{2}$ regardless of the horizon h , since

$$\begin{aligned}\frac{\text{cov}[FE_t, FR_t|\mu]}{\text{var}[y_t|\mu]} &= -\frac{\text{var}[\hat{\mu}_t|\mu]}{\text{var}[y_t|\mu]} + \frac{\text{cov}[\hat{\mu}_t, \hat{\mu}_{t-1}|\mu]}{\text{var}[y_t|\mu]} \\ \frac{\text{var}[FR_t|\mu]}{\text{var}[y_t|\mu]} &= 2\frac{\text{var}[\hat{\mu}_t|\mu]}{\text{var}[y_t|\mu]} - 2\frac{\text{cov}[\hat{\mu}_t, \hat{\mu}_{t-1}|\mu]}{\text{var}[y_t|\mu]} = -2\left(\frac{\text{cov}[FE_t, FR_t|\mu]}{\text{var}[y_t|\mu]}\right).\end{aligned}$$

Finally, to evaluate the expressions needed in order to compute the predicted value of b , in general we need to be able to compute $\text{cov}[\hat{\mu}_t, \hat{\mu}_{t-1}|\mu]$. For this, we need to solve for the long-run stationary fluctuations in $\hat{\mu}_t$. Using the notation X_∞ and v_∞ for the long-run stationary values of X_t and v_t , we can derive the long-run dynamics of $\hat{\mu}_t$ as follows.

$$\begin{aligned}\hat{\mu}_t &= (e'_1 - \gamma_1 c')\bar{m}_t + \gamma_1 y_t \\ &= (e'_1 - \gamma_1 c')\left((\lambda_\infty X_\infty v_\infty v'_\infty)\begin{pmatrix} \hat{\mu}_{t-1} \\ y_{t-1} \end{pmatrix} + \bar{\omega}_t\right) + \gamma_1 y_t \\ &= (e'_1 - \gamma_1 c')(\lambda_\infty X_\infty v_\infty (e'_1 v_\infty \hat{\mu}_{t-1} + e'_2 v_\infty y_{t-1}) + \bar{\omega}_t) + \gamma_1 (\rho y_{t-1} + \epsilon_t) \\ &= \rho_{\hat{\mu}} \hat{\mu}_{t-1} + \rho_{\hat{\mu}, y} y_{t-1} + \epsilon_{\hat{\mu}t},\end{aligned}$$

where $\rho_{\hat{\mu}} \equiv \lambda_\infty e'_1 v_\infty (e'_1 - \gamma_1 c') X_\infty v_\infty$, $\rho_{\hat{\mu}, y} \equiv \lambda_\infty e'_2 v_\infty (e'_1 - \gamma_1 c') X_\infty v_\infty + \gamma_1 \rho$, and $\epsilon_{\hat{\mu}t} = (e'_1 - \gamma_1 c')\bar{\omega}_t + \gamma_1 \epsilon_{y,t}$. From this, we can derive

$$\begin{aligned}\text{cov}[\hat{\mu}_t, \hat{\mu}_{t-1}|\mu] &= \text{cov}[\rho_{\hat{\mu}} \hat{\mu}_{t-1} + \rho_{\hat{\mu}, y} y_{t-1}, \hat{\mu}_{t-1}|\mu] \\ &= \rho_{\hat{\mu}} \text{var}[\hat{\mu}_t|\mu] + \rho_{\hat{\mu}, y} \text{cov}[\hat{\mu}_t, y_t|\mu].\end{aligned}$$

Using this result, together with the solution for $\text{var}[\hat{\mu}|\mu]$ derived in Appendix G.2, we can then evaluate both of the covariances (G.2) and (G.3), and use them to compute the predicted regression coefficient (G.4). These are the formulas used to compute the values shown in the left panel of Figure 8.

G.4 Predicted value of the Fuhrer regression coefficient γ

We wish to compute the value of

$$\gamma \equiv \frac{\text{cov}[FR_t^i, \Delta_{t-1}^i]}{\text{var}[\Delta_{t-1}^i]}, \quad (\text{G.5})$$

where FR_t^i is the forecast revision of an individual forecaster i ,

$$FR_t^i \equiv \hat{y}_{t+h|t}^i - \hat{y}_{t+h|t-1}^i,$$

and Δ_{t-1}^i is the difference between i 's forecast and the consensus forecast in period $t-1$. Individual forecasts are given by

$$\hat{y}_{t+h|t}^i = (1 - \rho^h)\hat{\mu}_t^i + \rho^h y_t,$$

where $\hat{\mu}_t^i$ is i 's estimate of μ at time t , given by

$$\hat{\mu}_t^i = \xi \tilde{m}_t^i + \gamma_1 y_t,$$

using (3.9).

We have shown in (4.6) that

$$\Delta_t^i = (1 - \rho^h) \xi \tilde{\omega}_t^{sum,i},$$

and correspondingly that

$$\Delta_{t-1}^i = (1 - \rho^{h+1}) \xi \tilde{\omega}_{t-1}^{sum,i},$$

where $\tilde{\omega}_t^{sum,i}$ is defined in (3.8). It follows that the denominator of (G.5) is equal to

$$\text{var}[\Delta_{t-1}^i] = (1 - \rho^{h+1})^2 \xi^2 \cdot \text{var}[\tilde{\omega}_{t-1}^{sum,i}].$$

Furthermore, the numerator of (3.8) depends only on the part of FR_t^i that is correlated with $\tilde{\omega}_{t-1}^{sum,i}$. It follows from (3.7) that

$$\tilde{m}_t^i = \tilde{m}_t^{avg} + \tilde{\omega}_{t+1}^{sum,i},$$

and hence that

$$\begin{aligned} \hat{y}_{t+h|t}^i &= \hat{y}_{t+h|t}^{avg} + (1 - \rho^h) [\hat{\mu}_t^i - \hat{\mu}_t^{avg}] \\ &= \hat{y}_{t+h|t}^{avg} + (1 - \rho^h) \xi [\tilde{m}_t^i - \tilde{m}_t^{avg}] \\ &= \hat{y}_{t+h|t}^{avg} + (1 - \rho^h) \xi \tilde{\omega}_t^{sum,i} \\ &= \hat{y}_{t+h|t}^{avg} + (1 - \rho^h) \xi \tilde{\omega}_t^i + (1 - \rho^h) \xi \rho_m \tilde{\omega}_{t-1}^{sum,i}. \end{aligned}$$

Similarly,

$$\begin{aligned} \hat{y}_{t+h|t-1}^i &= \hat{y}_{t+h|t-1}^{avg} + (1 - \rho^{h+1}) [\hat{\mu}_{t-1}^i - \hat{\mu}_{t-1}^{avg}] \\ &= \hat{y}_{t+h|t-1}^{avg} + (1 - \rho^{h+1}) \xi [\tilde{m}_{t-1}^i - \tilde{m}_{t-1}^{avg}] \\ &= \hat{y}_{t+h|t-1}^{avg} + (1 - \rho^{h+1}) \xi \tilde{\omega}_{t-1}^{sum,i}, \end{aligned}$$

so that

$$\begin{aligned} \text{cov}[FR_t^i, \Delta_{t-1}^i] &= \text{cov}[\hat{y}_{t+h|t}^i - \hat{y}_{t+h|t-1}^i, \Delta_{t-1}^i] \\ &= \text{cov}[(1 - \rho^h) \xi \rho_m \tilde{\omega}_{t-1}^{sum,i} - (1 - \rho^{h+1}) \xi \tilde{\omega}_{t-1}^{sum,i}, (1 - \rho^{h+1}) \xi \tilde{\omega}_{t-1}^{sum,i}] \\ &= [\rho_m (1 - \rho^h) (1 - \rho^{h+1}) - (1 - \rho^{h+1})^2] \xi^2 \cdot \text{var}[\tilde{\omega}_{t-1}^{sum,i}]. \end{aligned}$$

Substituting these results into (G.5), we obtain

$$\gamma = \frac{\rho_m (1 - \rho^h)}{1 - \rho^{h+1}} - 1. \quad (\text{G.6})$$

This is the formula that we use to compute the values shown in the right panel of Figure 8. Because of (G.1), this formula implies bounds on the value of the regression coefficient given by $\gamma > -1$ and

$$\gamma < \frac{1 - \rho^h}{1 - \rho^{h+1}} - 1 < 1 - 1 = 0.$$

Thus the model predicts that $-1 < \gamma < 0$, as stated in the text.