

Simple Forecasts and Paradigm Shifts

HARRISON HONG, JEREMY C. STEIN, and JIALIN YU*

ABSTRACT

We study the asset pricing implications of learning in an environment in which the true model of the world is a multivariate one, but agents update only over the class of simple univariate models. Thus, if a particular simple model does a poor job of forecasting over a period of time, it is discarded in favor of an alternative simple model. The theory yields a number of distinctive predictions for stock returns, generating forecastable variation in the magnitude of the value-glamour return differential, in volatility, and in the skewness of returns. We validate several of these predictions empirically.

IN ATTEMPTING TO MAKE EVEN THE MOST BASIC KINDS of forecasts, we can find ourselves inundated with a staggering amount of potentially relevant raw data. To take a specific example, suppose you are interested in forecasting how General Motors' stock will perform over the next year. The first place you might turn is to GM's annual report, which is instantly available online. GM's 2004 10-K filing is more than 100 pages long, and is filled with dozens of tables as well as a myriad of other facts, footnotes, and esoterica. And this is just the beginning. With a few more clicks, it is easy to find countless news stories about GM, assorted analyst reports, and so forth.

How is one to proceed in the face of all this information? Both common sense, as well as a large literature in psychology, suggest that people simplify the forecasting problem by focusing their attention on a small subset of the available data. One powerful way to simplify is with the aid of a theoretical model. A parsimonious model focuses the user's attention on those pieces of information that are deemed to be particularly relevant for the forecast at hand, and has her disregard the rest.

Of course, it need not be normatively inappropriate for people to use simple models, even exceedingly simple ones. There are several reasons why simplifying can be an optimal strategy. First, there are cognitive costs to encoding and processing the additional information required by a more complex model. Second, if the parameters of the model need to be estimated, the parsimony

*Harrison Hong is from Princeton University, Jeremy Stein is from Harvard University and the National Bureau of Economic Research, and Jialin Yu is from Columbia University. We are grateful to the National Science Foundation for research support. Thanks also to Patrick Bolton, John Campbell, Glenn Ellison, David Laibson, Sven Rady, Andrei Shleifer, Christopher Sims, Lara Tiedens, Pietro Veronesi, Jeff Wurgler, Lu Zhang, and seminar participants at Princeton, the University of Zurich, the University of Lausanne, the Stockholm School of Economics, the Norwegian School of Management, the Federal Reserve Board, the AFA meetings and the NBER corporate finance and behavioral finance meetings for helpful comments and suggestions.

inherent in a simple model improves statistical power: For a given amount of data, one can more precisely estimate the coefficient in a univariate regression than the coefficients in a regression with many right-hand-side variables. Thus, simplicity clearly has its normative virtues. However, a central theme in much of the psychology literature is that people generally do something other than just simplifying in an optimal way. Loosely speaking, it seems that rather than having the meta-understanding that the real world is in fact complex and that simplification is only a strategy to deal with this complexity, *people tend to behave as if their simple models provide an accurate depiction of reality*.¹

Theoretical work in behavioral economics and finance has begun to explore some of the consequences of such normatively inappropriate simplification. For example, in many recent papers about stock market trading, investors pay attention to their own signals and disregard the signals of others, even when these other signals can be inferred from prices. The labels for this type of behavior vary across papers: Sometimes it is called “overconfidence” (in the sense of investors overestimating the relative precision of their own signals), sometimes it is called “bounded rationality” (in the sense that it is cognitively difficult to extract others’ signals from prices), and sometimes it is called “limited attention.” But labels aside, the reduced forms often look quite similar.² The common thread is that, in all cases, agents make forecasts based on a subset of the information available to them, yet they behave as if these forecasts were based on complete information.

While this general approach is helpful in understanding a number of phenomena, it also has an important limitation, in that it typically takes as *exogenous and unchanging* the subset of available information that an agent restricts herself to. For example, it may be reasonable to posit that investors with limited attention have a general tendency to focus too heavily on a firm’s reported earnings, while ignoring other numbers and footnotes.³ At the same time, it seems hard to believe that even relatively naïve investors would not lose some of their faith in this sort of valuation model following the highly publicized accounting scandals at firms such as Enron, WorldCom, and Tyco. If so, new questions arise: How rapidly will investors move in the direction of a new model, one that pays less attention to reported earnings and more attention to numbers that may help flag accounting manipulation or other forms of misbehavior? And, what will be the implications of this learning for stock returns?

Our goal in this paper is to begin to address these kinds of questions. As in previous work, we start with the assumption that agents use simple models

¹ For textbook discussions, see, for example, Nisbett and Ross (1980) and Fiske and Taylor (1991). We review this and related work in more detail below.

² A partial list includes: (1) Miller (1977), Harrison and Kreps (1978), Varian (1989), Kandel and Pearson (1995), Morris (1996), Odean (1998), Kyle and Wang (1997), Daniel et al. (1998), Hong and Stein (2003a), and Scheinkman and Xiong (2003), all of whom couch their models in terms of either differences of opinion or overconfidence; (2) Hong and Stein (1999), who appeal to bounded rationality; and (3) Hirshleifer and Teoh (2003), Sims (2003), Peng and Xiong (2006), and Della Vigna and Pollet (2004), who invoke limited attention.

³ See, for example, Hirshleifer and Teoh (2003) for a discussion of this idea.

that consider only a subset of available information. But unlike this other work, we then go on to explicitly analyze the process of learning and model change. In particular, we assume that agents keep track of the forecast errors associated with their simple models. If a given model performs poorly over a period of time, it may be discarded in favor of an alternative model, albeit an equally oversimplified one, that would have done better over the same period.

To be more precise, our setup can be described as follows. Imagine a stock that at each date t pays a dividend of $D_t = A_t + B_t + \varepsilon_t$, where A_t and B_t can be thought of as two distinct sources of public information, and ε_t is random noise. The idea that an agent uses an oversimplified model of the world can be captured by assuming that her forecasts are based on either the premise that $D_t = A_t + \varepsilon_t$ (we refer to this as having an “A model”) or the premise that $D_t = B_t + \varepsilon_t$ (we refer to this as having a “B model”). Suppose the agent initially starts out with the A model, and focuses only on information about A_t in generating her forecasts of D_t . Over time, the agent keeps track of the forecast errors that she incurs with the A model and compares them to the errors she would have made had she used the B model instead. Eventually, if the A model performs poorly enough relative to the B model, we assume that the agent switches over to the B model; we term such a switch a “paradigm shift.”⁴

This type of learning is Bayesian in spirit, and we use much of the standard Bayesian apparatus to formalize the learning process. However, there is a critical sense in which our agents are not conventional fully rational Bayesians: We allow them to update only over the class of simple univariate models. That is, their priors assign zero probability to the correct multivariate model of the world, so that no matter how much data they see, they can never learn the true model.⁵

This assumption yields a range of empirical implications, which we develop in a stock market setting. Even before introducing learning effects, the premise that agents use oversimplified models, and hence do not pay attention to all available information, allows us to capture well-known “underreaction” phenomena such as momentum (Jegadeesh and Titman (1993)) and post-earnings announcement drift (Bernard and Thomas (1989, 1990)). Nevertheless, the primary contribution of the paper lies in delineating the additional effects that arise from our learning mechanism. We highlight five of these effects. First, learning generates a value-glamour differential, or book-to-market effect (Fama

⁴ Our rendition of the learning process is inspired in part by Thomas Kuhn’s (1962) classic, *The Structure of Scientific Revolutions*. Kuhn argues that scientific observation and reasoning is shaped by simplified models, which he refers to as paradigms. During the course of what Kuhn refers to as “normal science,” a single generally accepted paradigm is used to organize data collection and make predictions. Occasionally, however, a crisis emerges in a particular field, when it becomes clear that there are significant anomalies that cannot be rationalized within the context of the existing paradigm. According to Kuhn, such crises are ultimately resolved by revolutions, or changes of paradigm, in which an old model is discarded in favor of a new one that appears to provide a better fit to the data.

⁵ The idea that agents attempt to learn but assign zero probability to the true model of the world is also in Barberis, Shleifer, and Vishny (1998). We discuss the connection between our work and this paper below.

and French (1992), Lakonishok, Shleifer, and Vishny (1994)). Second, and more distinctively, there is substantial variation in the conditional expected returns to value and glamour stocks. For example, a high-priced glamour stock that has recently experienced a string of negative earnings surprises, a situation one might label “glamour with a negative catalyst,” has an increased probability of a paradigm shift that will tend to be accompanied by a large negative return. Thus, the conditional expected return on the stock is more strongly negative than would be anticipated on the basis of its high price alone. Symmetrically, a low-priced value stock has an expected return that is more positive when it has also experienced a recent series of positive earnings surprises, that is, when it can be characterized as “value with a positive catalyst.”

The same reasoning also yields our third and fourth implications: Even with symmetric and homoskedastic fundamentals, both the volatility and skewness of returns are stochastic, with movements that can be partially forecasted based on observables. In the above example of a glamour stock that has experienced a series of negative earnings shocks, the increased likelihood of a paradigm shift corresponds to elevated conditional volatility as well as to negative conditional skewness.

Finally, these episodes will be associated with a kind of *revisionism*: When there are paradigm shifts, investors will tend to look back at old, previously available public information and to draw very different inferences from it than they had before. In other words, when asked to explain a dramatic movement in a company's stock price, observers may point to data that have long been in plain view in the company's annual reports, but that were overlooked under the previous paradigm.

In developing our results, we consider two alternative descriptions of the market-wide learning process. First, we examine a setting in which there is a single representative agent who does the same thing that researchers in economics and many other scientific fields typically do when they make model-based forecasts: She engages in *model selection*, that is, she picks a single favorite model, as opposed to *model averaging*. The model-selection case is particularly helpful in drawing out the intuition for our results, so we discuss it in some detail. But this approach naturally raises the question of how well our conclusions stand up in the presence of heterogeneity across investors, each of whom may have a different favorite model at any point in time. Therefore, we also consider the case of model averaging, which can be motivated by thinking of a continuum of investors, each of whom practices model selection, but applies a different threshold when deciding whether to switch from one model to another. Interestingly, the qualitative predictions that emerge are very similar to those in the model-selection case. This suggests that the key to these results is not the distinction between model selection versus model averaging, but rather the fact that, in either case, we restrict the updating process to the space of simple univariate models.

The rest of the paper is organized as follows. Section I reviews some of the literature in psychology that is most relevant for our purposes. In Section II, we lay out our theory and use heuristic arguments to outline its qualitative implications for stock returns. In Section III, we run a series of simulations

in order to make more quantitatively precise predictions, which we then go on to examine empirically. In Section IV, we briefly discuss the recent history of Amazon.com in an effort to illustrate the phenomenon of revisionism. Section V looks at the connection between our work and several related papers, and Section VI concludes.

I. Some Evidence from Psychology

The idea that people use overly simplified models of the world is a fundamental one in the field of social cognition. According to the “cognitive miser” view, which has its roots in the work of Simon (1982), Bruner (1957), and Kahneman and Tversky (1973), humans have to confront an infinitely complex and ever-changing environment, yet are endowed with a limited amount of processing capacity. Thus, in order to conserve on scarce cognitive resources, they use theories, or schema, to organize data and make predictions.

Schank and Abelson (1977), Abelson (1978), and Taylor and Crocker (1980) review and classify these knowledge structures, and highlight some of their strengths and weaknesses. These authors argue that theory-driven/schematic reasoning helps people to improve their performance at a number of tasks, including the interpretation of new information, the storage and retrieval of information in memory, the filling-in of gaps due to missing information, and overall processing speed. At the same time, there are also several disadvantages, such as incorrect inferences (due, for example, to stereotyping), oversimplification, a tendency to discount disconfirming evidence, and incorrect memory retrieval.⁶ Fiske and Taylor (1991, p. 13) summarize the cognitive miser view as follows:

The idea is that people are limited in their capacity to process information, so they take shortcuts whenever they can . . . People adopt strategies that simplify complex problems; the strategies may not be normatively correct or produce normatively correct answers, but they emphasize efficiency.

Indeed, much of the psychology literature takes more or less for granted the idea that people will not use all the available information in making their forecasts. Instead, this literature focuses on the specific biases that shape *which kinds* of information are most likely to be attended to. For example, according to the well-known availability heuristic (Tversky and Kahneman (1973)), people tend to overweight information that is easily available in their memories, that is, information that is especially salient or vivid.

⁶ Kuhn (1962) discusses an experiment by Bruner and Postman (1949) in which individual subjects are shown to be extremely dependent on a priori models when encoding the most simple kinds of data. In particular, while subjects can reliably identify standard playing cards (such as a *black six of spades*) after these cards have been displayed for just an instant, they have great difficulty in identifying anomalous cards (such as a *red six of spades*) even when they are given an order of magnitude more time to do so. However, once they are aware of the existence of the anomalous cards, that is, once their model of the world is changed, subjects can identify them as easily as the standard cards.

Our theory relies on the general notion that agents disregard some relevant information when making forecasts, but importantly, it does not invoke an exogenous bias against any one type of information. Thus, in our setting, A_t and B_t can be thought of as two sources of public information that are a priori equally salient; only after an agent endogenously opts to use the A model can A_t be said to become more “available.”

Another prominent theme in the work on theories and schemas is that of theory maintenance. Simply put, people tend to resist changing their models, even in the face of evidence that, from a normative point of view, would appear to strongly contradict these models. Rabin and Schrag (1999) provide an overview of much of this work, including the classic contribution of Lord, Ross, and Lepper (1979). Nevertheless, even if people are stubborn about changing models, one probably does not want to take the extreme position that they never learn from the data. As Nisbett and Ross (1980, p. 189) write,

Children do eventually renounce their faith in Santa Claus; once popular political leaders do fall into disfavor . . . Even scientists sometimes change their views. . . . No one, certainly not the authors, would argue that new evidence or attacks on old evidence can never produce change. Our contention has simply been that generally there will be less change than would be demanded by logical or normative standards or that changes will occur more slowly than would result from an unbiased view of the accumulated evidence.

Our efforts below can be seen as very much in the spirit of this quote. That is, while we allow for the possibility that it might take a relatively large amount of data to get an agent to change models, our whole premise is that, eventually, enough disconfirming evidence will lead to the abandonment of a given model and to the adoption of a new one.

Although the idea of theory maintenance is well developed, the psychology literature seems to have produced less of a consensus as to when and how theories ultimately change. Lacking such an empirical foundation, our approach here is intended to be as axiomatically neutral as possible. We measure the accumulated evidence against a particular model like a Bayesian would, that is, as the updated probability (given the data and a set of priors) that the model is wrong. However, we do not impose any further biases in terms of which sorts of data get weighted more or less heavily in the course of the Bayesian-like updating.

II. Theory

A. Basic Ingredients

A.1. Linear Specification for Dividends

We consider the market for a single stock. There is an infinite horizon, and at each date t , the stock pays a dividend of $D_t = F_t + \varepsilon_t \equiv A_t + B_t + \varepsilon_t$, where A_t and B_t can be thought of as two distinct sources of public information, and

ε_t is random noise. Each of the sources of information follows an AR(1) process, so that $A_t = \rho A_{t-1} + a_t$ and $B_t = \rho B_{t-1} + b_t$, with $\rho < 1$. The random variables a_t , b_t , and ε_t are all independently normally distributed, with variances of v_a , v_b , and v_ε , respectively. For the sake of symmetry and simplicity, we restrict ourselves in what follows to the case in which $v_a = v_b$.

Immediately after the dividend is paid at time t , investors see the realizations of a_{t+1} and b_{t+1} , which they can use to estimate the next dividend, D_{t+1} . Assuming a constant discount rate of r , this dividend forecast can then be mapped directly into an ex-dividend present value of the stock at time t . For a fully rational investor who understands the true structure of the dividend process and who uses both sources of information, the ex-dividend value of the stock at time t , which we denote by V_t^R , is given by $V_t^R = k(A_{t+1} + B_{t+1})$, where $k = 1/(1 + r - \rho)$ is a dividend capitalization multiple.

By contrast, we assume that investors use overly simplified univariate models to forecast future dividends, and hence to value the stock. In particular, at any point in time, any individual investor believes that one of the following possibilities obtains: (1) the dividend process is $D_t = A_t + \varepsilon_t$ (we refer to this as the “A model”) or (2) the dividend process is $D_t = B_t + \varepsilon_t$ (we refer to this as the “B model”). Thus, an investor who uses the A model at time t has an ex-dividend valuation of the stock V_t^A , which satisfies $V_t^A = kA_{t+1}$, and an investor using the B model at time t has a valuation V_t^B , where $V_t^B = kB_{t+1}$.⁷

A.2. Log-Linear Specification for Dividends

The above linear specification for dividends has a number of attractive features. First and foremost, it lets us write down some very simple closed-form expressions that highlight the central economic mechanisms at work in our theory. At the same time, the linear specification is less than ideal from an empirical realism perspective; for example, it allows for the possibility of negative dividends and prices, and it forces us to work with dollar returns rather than percentage returns. So, while we use the linear specification to help build intuition in the remainder of this section, when we calibrate the model for testing purposes in Section III, we also consider a log-linear variant in which $\log(D_t) = A_t + B_t + \varepsilon_t$, but in which the stochastic processes for A_t , B_t , and ε_t are the same as described above. Appendix B gives the details of how prices and returns are computed in the log-linear case.

⁷ Note that another possible univariate model is to forecast future dividends based solely on observed values of past dividends. That is, one can imagine a “D model,” where $V_t^D = kD_t$. As a normative matter, the D model may be more accurate than either the A or the B models. (This happens when v_ε is small relative to the variances of A_t and B_t .) But given their mistaken beliefs about the structure of the dividend process, agents will always consider the D model to be dominated by both the A and the B models.

B. Benchmark Case: No Learning

In order to have a benchmark against which to compare our subsequent results, we begin with a simple no-learning case in the context of the linear specification for dividends. Assume that there is a single investor who always uses the A model, so that the stock price at time t , P_t , is given by $P_t = V_t^A = kA_{t+1}$. The (simple) excess return from $t - 1$ to t , which we denote by R_t , is defined by $R_t = D_t + P_t - (1 + r)P_{t-1}$.⁸ It is straightforward to show that we can rewrite R_t as $R_t = z_t^A + ka_{t+1}$, where z_t^A is the forecast error associated with trying to predict the time- t dividend using model A, that is, where $z_t^A = B_t + \varepsilon_t$. In other words, under the A model, the excess return at time t has two components, (1) the forecast error z_t^A , and (2) the incremental A-news about future dividends, ka_{t+1} .

With these variables in hand, some basic properties of stock returns can be immediately established. Consider first the autocovariance of returns at times t and $t - 1$. We have that $\text{cov}(R_t, R_{t-1}) = \text{cov}(z_t^A, z_{t-1}^A) + k\text{cov}(z_t^A, a_t)$. With a little manipulation, this yields

$$\text{cov}(R_t, R_{t-1}) = \rho v_b / (1 - \rho^2). \quad (1)$$

This expression reflects the positive short-run momentum in returns that arises from a “repeating-the-same-mistake” effect. Since the investor uses the same wrong model to make forecasts for times $t - 1$ and t , in both cases ignoring the persistent B information, her forecast errors z_{t-1}^A and z_t^A are positively correlated, which tends to induce positive autocovariance in returns.

Another item of interest is the covariance between the price level and future returns, $\text{cov}(R_t, P_{t-1})$. Since all dividends are paid out immediately as realized (there are no retained earnings), and since the scale of the dividend process never changes over time, it makes sense to think of the stock as a claim on an asset with a constant underlying book value. Thus, one can interpret the price of the stock, which is stationary in our model, as an analog to the market-to-book ratio, and $\text{cov}(R_t, P_{t-1})$ as a measure of how strongly this ratio forecasts returns. With no learning, it is easy to show that $\text{cov}(R_t, P_{t-1}) = 0$.

Thus, absent any learning considerations, the linear specification for dividends delivers a momentum-like pattern in stock returns, but nothing else. In particular, there is no value-glamour effect, and returns are symmetrically and homoskedastically distributed.⁹

⁸ Again, when using the linear dividend specification, it is easier to work with arithmetic returns as opposed to percentage returns. Given that the price level is stationary in our setting, this is a relatively innocuous choice.

⁹ The no-learning case can be enriched by allowing for heterogeneity among investors. Suppose a fraction f of the population uses model A, and $(1 - f)$ uses model B. We can demonstrate that this setup still generates momentum in stock returns. More interestingly, momentum is strongest when there is maximal heterogeneity among investors, that is, when $f = 1/2$. Since such heterogeneity also generates trading volume, we have the prediction that momentum will be greater when there is more trading volume, which fits nicely with the empirical findings of Lee and Swaminathan (2000). Although this extension of the no-learning case strikes us as promising, we do not pursue it in detail here, as our main goal is to draw out the implications of our particular learning mechanism.

C. Learning: Further Ingredients

To introduce learning, we must specify several further assumptions. The first of these is that at any point in time t , an agent believes that the dividend process is governed by either the A model or the B model, that is, she believes that either $D_t = A_t + \varepsilon_t$ or $D_t = B_t + \varepsilon_t$. The crucial point is that the agent always wrongly thinks that the true process is a univariate one and attaches zero probability to the correct, bivariate model of the world.

For the purposes of a general analytical treatment, we allow for the possibility that the agent might believe that the underlying dividend process switches over time, between being driven by the A model versus the B model, according to a Markov chain. Let π_A be the conditional probability that the agent attaches to dividends being generated by the A model in the next period given that they are being generated by the A model in the current period, and define π_B symmetrically. Finally, to keep things simple, set $\pi_A = \pi_B = \pi$.

In our simulations, we focus on the limiting scenario of $\pi = 1$, in which the agent (correctly) thinks that nature is unchanging, so that there is only a single model that applies for all time. This strikes us as the simplest and most naturally motivated account of beliefs. The only technical issue it raises is that with $\pi = 1$, the extent of learning is nonstationary, and depends on the length of the sample period: After a long stretch of time, there is a high probability that the agent will be almost convinced by one of the two models, thereby making further paradigm shifts extremely unlikely.

Alternatively, if one is interested in making the learning process stationary, and thereby giving our results a more steady-state flavor, one can assume that $1/2 < \pi < 1$, which means that the agent thinks that both states are persistent but not perfectly absorbing. As a practical matter, it turns out that when we simulate stock prices and learning over empirically plausible horizons, we obtain very similar results either way, so the fundamental predictions of the model do not turn on whether we assume $\pi = 1$ or $\pi < 1$.

With the assumptions in place, a first step is to describe how Bayesian updating works, given the structure and the set of priors that we have specified. It is important to stress that in our setting, one does not want to interpret such Bayesian updating as corresponding to the behavior of a fully rational agent, since we have restricted the priors in such a way that no weight can ever be attached to the correct model of the world. Let p_t be the probability weight on the A model going into period t . To calculate the posterior going into period $t + 1$, recall that for each model, we can construct an associated forecast error, with $z_t^A = B_t + \varepsilon_t$ being the error from the A model and $z_t^B = A_t + \varepsilon_t$ being the error from the B model. Intuitively, the updating process should tilt more in the direction of model A after period t if z_t^A is smaller than z_t^B in absolute value, and vice versa.

More precisely, conditional on the A model as well as on the realization of A_t , D_t has a normal density with mean A_t and variance v_ε , which we denote by $f_A(D_t | A_t)$, and that satisfies

$$f_A(D_t | A_t) = \frac{1}{\sigma_\varepsilon} \phi \left(\frac{D_t - A_t}{\sigma_\varepsilon} \right) = \frac{1}{\sigma_\varepsilon} \phi \left(\frac{z_t^A}{\sigma_\varepsilon} \right), \quad (2)$$

where $\phi(\cdot)$ is the standard normal density and σ_ε is the square root of v_ε . Similarly, conditional on the B model as well as on the realization of B_t , D_t has a normal density with mean B_t and variance v_ε , which we denote by $f_B(D_t | B_t)$, and that satisfies

$$f_B(D_t | B_t) = \frac{1}{\sigma_\varepsilon} \phi \left(\frac{D_t - B_t}{\sigma_\varepsilon} \right) = \frac{1}{\sigma_\varepsilon} \phi \left(\frac{z_t^B}{\sigma_\varepsilon} \right). \quad (3)$$

Next, we define the variable x_{t+1} as follows:

$$x_{t+1} = p_t L_z / (p_t L_z + (1 - p_t)), \quad (4)$$

where L_z is the likelihood ratio given by

$$L_z = f_A(D_t | A_t) / f_B(D_t | B_t) = \exp \left(-[(z_t^A)^2 - (z_t^B)^2] / 2v_\varepsilon \right). \quad (5)$$

Note that the likelihood ratio is always nonnegative, and increases the smaller is z_t^A relative to z_t^B in absolute value. With these definitions in place, standard arguments can be used to show that the Bayesian posterior going into period $t + 1$ is given by (see, e.g., Barberis, Shleifer, and Vishny (1998), or Hong and Rady (2002))

$$p_{t+1} = p^* + (\pi_A + \pi_B - 1)(x_{t+1} - p^*), \quad (6)$$

where $p^* = (1 - \pi_B) / (2 - \pi_A - \pi_B)$ is the fraction of the time that the dividend process is expected to spend in the A-model state over the long run. Given our assumption that $\pi_A = \pi_B$, it follows that $p^* = 1/2$, and (6) reduces to

$$p_{t+1} = 1/2 + (2\pi - 1)(x_{t+1} - 1/2). \quad (7)$$

Observe that in the limiting case in which $\pi = 1$, we have that $p_{t+1} = x_{t+1}$. This is the point mentioned earlier: Bayesian beliefs in this case are nonstationary, and eventually drift toward a value of either zero or 1. In contrast, if $\pi < 1$, Bayesian beliefs are stationary, with a long-run mean weight of $1/2$ being attached to the A model. In either case, however, it is clear that the updating process leans more toward the A model after period t if z_t^A is smaller than z_t^B in absolute value, and vice versa.

An essential piece of intuition for understanding the results that follow comes from asking how the *speed of learning* varies over time. Heuristically, the speed of learning measures the rate at which p_t adjusts toward either 1 (perfect certainty in the A model) or zero (certainty in the B model). O'Hara (1995) establishes that the speed of learning is proportional to relative entropy. In our setting, the relative entropy Ψ_t is given by

$$\Psi_t = \int_{-\infty}^{\infty} f_A(D_t | A_t) \log \frac{f_A(D_t | A_t)}{f_B(D_t | B_t)} dD_t. \quad (8)$$

Straightforward calculation based on (8) yields

$$\Psi_t = \frac{(A_t - B_t)^2}{2v_\varepsilon}. \quad (9)$$

Equation (9) says that there is more rapid learning in period t when A_t and B_t are further apart. This makes intuitive sense. In the limit, if $A_t = B_t$, the two models generate exactly the same forecasts, so there is no scope for distinguishing them in the data. In contrast, when the two models generate widely divergent forecasts, the next realization of dividends has the potential to discriminate strongly in favor of one or the other.

This observation gets to the heart of why there can be predictable variation in various moments of stock returns in our framework. Consider as an example volatility. If an econometrician can infer when A_t and B_t are relatively far apart, then, according to (9), he will be able to estimate when the potential for learning is high, and by extension, when stock return volatility is likely to be above its unconditional average.

D. Model Selection

As noted above, one way to proceed is to think of the market as a whole in terms of a single representative investor and to assume that this representative investor practices model selection. In other words, at time t , the representative investor has a preferred null model that she uses exclusively. Moreover, as long as the accumulated evidence against the null model is not too strong, it is carried over to time $t + 1$.

To be more precise, we define the indicator variable I_t^A to be equal to 1 (zero) if the investor's null model at time t is the A (B) model. We then assume the following dynamics for I_t^A

$$\text{If } I_t^A = 1, \quad \text{then } I_{t+1}^A = 1, \quad \text{unless } p_{t+1} < h. \quad (10)$$

$$\text{If } I_t^A = 0, \quad \text{then } I_{t+1}^A = 0, \quad \text{unless } p_{t+1} > (1 - h). \quad (11)$$

Here, h is a critical value that is less than one half. Thus, the investor maintains a given null model for the purposes of making forecasts until the updated (Bayesian) probability of it being correct falls below the critical value. So, for example, if her original null is the A model and $h = 0.05$, she continues to make forecasts exclusively with it until it is rejected at the 5% confidence level. Once this happens, the B model assumes the status of the null model and it is then used exclusively until it too is rejected at the 5% confidence level. Clearly, the smaller is h , the stronger is the degree of resistance to model change; the psychological literature on theory maintenance discussed above can therefore be thought of as suggesting a value of h relatively close to zero.

This formulation raises an important issue of interpretation. On the one hand, we motivate the assumption that the investor uses a univariate forecasting model at any point in time by appealing to limited cognitive resources, the

idea being that it is too difficult to simultaneously process both the A and B sources of information for the purposes of making a forecast. Yet, the investor does use both the A and B sources of information when deciding whether to abandon her null model: The Bayesian updating process for p_t that underlies her model selection criterion depends on both z_t^A and z_t^B . In other words, the investor is capable of performing quite sophisticated multivariate operations when evaluating which model is better, but is unable to make dividend forecasts based on more than a single variable at a time, which sounds somewhat schizophrenic.

One resolution to this apparent paradox relies on the observation that, in spite of the way in which we formalize things, it is not necessary for our results that the representative investor actively reviews her choice of models as frequently as once every period. Indeed, it is more plausible to think of the two basic tasks that the investor undertakes—forecasting and model selection—as happening on different time scales, and therefore involving different trade-offs of cognitive costs and benefits. For an active stock market participant, dividend forecasts have to be updated continuously as new information comes in. Thus, the model that generates these forecasts needs to be simple and not too cognitively burdensome, or it will be impractical to use in real time.¹⁰

In contrast, it may well be that the investor steps back from the ongoing task of forecasting and does systematic model evaluation only once in a long while; as a result, it might be feasible for this process to be more data intensive.¹¹ Indeed, it is not difficult to incorporate this sort of timing feature explicitly into our analysis, by allowing the investor to engage in model evaluation only once every m periods, with m relatively large. Our limited efforts at experimentation suggest that this approach yields results that are qualitatively similar to those we report below.

E. Model Averaging

As will become clear, the representative-investor/model-selection approach described above provides a useful way to communicate the main intuition behind our results. But it is important to underscore that these results do not hinge on the discreteness associated with the model selection mechanism. To illustrate this point, we also consider the “smoother” case in which the market

¹⁰ This is why we are reluctant to assume that any individual agent acts as a model averager. If a model averager assigns a probability p_t to the A model at time t , her forecast of the next dividend would be $p_t A_{t+1} + (1 - p_t) B_{t+1}$. However, such a forecast is no longer a cognitively simple one to make in real time, as it requires the agent to make use of both sources of information simultaneously. And if we are going to endow the agent with this much high-frequency processing power, it is less clear how one motivates the assumption that she does not consider more complicated models in her set of priors.

¹¹ Moreover, much of this low-frequency model evaluation may happen at the level of an entire investment community, rather than at the level of any single investor. For example, each investor may need to work alone with a given simple model to generate her own high-frequency forecasts, but may once in a while change models based on what she reads in the press, hears from fellow investors, etc. Again, the point to be made is that no single investor is literally going to be engaging in cognitively costly model evaluation on a continuous basis.

price is based on model averaging, that is, where $P_t = p_t kA_{t+1} + (1 - p_t) kB_{t+1}$. One way to motivate such model averaging is by appealing to a particular form of heterogeneity across investors.

To see this, suppose that there is a continuum of investors distributed uniformly across the interval $(0, 1)$, each of whom individually practices model selection. All investors share the same underlying Bayesian update p_t of the probability of the A model being correct at time t , with p_t evolving as before. But now, each investor has her own fixed threshold for determining when to use the A model as opposed to the B model: The investor located at point i on the interval uses the A model if and only if $p_t > i$.¹² This implies that the fraction of investors in the population using the A model at time t is given by p_t . To the extent that the market price is just the weighted average of individual investors' estimates of fundamental value, this in turn implies that $P_t = p_t kA_{t+1} + (1 - p_t) kB_{t+1}$.¹³

F. Implications for Stock Returns: Some Intuition

In Section III below, we use a series of simulations to provide a full-blown quantitative analysis that covers both the linear and log-linear specifications for dividends, as well as the cases of model selection and model averaging. But before doing so, we attempt to provide a heuristic sense for the mechanism driving our results. This is most transparently done in the context of the linear specification with model selection, so we focus exclusively on this one combination for the remainder of this section.

Assuming that we are in a model-selection world, suppose for the moment that the representative investor is using the A model at time $t - 1$, so that $P_{t-1} = kA_t$. There are two possibilities at time t . The first is that there will be no paradigm shift, so that the investor continues to use the A model. In this case, $P_t = kA_{t+1}$, and the return at time t , which we denote by R_t^N , is given by

$$R_t^N = z_t^A + ka_{t+1} = B_t + \varepsilon_t + ka_{t+1}. \quad (12)$$

Alternatively, if there is a paradigm shift at time t , the investor switches over to using the B model, in which case the price is $P_t = kB_{t+1}$, and the return, denoted by R_t^S , is

¹² One can interpret investors with low thresholds as those who have an innate preference for the A model.

¹³ This motivation is admittedly loose. In a dynamic model, it is not generally true that price simply equals the weighted average estimate of fundamental value as short-term trading considerations may arise; for example, investors may try to forecast the forecasts of others. Nevertheless, since we just want to demonstrate that our results are not wholly dependent on model selection, the model-averaging case is a natural point of comparison. An alternative way to motivate model averaging is in terms of a single representative investor who is a classical Bayesian (given the set of priors described above) and who therefore puts weight p_t on the A model at time t . Another advantage of this interpretation is that it avoids the "schizophrenia" problem alluded to above, since the representative investor now uses both sources of information in making her forecasts at any point in time. The disadvantage is that it is no longer the case that every individual actor makes forecasts that are simple in nature.

$$R_t^S = z_t^A + kb_{t+1} + \rho k(B_t - A_t) = B_t + \varepsilon_t + kb_{t+1} + \rho k(B_t - A_t). \quad (13)$$

Observe that $R_t^S = R_t^N + k(b_{t+1} - a_{t+1}) + \rho k(B_t - A_t)$. Simply put, the return in the paradigm shift case differs from that in the no-shift case as a result of current and lagged A-information being discarded from the price and replaced with B-information.

Let us begin by revisiting the magnitude of the value-glamour effect, as proxied for by $\text{cov}(R_t, P_{t-1})$. (Recall that we have $\text{cov}(R_t, P_{t-1}) = 0$ in the no-learning case.) In Appendix A, we show that irrespective of whether we begin in the A regime or the B regime, $\text{cov}(R_t, P_{t-1})$ can be decomposed as follows

$$\begin{aligned} \text{cov}(R_t, P_{t-1}) &= \text{cov}(R_t^S, P_{t-1}/\text{shift}) * \text{prob}(\text{shift}) \\ &\quad + \text{cov}(R_t^N, P_{t-1}/\text{no shift}) * \text{prob}(\text{no shift}). \end{aligned} \quad (14)$$

Substituting in the definitions of R_t^N and R_t^S from (12) and (13), and simplifying, we can rewrite (14) as

$$\begin{aligned} \text{cov}(R_t, P_{t-1}) &= k\{\text{cov}(\varepsilon_t, A_t) + \text{cov}(A_t, B_t)\} \\ &\quad + \rho k^2\{\text{cov}(A_t, B_t/\text{shift}) - \text{var}(A_t/\text{shift})\} * \text{prob}(\text{shift}). \end{aligned} \quad (15)$$

Note that both the $\text{cov}(\varepsilon_t, A_t)$ term as well as the first $\text{cov}(A_t, B_t)$ term in (15) are *unconditional* covariances. We have been assuming all along that these unconditional covariances are zero. Thus, (15) can be further reduced to

$$\text{cov}(R_t, P_{t-1}) = \rho k^2\{\text{cov}(A_t, B_t/\text{shift}) - \text{var}(A_t/\text{shift})\} * \text{prob}(\text{shift}). \quad (16)$$

Equation (16) clarifies the way in which a value-glamour effect arises in the presence of learning. A preliminary observation is that $\text{cov}(R_t, P_{t-1})$ can only be nonzero to the extent that the probability of a paradigm shift, $\text{prob}(\text{shift})$, is nonzero: As we have already seen, there is no value-glamour effect absent learning. When $\text{prob}(\text{shift}) > 0$, two distinct mechanisms are at work. First, there is a negative contribution from the $-\text{var}(A_t/\text{shift})$ term. This term reflects the fact that A-information is abruptly removed from the price at the time of a paradigm shift. This tends to induce a negative covariance between the price level and future returns, since, for example, a positive value of A_t at time $t - 1$ will lead to a high price at this time, and then to a large negative return when this information is discarded from the price at time t .

Second, and more subtly, there is the $\text{cov}(A_t, B_t/\text{shift})$ term. While the unconditional covariance between A_t and B_t is zero, the covariance *conditional on a paradigm shift* is not. To see why, think about the circumstances in which a shift from the A model to the B model is most likely to occur. Such a shift will tend to happen when the underlying Bayesian posterior p_t moves sharply, that is, when there is a lot of Bayesian learning. According to equation (9), the relative entropy Ψ_t , and hence the speed of learning, is greatest when A_t and B_t are far apart. Put differently, if $A_t = B_t$, there is no scope for Bayesian learning, and hence no possibility of a paradigm shift.

This line of reasoning suggests that $\text{cov}(A_t, B_t/\text{shift}) < 0$, which in turn makes the overall value of $\text{cov}(R_t, P_{t-1})$ in (16) even more negative, thereby strengthening the value-glamour differential.¹⁴ When a paradigm shift occurs, not only is A-information discarded from the price, it is also replaced with B-information. And conditional on a shift occurring, these two pieces of information tend to be pointing in opposite directions. So, if a positive value of A_t at $t - 1$ has led to a high price at this time, there will tend to be an extra negative impact on returns in the event of a paradigm shift at t (above and beyond that associated with just the discarding of A_t), when B_t enters into the price for the first time.

Importantly, in our setting learning generates more than just a simple time-invariant value-glamour effect: It also creates predictable variation in the expected returns to value and glamour stocks. To see why, recall that return predictability based on price levels is entirely concentrated in those periods in which paradigm shifts occur. Thus, if an econometrician can track variation over time in the probability of a paradigm shift, he will also be able to forecast when such predictability is likely to be the greatest.

Again, the key piece of insight comes from the expression for relative entropy Ψ_t in (9), which tells us that there is more potential for learning when the A model and the B model produce divergent forecasts. What does this mean in terms of observables? To be specific, think of a situation in which A_t is very positive, so the stock is a high-priced glamour stock. Going forward, there will be more scope for learning if, in addition, B_t is negative. This will tend to show up as negative values of the forecast error z_t^A , since $z_t^A = B_t + \varepsilon_t$. In other words, if a high-priced stock is experiencing negative forecast errors, this is a clue that the two models are at odds with one another.

A sharper prediction of our theory, therefore, is that a high-priced glamour stock will be particularly vulnerable to a paradigm shift, and hence to a sharp decline in prices, after a series of negative z -surprises about fundamentals. One might term such an especially bearish situation “glamour with a negative catalyst.” The conversely bullish scenario, “value with a positive catalyst,” involves a low-priced value stock and a series of positive z -surprises.¹⁵ The closest empirical analog to such z -surprises would probably be either: (1) a measure of realized earnings in a given quarter relative to the median analyst’s earnings forecast, or (2) the stock price response to an earnings announcement. In our empirical work, we use the latter of these two variables as a proxy for z -surprises.

¹⁴ We are able to prove analytically that $\text{cov}(A_t, B_t/\text{shift}) < 0$ for the limiting case in which the persistence parameter ρ approaches zero. (The proof is available on request). In addition, we exhaustively simulate the model over the entire parameter space to verify that this condition holds everywhere else.

¹⁵ The idea that value and/or glamour effects are more pronounced in the presence of such catalysts has some currency among practitioners. For example, the Bernstein Quantitative Handbook (2004) presents a variety of quantitative screens that “we believe lead to outperformance.” One of these screens, labeled “Value With a Catalyst,” is chosen to select “undervalued stocks reporting a positive earnings surprise” (pp. 22–23).

When we say that a glamour stock has more negative expected returns conditional on a recent string of disappointing earnings surprises, we need to stress a crucial distinction. This phenomenon is *not* simply a result of adding together the unconditional value-glamour and momentum effects. Rather, in the context of a regression model that forecasts future returns, our theory predicts that not only should there be book-to-market and momentum variables, but also *interaction terms* that represent the product of book-to-market with proxies for the direction of recent earnings surprises. In other words, we would expect an interaction term for glamour and bad news to attract a negative coefficient, and an interaction term for value and good news to attract a positive coefficient. We highlight this prediction in both our simulations and our empirical work below.

The same basic mechanisms produce forecastable movements in stock return volatility and skewness. As a comparison of equations (12) and (13) makes clear, volatility is inherently stochastic in our setting because returns have more variance at times of paradigm shifts than at other times. Moreover, these movements in volatility can be partially forecasted by an econometrician using exactly the same logic as above. For example, a high-priced glamour stock is more apt to experience a paradigm shift, which will manifest itself not only as a negative return, but also as an unusually large absolute price movement, after a sequence of negative fundamental surprises. Again, this is because such negative surprises are an indicator that the A and B models are in disagreement, which, according to the relative entropy formula in (9), raises the potential for learning.

Analogous arguments apply for conditional skewness. First, glamour stocks will tend to have more negatively skewed returns than value stocks. This is because the very largest movements in glamour stocks, those associated with paradigm shifts, will on average be negative, and conversely for value stocks. This feature of our theory is reminiscent of classic accounts of bubbles, as the potential for the sudden popping of a bubble in a high-priced glamour stock similarly generates negative conditional skewness. However, while the popping of a bubble is exogenous in Blanchard and Watson (1982), our theory endogenizes it.¹⁶ Moreover, we have the sharper prediction that these general skewness effects will be more pronounced if one further conditions on recent news. Thus, the negative skewness in a glamour stock will be strongest after it has experienced a recent string of bad news, and the positive skewness in a value stock will be greatest after a string of good news.

Although we have focused our discussion on the model-selection case, the intuition for the model-averaging case is very similar. With model selection, the notion of effective learning at the market level is dichotomous: Either there is a paradigm shift in a given period, or there is not. But this discreteness is not what is driving the results. Rather, what matters for the various asset pricing patterns is that an econometrician can forecast when there is likely to be a lot of learning, that is, he can tell when the A and B models are pointing in opposite directions. With model averaging, the amount of market-wide learning that

¹⁶ Abreu and Brunnermeier (2003) can also be thought of as a theory that endogenizes the collapse of bubbles.

takes place is a continuous variable, but the econometrician can still partially forecast it for the same reason as before. In particular, when a glamour stock is observed to have a series of negative earnings surprises, this suggests that there is a divergence between the A and B models, which, according to equation (9), tells us that the relative entropy, and hence the speed of learning, is likely to be high. The implications for conditional variation in value and glamour return premia, in volatility, and in skewness all follow from this ability to anticipate variation in the intensity of learning over time.

III. Simulations and Empirical Tests

In order to flesh out the implications of the theory more fully, and to assess their quantitative importance, we now turn to a series of simulations. The simulations cover both the linear and log-linear dividend specifications, as well as the model-selection and model-averaging cases. However, before turning to the details, we should stress an important general caveat. When we generate a panel of stock returns, we do so by applying our learning model to each individual stock in the panel *independently*. In other words, we assume that all learning happens at the stock level and is uncorrelated across stocks. This may well not be the most attractive assumption; for example, it may make more sense to posit that investors apply a common paradigm to all stocks in the same industry.

We do not explore the implications of such correlated learning for stock returns, but depending on exactly how it is modeled, it would appear to have the potential to introduce a variety of further complexities. To take just one example, correlated learning will tend to make all stocks in an industry co-move together strongly. This raises the possibility that some of what we are currently interpreting as a value-glamour effect might be “explained away” by differences in factor loadings of one sort or another.

This caveat must be borne in mind when comparing our simulation results to the data. To the extent that our current formulation of the learning process omits some potentially important elements, the empirical analysis should not be thought of as an attempt to test the broader theory in a quantitatively precise fashion. Rather, the goal is to see if a first-generation version of the theory can deliver effects of an economically interesting magnitude, and to highlight the dimensions along which the current version appears to fall short.

A. Calibration

In each of our simulations, we create a panel of 2,500 independent stocks, which we then track for 100 periods. When we calibrate the parameters, we treat each period as corresponding to one calendar quarter, so that with 100 periods, we have a 25-year panel. (This matches up closely with the length of our empirical sample period, which runs from 1971 to 2004.) We then repeat each of these 2,500-stock-by-100-quarter exercises 100 times. As will become clear, this appears to be more than sufficient to generate precise estimates of the moments of interest.

The simulations require that we specify the following parameters: The variances v_a , v_b , and v_ε , the persistence parameter ρ , the discount rate r , the Markov transition parameter π , and the model rejection critical value h . Note, however, that h only plays a role in the case of model selection—we do not need to specify a value of h for the model-averaging case.

We begin by setting $\pi = 1$ and $h = 0.05$. The former assumption corresponds to the scenario in which agents believe that there is a single simple model that is correct for all time, that is, agents do not believe that there is regime shifting with respect to the underlying model of the world.¹⁷ The latter assumption implies that the status quo model is discarded when the updated probability of it being correct falls below 5%. We set the discount rate to $r = 0.015$, which corresponds to an annualized value of 6%. We also simplify things by assuming that all the variances are the same, that is, that $v_a = v_b = v_\varepsilon = v$. Our task then boils down to coming up with empirically realistic values of v and of the persistence parameter ρ .

We pick these two parameters so as to roughly match observed levels of earnings persistence and stock return volatility. Given the assumption that $v_a = v_b = v_\varepsilon = v$, the autocorrelation properties of dividends in our model are entirely pinned down by the persistence parameter ρ . (See Appendix C for details.) We set $\rho = 0.97$, which implies a first-order autocorrelation of log dividends (in the log-linear specification) of 0.94. This lines up closely with the value of the first-order autocorrelation coefficient of 0.96 that we estimate using quarterly data on the log of real S&P operating earnings over the period 1988 to 2004.¹⁸

Once all the other parameters have been chosen, there is a one-to-one mapping between v and stock return volatility, although this mapping depends on the nature of the learning process (i.e., model selection vs. model averaging) and is not something that we can express in closed form. After some experimentation, we set v to 0.00001 in the linear specification and to 0.045 in the log specification. As we will see momentarily, these values lead to annualized stock return volatilities in the neighborhood of 30%.

B. Simulation Results: Linear Dividend Specification

Table I displays our simulation results for the linear dividend specification. The table contains three panels: Panel A for the no-learning benchmark case, Panel B for the case of model selection, and Panel C for the case of model averaging. Within each panel, we display two sets of three regressions each; these are simply Fama–MacBeth (1973) regressions that we run on the simulated data samples. Again, recall that the samples are 2,500-stock-by-100-quarter panels.

¹⁷ However, as a robustness check, we redo all of the simulations below with a steady-state version of the model in which π is reset to 0.95. Given our 25-year simulation horizon, the results are very similar, both qualitatively and quantitatively, to those with $\pi = 1$.

¹⁸ We use data on operating earnings, rather than dividends, for calibration purposes. This is because unlike in our theoretical setting, real-world dividends are not exogenous, but rather are heavily smoothed by managers. Thus, observed earnings arguably provide a better match for the theoretical construct of “dividends.”

Table I
Simulation Results for the Linear Model

We simulate the dividends for a cross section of 2,500 stocks for 100 quarters. The variances of the shocks are set to $v_a = v_b = v_e = 0.00001$. The autocorrelation coefficient of the processes A_t and B_t is set to $\rho = 0.97$. The quarterly interest rate is set to $r = 0.015$. The Markov transition parameters $\pi_A = \pi_B = 1$ and $h = 0.051$. We use the linear pricing model to generate stock prices for the three cases of No Learning, Model Selection, and Model Averaging. For each of these three cases, we run three sets of Fama–MacBeth (1973) regressions, to forecast stock returns, return volatility, and return skewness. The dependent variables include the following. RET is the annualized return for quarter t . VOL is the annualized volatility calculated using four quarters of returns (t to $t + 3$). SKEW is the skewness coefficient calculated using four quarters of returns (t to $t + 3$). The independent variables include a constant term (not reported) and the following. $Price(t - 1)$ is the lagged stock price, normalized to have zero mean and unit standard deviation. $News(t - 4, t - 1)$ is the cumulative z -surprise over the previous four quarters, normalized to have zero mean and unit standard deviation. For the 2×2 sort, $Value * GoodNews$ is a dummy variable that equals 1 if $Price(t - 1)$ is below the median value and $News(t - 4, t - 1)$ is above the median value for that quarter, while $Glamour * BadNews$ is a dummy variable that equals 1 if $Price(t - 1)$ is above the median value and $News(t - 4, t - 1)$ is below the median value for that quarter. For the 3×3 sort, $Value * GoodNews$ is a dummy variable that equals one if $Price(t - 1)$ is in the lowest one third of values and $News(t - 4, t - 1)$ is in the top one third of values for that quarter, while $Glamour * BadNews$ is a dummy variable that equals 1 if $Price(t - 1)$ is in the top one third of values and $News(t - 4, t - 1)$ is in the bottom one third of values for that quarter. For each simulation, we recover and save the Fama–MacBeth regression coefficients (the time-series average of the cross-sectional regression coefficients). We run 100 simulations and then take the average of the Fama–MacBeth coefficients, which are reported in the panels below. The standard errors are calculated simply as the standard deviation of these coefficients across simulations, divided by the square root of 100, and the associated t -statistics are reported in parentheses.

Panel A: No Learning (Unconditional Annualized Stock Return Volatility of 16.9%)						
	2 × 2 Sort			3 × 3 Sort		
	RET	VOL	SKEW	RET	VOL	SKEW
<i>Value * GoodNews</i>	0.0001 (0.34)	−0.0001 (1.24)	−0.0004 (0.75)	−0.0001 (0.34)	0.0000 (0.51)	0.0000 (0.06)
<i>Glamour * BadNews</i>	−0.0002 (1.30)	0.0001 (1.07)	0.0001 (0.23)	−0.0005 (2.15)	0.0000 (0.26)	0.0008 (1.53)
<i>Price(t − 1)</i>	0.0001 (1.58)	−0.0001 (1.69)	−0.0002 (0.81)	0.0001 (1.67)	0.0000 (0.48)	−0.0002 (0.98)
<i>News(t − 4, t − 1)</i>	0.0579 (612.46)	0.0000 (1.22)	0.0001 (0.44)	0.0579 (613.73)	0.0000 (0.10)	0.0001 (0.64)
<i>(continued)</i>						

(continued)

Table I—Continued

Panel B: Model Selection (Unconditional Annualized Stock Return Volatility of 30.4%)						
	2 × 2 Sort			3 × 3 Sort		
	RET	VOL	SKEW	RET	VOL	SKEW
<i>Value * GoodNews</i>	0.0760 (158.55)	0.0598 (210.23)	0.0954 (161.75)	0.0095 (12.45)	0.0906 (230.26)	−0.0148 (16.27)
<i>Glamour * BadNews</i>	−0.0757 (157.49)	0.0597 (221.31)	−0.0961 (162.19)	−0.0094 (13.78)	0.0906 (186.75)	0.0128 (14.43)
<i>Price(t − 1)</i>	−0.0407 (148.35)	0.0000 (0.06)	−0.0087 (34.71)	−0.0710 (337.21)	−0.0001 (0.59)	−0.0539 (229.67)
<i>News(t − 4, t − 1)</i>	0.0760 (342.38)	−0.0001 (0.42)	0.0364 (155.88)	0.1020 (513.02)	0.0000 (0.26)	0.0758 (303.57)
Panel C: Model Averaging (Unconditional Annualized Stock Return Volatility of 31.1%)						
<i>Value * GoodNews</i>	0.0853 (201.06)	0.0657 (212.50)	0.0610 (118.35)	0.0564 (63.41)	0.1003 (184.66)	0.0208 (27.00)
<i>Glamour * BadNews</i>	−0.0854 (157.69)	0.0659 (232.55)	−0.0623 (96.26)	−0.0564 (73.49)	0.1007 (213.66)	−0.0225 (24.73)
<i>Price(t − 1)</i>	−0.0386 (137.52)	−0.0001 (0.41)	−0.0114 (46.14)	−0.0600 (281.84)	−0.0001 (0.91)	−0.0321 (145.44)
<i>News(t − 4, t − 1)</i>	0.0617 (275.11)	0.0000 (0.09)	0.0198 (90.33)	0.0786 (345.85)	0.0001 (0.29)	0.0368 (159.53)

The numbers reported in the tables are the mean coefficients across the 100 trials of each panel regression, along with the t -statistics associated with these means.

In the first regression of each set, we forecast (annualized) returns in quarter t based on four variables: (1) A value-glamour proxy, namely, the price level at the end of quarter $t - 1$; (2) a recent news proxy, namely, the sum of the z -surprises over quarters $t - 4$ through $t - 1$; (3) a *Value * GoodNews* interaction term; and (4) a *Glamour * BadNews* interaction term. The price level and news variables are continuous, and are standardized so as to have zero mean and unit standard deviation. The interaction terms are dummy variables. In the so-called “ 2×2 sort,” *Value * GoodNews* takes the value of 1 if and only if the price level is below the median value and the news proxy is above the median value for that quarter. Similarly, *Glamour * BadNews* takes the value of 1 if and only if the price level is above the median value and the news proxy is below the median value for that quarter.

The second and third regressions are identical to the first, except that instead of forecasting returns over the next quarter, we forecast the (annualized) volatility and skewness of returns over the next four quarters, from t through $t + 3$. Note that we need to do the forecasting over more than one quarter simply because we cannot compute volatility and skewness using just a single quarterly return.

The second set of three regressions in each panel is similar, except that we use a “ 3×3 ” sort. This means that the *Value * GoodNews* dummy only takes the value of 1 if the price level is in the lowest one third of values and the news proxy is in the highest one third of values for that quarter; the *Glamour * BadNews* dummy is defined analogously. In other words, with the 3×3 sort, we assign a smaller and more extreme set of stocks to both the *Value * GoodNews* and *Glamour * BadNews* portfolios each quarter.

The results in Panel A for the no-learning benchmark case confirm what we establish above analytically. When predicting returns, the only variable that enters significantly is the news proxy, which attracts a coefficient of 0.0579, meaning that a one-standard deviation increase in the value of past z -surprises increases expected returns by 5.79% in annualized terms. There is no value-glamour effect, nor any interaction of value or glamour with the news proxy. When predicting volatility and skewness, none of the variables has a meaningful effect, that is, volatility and skewness are simply constants.

Things get more interesting when we move to Panels B and C, which cover the cases of model selection and model averaging. Because the basic thrust of the results is similar across these two panels, as well as across the 2×2 and 3×3 sorts, we focus our discussion on the model-selection case with a 2×2 sort. Consider first the regression that forecasts returns. The coefficient on the news proxy is similar to before, at 0.0760. But now, there is also an unconditional value-glamour effect, as seen in the coefficient on the price variable of -0.0407 . This implies that all else equal, a one-standard deviation increase in price reduces expected returns by 4.07% on an annualized basis.

Moreover, the *Value * GoodNews* and *Glamour * BadNews* terms attract significant coefficients of 0.0760 and -0.0757 , respectively. In other words,

controlling for the price level and past news, a stock that is in the *Value * GoodNews* quadrant has an additional expected return of 7.60% on an annualized basis, while a stock that is in the *Glamour * BadNews* quadrant has an expected return that is reduced by 7.57%. Again, these interaction effects are the key differentiating prediction of our theory.

Turning to the regression that forecasts volatility, we find that the only two significant predictor variables are the *Value * GoodNews* dummy and the *Glamour * BadNews* dummy, each of which attracts a positive coefficient of 0.0598. Thus, when a stock is in either of these quadrants, annualized volatility is increased by 5.98 percentage points. As we have seen above, this is because the potential for learning is elevated in these situations.

With respect to skewness, *Value * GoodNews* forecasts positive skewness, and *Glamour * BadNews* forecasts negative skewness, as anticipated in our intuitive discussion. In addition, the price level has a negative impact on future skewness—this is the “bubble-popping” effect mentioned above—while the news proxy has a positive impact.

In addition to the results shown in Table I, we also examine in the linear setting an alternative “rational-learning” benchmark. In this variant, investors update just as in the model-averaging case of Panel C, but the objective reality is that dividends are either generated by the simple A model or by the simple B model. In other words, investors’ perception of the environment now coincides with objective reality, so they can be thought of as standard rational Bayesians.

Volatility is stochastic in this setting, since the intensity of learning varies over time. However, none of the distinctive predictions that obtain in Panels B and C emerge with fully rational learning. Instead, we get an outcome that exactly mirrors Panel A: Neither returns, nor volatility, nor skewness is at all forecastable based on the value-glamour proxy, the news proxy, or any of their interactions. We therefore conclude that rational learning per se is not sufficient to generate the effects that we emphasize, even those for volatility, and that these effects are attributable to our particular rendition of the learning process.

C. Simulation Results: Log-Linear Dividend Specification

Table II presents the simulation results for the log-linear dividend specification. The format is identical to Table I, with the following exceptions. First, we omit the panel corresponding to the no-learning benchmark, and show only the model-selection and model-averaging cases.¹⁹ Second, all returns are in percentage terms, rather than in dollars. And third, when we compute skewness, this now refers to the skewness of log returns.²⁰

¹⁹ It turns out that the no-learning benchmark is not quite as clean in the log-linear case due to second-order Jensen’s inequality effects that arise. In particular, many of the regression coefficients that were almost exactly zero in the linear no-learning case are now statistically different from zero, albeit still small in economic terms.

²⁰ This is natural, since absent learning log returns should be symmetrically distributed in this log-linear setting.

Table II
Simulation Results for the Log-Linear Model

We simulate the dividends for a cross section of 2,500 stocks for 100 quarters. The variances of the shocks are set to $v_a = v_b = v_\varepsilon = 0.045$. The autocorrelation coefficient of the processes A_t and B_t is set to $\rho = 0.97$. The quarterly interest rate is set to $r = 0.015$. The Markov transition parameters $\pi_A = \pi_B = 1$ and $h = 0.051$. We use the log-linear pricing model to generate stock prices for the three cases of No Learning, Model Selection, and Model Averaging. For each of these three cases, we run three sets of Fama–MacBeth (1973) regressions, to forecast stock returns, return volatility, and return skewness. The dependent variables include the following. RET is the annualized return for quarter t . VOL is the annualized volatility calculated using four quarters of returns (t to $t + 3$). SKEW is the skewness coefficient for log returns, calculated using four quarters of returns (t to $t + 3$). The independent variables include a constant term (not reported) and the following. $Price(t - 1)$ is the lagged stock price, normalized to have zero mean and unit standard deviation. $News(t - 4, t - 1)$ is the cumulative z -surprise over the previous four quarters, normalized to have zero mean and unit standard deviation. For the 2×2 sort, $Value * GoodNews$ is a dummy variable that equals 1 if $Price(t - 1)$ is below the median value and $News(t - 4, t - 1)$ is above the median value for that quarter, while $Glamour * BadNews$ is a dummy variable that equals 1 if $Price(t - 1)$ is above the median value and $News(t - 4, t - 1)$ is below the median value for that quarter. For the 3×3 sort, $Value * GoodNews$ is a dummy variable that equals 1 if $Price(t - 1)$ is in the lowest one third of values and $News(t - 4, t - 1)$ is in the top one third of values for that quarter, while $Glamour * BadNews$ is a dummy variable that equals 1 if $Price(t - 1)$ is in the top one third of values and $News(t - 4, t - 1)$ is in the bottom one third of values for that quarter. For each simulation, we recover and save the Fama–MacBeth regression coefficients (the time-series average of the cross-sectional regression coefficients). We run 100 simulations and then take the average of the Fama–MacBeth coefficients, which are reported in the panels below. The standard errors are calculated simply as the standard deviation of these coefficients across simulations, divided by the square root of 100, and the associated t -statistics are reported in parentheses.

Panel A: Model Selection (Unconditional Annualized Stock Return Volatility of 27.8%)						
	2 × 2 Sort			3 × 3 Sort		
	RET	VOL	SKEW	RET	VOL	SKEW
<i>Value * GoodNews</i>	0.0712 (162.54)	0.0517 (205.70)	0.1014 (160.04)	0.0589 (70.55)	0.0754 (147.85)	0.0309 (32.25)
<i>Glamour * BadNews</i>	−0.0877 (205.44)	0.0528 (259.30)	−0.1021 (163.47)	−0.0376 (61.09)	0.0910 (243.84)	−0.0111 (12.03)
<i>Price(t − 1)</i>	−0.0140 (55.18)	0.0156 (128.92)	−0.0025 (9.59)	−0.0330 (133.71)	0.0114 (112.46)	−0.0366 (147.42)
<i>News(t − 4, t − 1)</i>	0.0644 (326.37)	0.0149 (90.40)	0.0339 (140.75)	0.0809 (407.61)	0.0170 (116.35)	0.0649 (248.95)

(continued)

Table II—Continued

Panel B: Model Averaging (Unconditional Annualized Stock Return Volatility of 28.2%)						
	2 × 2 Sort			3 × 3 Sort		
	RET	VOL	SKEW	RET	VOL	SKEW
<i>Value * GoodNews</i>	0.0772 (226.68)	0.0535 (199.56)	0.0706 (126.38)	0.0904 (99.65)	0.0821 (129.03)	0.0550 (72.79)
<i>Glamour * BadNews</i>	−0.0952 (224.10)	0.0596 (266.76)	−0.0669 (100.58)	−0.0741 (126.82)	0.0966 (248.52)	−0.0325 (36.51)
<i>Price(t − 1)</i>	−0.0121 (48.83)	0.0166 (121.81)	−0.0047 (17.54)	−0.0251 (104.11)	0.0133 (122.36)	−0.0205 (92.54)
<i>News(t − 4, t − 1)</i>	0.0525 (262.25)	0.0131 (72.72)	0.0168 (66.32)	0.0625 (285.21)	0.0143 (92.14)	0.0300 (121.15)

The qualitative results run closely parallel to those in Table I, and the economic magnitudes are generally similar. As before, consider the model-selection case with a 2×2 sort as a concrete example. Now when forecasting returns, the *Value * GoodNews* and *Glamour * BadNews* terms attract coefficients of 0.0712 and -0.0877 , respectively. When forecasting volatility, the corresponding coefficients are 0.0517 and 0.0528. Again, these would seem to be economically interesting magnitudes.

Finally, we should underscore that for the parameter values used in Table II, the model-selection case generates an unconditional annualized volatility of 27.8%, while the model-averaging case generates a volatility of 28.2%, both of which are realistic values for individual stocks. Thus, it appears that we can obtain economically interesting predictions without having to crank up the underlying variances in our model to implausible levels.

D. Empirical Results

Tables I and II embody the quantitative predictions of our theory. In Table III, we investigate these predictions empirically. Our empirical analysis is motivated in part by the observation that, in spite of the enormous literature on value-glamour effects, momentum, and post-earnings announcement drift, little work focuses on the interaction effects that are at the heart of our theory. The two exceptions that we are aware of are Asness (1997) and Swaminathan and Lee (2000), both of which we discuss further below. In any event, it would seem that there is room for much more work in this area, and our efforts here should be thought of as just a brief first cut.

We use CRSP stock return data and earnings announcement dates from Compustat over the period 1971 to 2004 to create a direct empirical analog to Table II.²¹ Our methodology is as follows. First, in place of the “price” variable in the simulations, we use the log of the market-to-book ratio, $\log(M/B)$. As in the simulations, this variable is normalized to have zero mean and unit standard deviation in any given cross section so as to make the magnitudes of the empirical and simulated coefficients directly comparable.

Second, in place of the “news” variable, we use the sum of the earnings announcement returns from the prior four quarters, with each return based on the 3-day interval (-1 to 1) around the announcement. Again, this variable is normalized to have zero mean and unit standard deviation in any cross section. As we stress above, announcement returns are the closest analog to the z -surprises that we use in the simulations, since absent a paradigm shift, the stock return at the moment of a dividend realization is exactly equal to the z -surprise.

With these two proxies in hand, we can define the *Value * GoodNews* and *Glamour * BadNews* dummies exactly as before for either the 2×2 or

²¹ Our sample includes all firms for which we have data on returns and market capitalization from CRSP, and data on earnings dates and book value from Compustat. We also require that book value be positive. The earnings dates are only available from Compustat beginning in the 1970s, which explains our sample period.

Table III
Empirical Results

Using CRSP stock-return data and earnings dates from Compustat for the period 1971 to 2004, we run the empirical analogs to the forecasting regressions described in Table II. We run three sets of Fama-MacBeth (1973) regressions, to forecast stock returns, return volatility, and return skewness. The dependent variables include the following. RET is the annualized return for quarter t . VOL is the annualized volatility for quarter t , calculated using daily returns. SKEW is the skewness coefficient for log returns for quarter t , calculated using daily returns. The independent variables include a constant term (not reported) and the following. $\text{Log}(M/B)$ is the lag of the log market-to-book ratio, normalized to have zero mean and unit standard deviation. $\text{News}(t - 4, t - 1)$ is the sum of the earnings announcement date returns (the average return from day -1 to day 1) in the previous four quarters, normalized to have zero mean and unit standard deviation. For the 2×2 sort, $\text{Value} * \text{GoodNews}$ is a dummy variable that equals 1 if $\text{Log}(M/B)$ is below the median value and $\text{News}(t - 4, t - 1)$ is above the median value for that quarter, while $\text{Glamour} * \text{BadNews}$ is a dummy variable that equals 1 if $\text{Log}(M/B)$ is above the median value and $\text{News}(t - 4, t - 1)$ is below the median value for that quarter. For the 3×3 sort, $\text{Value} * \text{GoodNews}$ is a dummy variable that equals one if $\text{Log}(M/B)$ is in the lowest one third of values and $\text{News}(t - 4, t - 1)$ is in the top one third of values for that quarter, while $\text{Glamour} * \text{BadNews}$ is a dummy variable that equals 1 if $\text{Log}(M/B)$ is in the top one third of values and $\text{News}(t - 4, t - 1)$ is in the bottom one third of values for that quarter. We report Fama-MacBeth regression coefficients, along with t -statistics that are based on Newey-West (1987) standard errors with four lags.

	2 × 2 Sort			3 × 3 Sort		
	RET	VOL	SKEW	RET	VOL	SKEW
<i>Value * GoodNews</i>	0.0205 (2.57)	-0.0118 (3.36)	0.0462 (8.25)	0.0454 (4.31)	0.0824 (8.98)	0.0769 (7.85)
<i>Glamour * BadNews</i>	-0.0255 (2.17)	-0.0067 (1.27)	0.0029 (0.17)	-0.0237 (2.07)	0.0424 (8.33)	0.0053 (0.32)
<i>Log(M/B)</i>	-0.0140 (1.31)	-0.0405 (6.87)	-0.0617 (11.80)	-0.0138 (1.24)	-0.0351 (6.22)	-0.0630 (11.96)
<i>News(t - 4, t - 1)</i>	0.0328 (5.07)	-0.0140 (4.09)	-0.0027 (0.42)	0.0321 (5.33)	-0.0203 (5.84)	-0.0019 (0.29)

3×3 sort. Finally, we run Fama and MacBeth (1973) regressions to forecast returns, volatility, and skewness based on the four predictors, just as in the simulations.²²

The results for returns line up remarkably well with our theoretical predictions. As would be expected based on previous research, the coefficient on $\log(M/B)$ is negative, and the coefficient on the news variable is positive. More strikingly from the perspective of our theory, the coefficient on $Value * GoodNews$ is significantly positive, while that on $Glamour * BadNews$ is significantly negative. This is true in both the 2×2 and 3×3 sorts.²³ The economic magnitudes are also in the same ballpark as, albeit somewhat smaller than, those from the log-linear simulations in Table II. In the 2×2 sort, the coefficient on $Value * GoodNews$ is 0.0205, while that on $Glamour * BadNews$ is -0.0255 . In the 3×3 sort, the corresponding numbers are 0.0454 and -0.0237 .

The results for volatility and skewness are more mixed. In the 3×3 sort, our theoretical predictions for volatility emerge strongly, with coefficients on $Value * GoodNews$ and $Glamour * BadNews$ of 0.0824 and 0.0424, respectively. But in the 2×2 sort, the coefficients on these interaction terms are much smaller, and of the wrong sign. In the skewness regressions, the coefficient on $Value * GoodNews$ is significantly positive, as predicted, in both the 2×2 and 3×3 sorts. But the coefficient on $Glamour * BadNews$ is very close to zero in both cases. Finally, consistent with both our theory and with previous empirical work by Chen, Hong, and Stein (2001), skewness is significantly more negative for high market-to-book stocks.

Overall, we draw the following conclusions from the work reported in this section. First, when calibrated with realistic parameter values, our theory delivers quantitative predictions that are of an economically interesting order of magnitude. In other words, the conditional variation in expected returns, volatility, and skewness generated by the theory is of first-order importance relative to the unconditional values of these moments. Second, the directional predictions of the theory for expected returns (most notably, the novel predictions regarding the effects of the interaction terms $Value * GoodNews$ and $Glamour * BadNews$)

²² One difference between the empirical setting and the simulations is that in the former, we can take advantage of daily data to more precisely estimate volatility and skewness. This is what we do in Table III, where volatility and skewness are estimated based on one quarter's worth of daily returns. However, we get similar results if instead we estimate volatility and skewness based on four quarters' worth of quarterly returns, as in the simulations.

²³ Swaminathan and Lee (2000) present closely related evidence, using double sorts rather than Fama–MacBeth regressions. Using data from 1974 to 1995, they do a 5×5 sort of stocks along two dimensions: book-to-market and earnings surprises. In the most negative earnings surprise quintile, glamour stocks (i.e., those in the lowest quintile of book-to-market) underperform moderately priced stocks (those in the middle quintile of book-to-market) by 4.71% per year. In contrast, in the highest earnings surprise quintile, the corresponding underperformance figure for glamour stocks is only 0.83% per year. With value stocks, the picture is reversed: They outperform moderately priced stocks by more when earnings surprises are in the upper quintile as opposed to the lower quintile, by 4.78% versus 1.55%. Also related are the findings of Asness (1997), who uses double sorts to study the interaction of book-to-market and price momentum.

are uniformly supported by the data. The theory also seems to have some explanatory power for movements in volatility and skewness, though not all of the predictions for these two moments come through as unambiguously.

IV. Revisionism: Equity Analysts and Amazon.com

In addition to its quantitative predictions for various moments of stock returns, our theory also implies the existence of a kind of *revisionism*: When there are paradigm shifts, investors will tend to look back at previously available public information and to draw different inferences from it than they had before. In an earlier version of this paper (Hong and Stein (2003b)), we illustrate the phenomenon of revisionism with a detailed account of equity analysts' reports on Amazon.com over the period 1997 to 2002, focusing both on the models that analysts use to arrive at their valuations for Amazon, and on how these models change over time. Here we just provide a brief summary of the narrative, and refer the interested reader to the working paper for details from the individual analyst reports.

In the period from its IPO in May 1997 up through its stock price peak in December of 1999, analysts offering valuations for Amazon repeatedly stressed its long-run revenue growth potential. At the same time, they explicitly dismissed the fact that Amazon's gross margins were much lower than those of its closest off-line retailing peers like Barnes & Noble. In fact, several analysts made a point of arguing that Barnes & Noble was the wrong analogy to draw, and that Amazon should be viewed as a fundamentally different type of business.

After a disappointing Christmas season in 1999, when Amazon's sales fell below expectations and the stock price began to drop precipitously, there appears to have been an abrupt shift in perspective. Many analysts began to point out the similarities between Amazon and off-line retailers, and started to emphasize gross margins in making their forecasts and recommendations. Indeed, a number of their post-1999 reports gave a lot of play to unfavorable data on Amazon's margins that had already been widely available for some time. And strikingly, some now used this stale data to justify downgrading the stock. This is just the sort of revisionism that our theory suggests.

V. Related Work

A large literature in game theory examines the implications of learning by less than fully rational agents.²⁴ While we share some of the same behavioral premises as this work, its goals are very different from ours, as for the most part, it seeks to understand the extent to which learning can *undo* the effects of agents' cognitive limitations.²⁵ For example, a commonly studied question in

²⁴ Early contributions to the learning-in-games literature include Robinson (1951), Miyasawa (1961), and Shapley (1964). For a survey of more recent work, see Fudenberg and Levine (1998).

²⁵ A similar comment can be made about the literature that asks whether learning by boundedly rational agents leads to convergence to rational expectations equilibria. See, for example, Cyert and DeGroot (1974), Blume, Bray, and Easley (1982), and Bray and Savin (1986).

this literature is whether learning will in the long run lead to convergence to Nash equilibrium.

Perhaps the paper that is closest to ours is that of Barberis et al. (1998), hereafter BSV.²⁶ As we do, BSV consider agents who attempt to learn, but who are restricted to updating over a class of incorrect models. In their setting, the models are specifically about the persistence of the earnings process—in one model shocks to earnings growth are relatively permanent, while in another model these shocks are more temporary in nature.²⁷ BSV's conclusions about under- and overreaction to earnings news then follow directly from the mistakes that agents make in estimating persistence.

In our theory, the notion of a model is considerably more abstract: A model is *any construct* that implies that one sort of information is more useful for forecasting than another. Thus, a model can be a metaphor such as “Amazon is just another Barnes & Noble,” which might imply that it is particularly important to study Amazon's gross margins. Or alternatively, a model can be “Company X seems a lot like Tyco,” which might suggest looking especially carefully at those footnotes in Company X's annual report where relocation loans to executives are disclosed. We view it as a strength of our approach that we are able to obtain a wide range of empirical implications without having to spell out such details.

The representative-agent/model-selection version of our theory is also reminiscent of Mullainathan's (2000) work on categorization. Indeed, our notion that individual agents practice model selection, instead of Bayesian model averaging, is essentially the same as Mullainathan's treatment of categorization. In spite of this apparent similarity, however, it is important to reiterate that our main empirical predictions do not follow from a discrete category-switching mechanism as in Mullainathan (2000), but rather from the fact that agents restrict their updating to the class of simple models, which in turn enables an econometrician to forecast variations in the intensity of learning over time.

VI. Conclusions

This paper can be seen as an attempt to integrate learning considerations into a behavioral setting in which agents are predisposed to using overly simplified forecasting models. The key assumption underlying our approach is that

²⁶ Other recent papers on the effects of learning for asset prices include Timmerman (1993), Wang (1993), Veronesi (1999), and Lewellen and Shanken (2002). In contrast to our setting or that of BSV, these papers consider a rational expectations setting and look at how learning about a hidden and time-varying growth rate for dividends leads to stock market predictability and excess volatility.

²⁷ In BSV, agents put zero weight on the model with the correct persistence parameter. One might argue that this assumption is hard to motivate, since the correct model is no more complicated or unnatural than the incorrect models that agents entertain. By contrast, in our setting the correct multivariate model *is* more complicated than the simple univariate models that agents actually update over.

agents update only over the class of simple models, placing zero weight on the correct, more complicated model of the world. As we demonstrate, this assumption yields a fairly rich set of empirical implications, many of which are supported in the data. Moreover, these implications seem to be robust to aggregation. That is, they come through either when there is a single representative agent who practices model selection, or when there is a market comprising heterogeneous agents, in which case the market can be said to practice a form of model averaging.

Appendix A. Derivation of Equation (14)

First, observe that

$$\text{cov}(R_t, P_{t-1}) = E(R_t P_{t-1}) - E(R_t)E(P_{t-1}). \quad (\text{A1})$$

Similarly,

$$\begin{aligned} \text{cov}(R_t, P_{t-1} | \text{shift}) &= E(R_t P_{t-1} | \text{shift}) \\ &\quad - E(R_t | \text{shift})E(P_{t-1} | \text{shift}), \end{aligned} \quad (\text{A2})$$

and

$$\begin{aligned} \text{cov}(R_t, P_{t-1} | \text{no shift}) &= E(R_t P_{t-1} | \text{no shift}) \\ &\quad - E(R_t | \text{no shift})E(P_{t-1} | \text{no shift}). \end{aligned} \quad (\text{A3})$$

We can also decompose $E(R_t P_{t-1})$ as

$$\begin{aligned} E(R_t P_{t-1}) &= E(R_t P_{t-1} | \text{shift}) \Pr(\text{shift}) \\ &\quad + E(R_t P_{t-1} | \text{no shift}) \Pr(\text{no shift}). \end{aligned} \quad (\text{A4})$$

Therefore, to establish equation (14), it suffices to prove that

$$E(P_{t-1} | \text{shift}) = E(P_{t-1}) = E(P_{t-1} | \text{no shift}) = 0. \quad (\text{A5})$$

For $E(P_{t-1} | \text{shift})$ we can write

$$\begin{aligned} E(P_{t-1} | \text{shift}) &= \int P_{t-1} f(P_{t-1} | \text{shift}) dP_{t-1} \\ &= \frac{1}{\Pr(\text{shift})} \int P_{t-1} \Pr(\text{shift} | P_{t-1}) f(P_{t-1}) dP_{t-1}, \end{aligned} \quad (\text{A6})$$

where the latter equality follows from an analog to Bayes's rule (a detailed proof of which is available upon request). Next, note that

$$\Pr(\text{shift} | P_{t-1} = x) = \Pr(\text{shift} | P_{t-1} = -x). \quad (\text{A7})$$

This property holds because of the symmetry of the normal learning process in equation (5) around zero. We also know that the unconditional distribution

$f(P_{t-1})$ is symmetric around zero. Therefore, it follows that $E(P_{t-1} | \text{shift}) = 0$, since for any function $g(x)$ that is symmetric around zero, $\int x g(x) dx = 0$. Identical logic establishes that $E(P_{t-1} | \text{noshift}) = 0$, and hence that $E(P_{t-1}) = 0$.

Appendix B: Stock Prices in the Log-Linear Case

We assume that dividends follow the process $D_t = \exp(A_t + B_t + \varepsilon_t)$, where $A_t = \rho A_{t-1} + a_t$, $B_t = \rho B_{t-1} + b_t$, $a_t \sim N(0, v_a)$, $b_t \sim N(0, v_b)$, and $\varepsilon_t \sim N(0, v_\varepsilon)$. We begin by calculating the stock price for an investor who understands the true model. This is given in Proposition A1, which is an application of the main result in Ang and Liu (2004).

PROPOSITION A1: *When $r > 0$, the rational stock valuation V_t^R for an investor who understands the true model is*

$$V_t^R = \sum_{s=1}^{\infty} \exp\left(-rs + \rho^{s-1}(A_{t+1} + B_{t+1}) + \frac{1}{2}(v_a + v_b) \frac{1 - \rho^{2(s-1)}}{1 - \rho^2} + \frac{1}{2}v_\varepsilon\right). \quad (\text{B1})$$

Proof: Observe that the rational stock valuation is simply the expected present value of future dividends (assuming the true dividend process)

$$V_t^R = E_t \left[\sum_{s=1}^{\infty} \exp(-rs) D_{t+s} \right] = \sum_{s=1}^{\infty} \exp(-rs) E_t \exp(A_{t+s} + B_{t+s} + \varepsilon_{t+s}). \quad (\text{B2})$$

We expand the three terms inside the second exponential function of the rational stock valuation as follows

$$\begin{aligned} & A_{t+s} + B_{t+s} + \varepsilon_{t+s} \\ &= \rho(A_{t+s-1} + B_{t+s-1}) + a_{t+s} + b_{t+s} + \varepsilon_{t+s} \\ &= \rho^2(A_{t+s-2} + B_{t+s-2}) + \rho(a_{t+s-1} + b_{t+s-1}) + a_{t+s} + b_{t+s} + \varepsilon_{t+s} \\ &= \dots \\ &= \rho^{s-1}(A_{t+1} + B_{t+1}) + \rho^{s-2}(a_{t+2} + b_{t+2}) \\ &\quad + \dots + \rho(a_{t+s-1} + b_{t+s-1}) + a_{t+s} + b_{t+s} + \varepsilon_{t+s}. \end{aligned} \quad (\text{B3})$$

Substituting this expansion inside the exponential function and taking expectations gives us

$$\begin{aligned}
 & E_t \exp(A_{t+s} + B_{t+s} + \varepsilon_{t+s}) \\
 &= \exp \left[\rho^{s-1}(A_{t+1} + B_{t+1}) + \frac{1}{2}(v_a + v_b) \sum_{n=0}^{s-2} \rho^{2n} + \frac{1}{2}v_\varepsilon \right] \\
 &= \exp \left[\rho^{s-1}(A_{t+1} + B_{t+1}) + \frac{1}{2}(v_a + v_b) \frac{1 - \rho^{2(s-1)}}{1 - \rho^2} + \frac{1}{2}v_\varepsilon \right]. \quad (\text{B4})
 \end{aligned}$$

Substituting the expression in (B4) back into the rational valuation formula above yields

$$V_t^R = \sum_{s=1}^{\infty} \exp \left(-rs + \rho^{s-1}(A_{t+1} + B_{t+1}) + \frac{1}{2}(v_a + v_b) \frac{1 - \rho^{2(s-1)}}{1 - \rho^2} + \frac{1}{2}v_\varepsilon \right). \quad (\text{B5})$$

Q.E.D.

The rational stock valuation depends on an infinite sum. Ang and Liu (2004) point out that when $r > 0$, successive terms in the summation decrease exponentially fast and V_t^R can be approximated via the first m terms in the summation for some large m . In our simulations, we set $m = 1,000$.

With this rational stock valuation in hand, we can then work out the prices for the cases of no learning, model selection, and model averaging. For an investor who uses model A and ignores signal B , his valuation is

$$V_t^A = \sum_{s=1}^{\infty} \exp \left(-rs + \rho^{s-1}A_{t+1} + \frac{1}{2}v_a \frac{1 - \rho^{2(s-1)}}{1 - \rho^2} + \frac{1}{2}v_\varepsilon \right), \quad (\text{B6})$$

which we derive by assuming that $D_t = \exp(A_t + \varepsilon_t)$, where $A_t = \rho A_{t-1} + a_t$, $a_t \sim N(0, v_a)$, and $\varepsilon_t \sim N(0, v_\varepsilon)$, and applying Proposition A1. Similarly, for an investor using model B, his valuation is

$$V_t^B = \sum_{s=1}^{\infty} \exp \left(-rs + \rho^{s-1}B_{t+1} + \frac{1}{2}v_b \frac{1 - \rho^{2(s-1)}}{1 - \rho^2} + \frac{1}{2}v_\varepsilon \right), \quad (\text{B7})$$

which we derive by assuming that $D_t = \exp(B_t + \varepsilon_t)$, where $B_t = \rho B_{t-1} + b_t$, $b_t \sim N(0, v_b)$, and $\varepsilon_t \sim N(0, v_\varepsilon)$, and applying Proposition A1.

We determine the stock price at time t for the three different cases (no learning, model selection, and model averaging) in the following way. In the no-learning case, we assume the investor sticks to model A and the stock price is given by

$$P_t = V_t^A. \quad (\text{B8})$$

Under the model-selection case, the stock price is determined by the current model

$$P_t = \begin{cases} V_t^A & \text{if model A} \\ V_t^B & \text{if model B} \end{cases}. \quad (\text{B9})$$

In the model-averaging case, the stock price is given by the average of the valuations under models A and B, weighted by the proportion of investors (p_t) using each model, that is,

$$P_t = p_t V_t^A + (1 - p_t) V_t^B. \quad (\text{B10})$$

In each of these cases, the stock return is calculated simply as

$$R_t = \frac{P_t + D_t}{P_{t-1}} - 1, \quad (\text{B11})$$

where P_t is given by one of the three cases (no learning, model selection, and model averaging) and D_t follows the true log-linear specification given above.

Appendix C: Calibration

We now provide calculations of the first-order autocorrelation of log dividends that is useful in the calibration of our model. We set $v_a = v_b = v_\varepsilon = v$ and given that $\log D_t = A_t + B_t + \varepsilon_t$, we compute the variance of this process as

$$V_0 = \text{var}(\log D_t) = \frac{2v}{1 - \rho^2} + v = \left(\frac{2}{1 - \rho^2} + 1 \right) v. \quad (\text{C1})$$

The first-order auto-covariance of this process is

$$\begin{aligned} V_1 &= \text{cov}(\log D_t, \log D_{t-1}) \\ &= \text{cov}(A_t + B_t + \varepsilon_t, A_{t-1} + B_{t-1} + \varepsilon_{t-1}) \\ &= \text{cov}(A_t, A_{t-1}) + \text{cov}(B_t, B_{t-1}) \\ &= \frac{2\rho}{1 - \rho^2} v. \end{aligned} \quad (\text{C2})$$

The first-order autocorrelation of the log dividends implied by the log-linear specification is

$$\frac{V_1}{V_0} = \frac{\frac{2\rho}{1 - \rho^2}}{\frac{2}{1 - \rho^2} + 1} = \frac{2\rho}{3 - \rho^2}. \quad (\text{C3})$$

The parameter ρ uniquely determines the serial correlation of log dividends. When $\rho = 0.97$, the implied first-order autocorrelation is 0.94, roughly

matching the first-order autocorrelation for S&P 500 quarterly log real operating earnings during the period of 1988 to 2004 (which we calculate to be 0.96).²⁸

REFERENCES

- Abelson, Robert P., 1978, *Scripts*, Invited address to the Midwestern Psychological Association, Chicago.
- Abreu, Dilip, and Markus Brunnermeier, 2003, Bubbles and crashes, *Econometrica* 71, 173–204.
- Ang, Andrew, and Jun Liu, 2004, How to discount cashflows with time-varying expected returns, *Journal of Finance* 59, 2745–2783.
- Asness, Clifford S., 1997, The interaction of value and momentum strategies, *Financial Analysts Journal* 53, 29–39.
- Barberis, Nicholas, Andrei Shleifer, and Robert W. Vishny, 1998, A model of investor sentiment, *Journal of Financial Economics* 49, 307–343.
- Bernard, Victor L., and J. Thomas, 1989, Post-earnings announcement drift: Delayed price response or risk premium? *Journal of Accounting Research* 27, 1–48.
- Bernard, Victor L., and J. Thomas, 1990, Evidence that stock prices do not fully reflect the implications of current earnings for future earnings, *Journal of Accounting and Economics* 13, 305–340.
- Bernstein Quantitative Handbook, 2004, (Sanford C. Bernstein & Co. LLC, a subsidiary of Alliance Capital Management LP, New York, NY).
- Blanchard, Olivier J., and Mark W. Watson, 1982, Bubbles, rational expectations, and financial markets, in Paul Wachtel, ed. *Crises in Economic and Financial Structure* (Lexington Books, Lexington, MA).
- Blume, Lawrence, Margaret M. Bray, and David Easley, 1982, Introduction to stability of rational expectations equilibrium, *Journal of Economic Theory* 26, 313–317.
- Bray, Margaret M., and Nathan E. Savin, 1986, Rational expectations equilibrium, learning and model specification, *Econometrica* 54, 1129–1160.
- Bruner, Jerome S., 1957, Going beyond the information given, in H. Gulber, and others, eds. *Contemporary Approaches to Cognition* (Harvard University Press, Cambridge, MA).
- Bruner, Jerome S., and Leo Postman 1949, On the perception of incongruity: A paradigm, *Journal of Personality* 18, 206–223.
- Chen, Joseph, Harrison Hong, and Jeremy C. Stein, 2001, Forecasting crashes: Trading volume, past returns and conditional skewness in stock prices, *Journal of Financial Economics* 61, 345–381.
- Cyert, Richard M., and Morris H. DeGroot, 1974, Rational expectations and Bayesian analysis, *Journal of Political Economy* 82, 521–536.
- Daniel, Kent D., David Hirshleifer, and Avanidhar Subrahmanyam, 1998, Investor psychology and security market under- and over-reactions, *Journal of Finance* 53, 1839–1885.
- Della Vigna, Stefano, and Joshua Pollet, 2004, Attention, demographics and the stock market, Working paper, Harvard University.
- Fama, Eugene F., and Kenneth R. French, 1992, The cross-section of expected returns, *Journal of Finance* 47, 427–465.
- Fama, Eugene F., and James D. MacBeth, 1973, Risk, return, and equilibrium: Empirical tests, *Journal of Political Economy*, 81, 607–636.

²⁸ We download the operating earnings data from Standard & Poor's website <http://www2.standardandpoors.com/spf/xls/index/SP500EPSEST.XLS>. The earnings data are deflated using the CPI before taking logs. We obtain CPI data from the Bureau of Labor Statistics website, <ftp://ftp.bls.gov/pub/special.requests/cpi/cpiat.txt>. We use the monthly CPI series for all urban consumers (series CPI-U), employing the 3-month average of CPI-U within a quarter to deflate the operating earnings corresponding to that quarter.

- Fiske, Susan T., and Shelley E. Taylor, 1991, *Social Cognition*, 2nd ed. (McGraw Hill, Inc., New York, NY).
- Fudenberg, Drew, and David K. Levine, 1998, Learning in games: Where do we stand? *European Economic Review* 42, 631–639.
- Harrison, J. Michael, and David M. Kreps, 1978, Speculative investor behavior in a stock market with heterogeneous expectations, *Quarterly Journal of Economics* 93, 323–336.
- Hirshleifer, David, and Siew Hong Teoh, 2003, Limited attention, information disclosure and financial reporting, *Journal of Accounting and Economics* 36, 337–386.
- Hong, Harrison, and Sven Rady, 2002, Strategic trading and learning about liquidity, *Journal of Financial Markets* 5, 419–450.
- Hong, Harrison, and Jeremy C. Stein, 1999, A unified theory of underreaction, momentum trading and overreaction in asset markets, *Journal of Finance* 54, 2143–2184.
- Hong, Harrison, and Jeremy C. Stein, 2003a, Differences of opinion, short-sales constraints and market crashes, *Review of Financial Studies* 16, 487–525.
- Hong, Harrison, and Jeremy C. Stein, 2003b, Simple forecasts and paradigm shifts, NBER Working paper 10013.
- Jegadeesh, Narasimhan, and Sheridan Titman, 1993, Returns to buying winners and selling losers: Implications for stock market efficiency, *Journal of Finance* 48, 93–130.
- Kahneman, Daniel, and Amos Tversky, 1973, On the psychology of prediction, *Psychological Review* 80, 237–251.
- Kandel, Eugene, and Neil D. Pearson, 1995, Differential interpretation of public signals and trade in speculative markets, *Journal of Political Economy* 103, 831–872.
- Kuhn, Thomas S., 1962, *The Structure of Scientific Revolutions* (University of Chicago Press, Chicago, IL).
- Kyle, Albert S., and F. Albert Wang, 1997, Speculation duopoly with agreement to disagree: Can overconfidence survive the market test? *Journal of Finance* 52, 2073–2090.
- Lakonishok, Josef, Andrei Shleifer, and Robert W. Vishny, 1994, Contrarian investment, extrapolation and risk, *Journal of Finance* 49, 1541–1578.
- Lee, Charles, and Bhaskaran Swaminathan, 2000, Price momentum and trading volume, *Journal of Finance* 55, 2017–2069.
- Lewellen, Jonathan, and Jay Shanken, 2002, Learning, asset-pricing tests and market efficiency, *Journal of Finance* 57, 1113–1145.
- Lord, Charles, Lee Ross, and Mark R. Lepper, 1979, Biased assimilation and attitude polarization: The effects of prior theories on subsequently considered evidence, *Journal of Personality and Social Psychology* 37, 2098–2109.
- Miller, Edward, 1977, Risk, uncertainty, and divergence of opinion, *Journal of Finance* 32, 1151–1168.
- Miyasawa, Koichi, 1961, On the convergence of learning processes in a 2×2 non-zero sum game, Research memorandum No. 33, Princeton University.
- Morris, Stephen, 1996, Speculative investor behavior and learning, *Quarterly Journal of Economics* 111, 1111–1133.
- Mullainathan, Sendhil, 2000, Thinking through categories, Working paper, MIT.
- Newey, Whitney K., and Kenneth D. West, 1987, A simple positive-definite heteroskedasticity and autocorrelation consistent covariance matrix, *Econometrica* 55, 703–708.
- Nisbett, Richard, and Lee Ross, 1980, *Human Inference: Strategies and Shortcomings of Social Judgment* (Prentice-Hall, Englewood Cliffs, NJ).
- Odean, Terrance, 1998, Volume, volatility, price and profit when all traders are above average, *Journal of Finance* 53, 1887–1934.
- O'Hara, Maureen, 1995, *Market Microstructure Theory* (Blackwell Publishers, Cambridge, MA).
- Peng, Lin, and Wei Xiong, 2006, Investor attention, overconfidence and category learning, *Journal of Financial Economics* 80, 563–602.
- Rabin, Matthew, and Joel L. Schrag, 1999, First impressions matter: A model of confirmatory bias, *Quarterly Journal of Economics* 114, 37–82.
- Robinson, Julia, 1951, An iterative method of solving a game, *Annals of Mathematics* 24, 296–301.

- Schank, Roger C., and Robert P. Abelson, 1977, *Scripts, Plans, Goals and Understanding: An Introduction into Human Knowledge Structures* (Lawrence Erlbaum Associates, Hillsdale, NJ).
- Scheinkman, Jose, and Wei Xiong, 2003, Overconfidence and speculative bubbles, *Journal of Political Economy* 111, 1183–1219.
- Shapley, Lloyd S., 1964, Some topics in two-person games, in M. Drescher, L. S. Shapley, and A. W. Tucker eds. *Advances in Game Theory* (Princeton University Press, Princeton, NJ).
- Simon, Herbert A., 1982, *Models of Bounded Rationality: Behavioral Economics and Business Organizations*, vol. 2 (MIT Press, Cambridge, MA).
- Sims, Christopher A., 2003, Implications of rational inattention, *Journal of Monetary Economics* 50, 665–690.
- Swaminathan, Bhaskaran, and Charles Lee, 2000, Do stock prices overreact to earnings news? Working paper, Cornell University.
- Taylor, Shelley E., and Jennifer C. Crocker, 1980, Schematic bases of social information processing, in E. Tory Higgins, C. Peter Herman, and Mark P. Zanna, eds. *The Ontario Symposium on Personality and Social Psychology*, vol. 1 (Lawrence Erlbaum, Hillsdale, NJ).
- Timmermann, Allan, 1993, How learning in financial markets generates excess volatility and predictability in stock prices, *Quarterly Journal of Economics* 108, 1135–1145.
- Tversky, Amos, and Daniel Kahneman, 1973, Availability: A heuristic for judging frequency and availability, *Cognitive Psychology* 5, 207–232.
- Varian, Hal R., 1989, Differences of opinion in financial markets, in C. C. Stone, ed. *Financial Risk: Theory, Evidence, and Implications: Proceedings of the 11th Annual Economic Policy Conference of the Federal Reserve Bank of St. Louis* (Kluwer Academic Publishers, Boston, MA).
- Veronesi, Pietro, 1999, Stock market overreaction to bad news in good times: A rational expectations equilibrium model, *Review of Financial Studies* 12, 975–1007.
- Wang, Jiang, 1993, A model of intertemporal asset prices under asymmetric information, *Review of Economic Studies* 60, 249–282.