

Wealth Evolution and Distorted Financial Forecasts

Blake LeBaron *

Brandeis University

December, 2007

Revised: March 2008

Abstract

Evolutionary metaphors have been prominent in both economics and finance. They are often used as basic foundations for rational behavior and efficient markets. Theoretically, a mechanism which selects for rational investors actually requires many caveats, and is far from generic. This paper tests wealth based evolution in a simple, stylized agent-based financial market. The setup borrows extensively from current research in finance that considers optimal behavior with some amount of return predictability. The results confirm that with a homogeneous world of log utility investors wealth will converge onto optimal adaptive forecasting parameters. However, in the case of utility functions which differ from log, wealth selection alone converges to parameters which are economically far from the optimal forecast parameters. This serves as a strong reminder that wealth selection and utility maximization are not the same thing. Therefore, suboptimal financial forecasting strategies may be difficult to drive out of a market, and may even do quite well for some time.

*International Business School, Brandeis University, 415 South Street, Mailstop 32, Waltham, MA 02453 - 2728, blebaron@brandeis.edu, www.brandeis.edu/~blebaron. The author is also a faculty research fellow at the National Bureau of Economic Research.

1 Introduction

Evolution has always played an important background role in both finance and economics. Many researchers have taken comfort in thinking that irrational trading strategies, or less than profitable firms would eventually be removed from the market.¹ The theoretical backing for this strong defense of rationality is not as definitive as its proponents would have us think. Different situations require different restrictions on behaviors for convergence to rationality. This paper explores this question in a very standard financial forecasting test case where stock returns have some weak predictability. In this world wealth evolution can select for distorted predictors which are economically far from the optimal true probabilities. This demonstrates that some irrational forecasters may be very difficult to remove from a market, and may even thrive.

Three difficult problems are often given to counter the simple argument that only rational strategies combined with correct probabilities will survive in the long run in a financial setting. First, evolution is taking place in a dynamic environment where prices react to strategies, and strategies react to each other. Most papers on evolutionary markets concentrate on this point, and work hard to make sure prices are endogenous.² This paper will actually take prices as exogenous and ignore this issue. This is done to concentrate on two other issues, the fact that wealth selects for growth maximizing, and not utility maximizing strategies, and that wealth evolution alone can be quite slow when examined at reasonable parameter values.

The fact that wealth growth and utility maximization are not the same thing is well known in finance. It generated a large debate in the 60's and 70's about the normative case for holding portfolios that maximized the growth rate of wealth.³ This paper, and also the modern literature on growth optimality, looks at the positive question of which strategies survive in a wealth evolutionary dynamical system. In this situation the growth optimal portfolio plays an important role. It is often the strategy which survives in the long run.⁴ In different situations, for different preferences, it may or may not be a utility maximizing strategy. This paper contributes to this question by looking at the case where returns are predictable.

The final key problem which is mentioned in connection to growth optimal strategies is the speed of

¹The early comments on this are in Alchian (1950) and Friedman (1953).

²Blume & Easley (1990), Sandroni (2000), and Kogan, Ross & Wang (2006) are good examples of evolution with endogenous prices. DeLong, Shleifer, Summers & Waldmann (1991) is an early paper which also examines the evolutionary stability of irrational beliefs. Also, Cecchetti, Lam & Mark (2000) show that distorted beliefs can help to explain aggregate stock returns. Computational agent-based approaches are primarily concerned with price dynamics and the interactions between strategies occurring through prices. See LeBaron (2006) for a survey of this literature.

³This known as the growth optimal portfolio. See Samuelson (1971) and Hakansson (1971) for the original debate. Also, Kelley (1956) and Breiman (1961) provide the theoretical foundations. A nice summary of this is in Markowitz (1976).

⁴Various theoretical papers have reached similar conclusions in different frameworks. These include Blume & Easley (1990) and Blume & Easley (2006) which analyze utility maximizing strategies with prices set endogenously. The latter paper proves that in a complete markets world the convergence to true beliefs will occur regardless of preference parameters. However, the authors point out that in an incomplete markets world this convergence is not guaranteed. Evstigneev, Hens & Schenk-Hoppe (2006) look at an incomplete markets world with endogenous prices. In their framework the growth optimal strategy will dominate any other competing strategy in terms of acquiring all wealth in the long run.

convergence. Several papers have concluded that asymptotic results may only be relevant at extremely long horizons.⁵ This paper will address this issue indirectly, since many of the results are based on small sample computer simulations and not asymptotics.

The results in this paper are important for examining learning and strategy construction in finance, and economics in general. Learning can take many forms, and its connection to evolution is often unclear. Learning may come from deductive thinking and explicit utility maximization on the part of agents. A weaker form of learning is related to some general form of adaptation. Agents examine strategies, and if they see others performing better subject to some objective, they shift. They might perform these experiments themselves, or might look across a set of strategies currently in use. Adaptive learning is intuitively appealing, but difficult to implement. There are many ways to model this. Also, results can be very sensitive to the frequency with which adaptation takes place.⁶

These first two forms of learning should be viewed as active learning in which the agent is taking explicit actions to improve their outcomes. Wealth evolution should be viewed as passive. No one in these models is trying to improve on strategies. The only form of learning taking place comes from wealth shifting to the relatively successful strategies. There are two key issues which make wealth evolution important. First, there really are no questions about how to model this. Wealth is accumulated based on realized returns, and this gives a well defined dynamic for wealth shares. Second, wealth evolution, for all its faults, is somewhat foundational. It must be present in any sensible financial model even built off other learning mechanisms. Therefore, it's directions and biases are important to understand.

The modeling strategies used in this paper will draw heavily on the financial forecasting literature, and results on dynamic portfolio construction. Specifically, the underlying economic structure will be based both on Campbell & Viceira (1999) and forecasting rules motivated by state space models such as those used in Pastor & Stambaugh (2006). Section 2 will describe the modeling structure in detail. Section 3 presents the results. Section 4 performs some extensions and robustness checks, and section 5 concludes.

2 Model structure

2.1 Return Dynamics

The economy considered here is a partial equilibrium one where security prices are set exogenously, and are not influenced by changes in wealth. There are two assets in the market. A risk free asset which pays a

⁵See Rubinstein (1991) for some early tests directed at the normative side of this question. Also, Figlewski (1978) provides evidence in a market with endogenous price setting. More recently Berrada (2006) and Yan (2006) address this in terms of the speed of convergence of wealth fractions.

⁶For examples of this see LeBaron (2001).

fixed return, and a risky asset paying a stochastic return with a small predictable component. Returns will be generated at a weekly frequency, and all portfolio rebalancing decisions will be made on a weekly basis.

The parameters are calibrated to well known results from financial markets to look reasonable.⁷ The risk free return is given as R_f , with $r_f = \log(1 + R_f)$. The return on the single risky asset is given by R_t with $r_t = \log(1 + R_t)$.

The dynamics of r_t are given by

$$r_{t+1} = x_{t+1} + e_{t+1} \quad (1)$$

$$x_{t+1} = \mu + \rho(x_t - \mu) + \eta_{t+1}. \quad (2)$$

This representation follows Campbell & Viceira (1999), and is a reasonable benchmark for financial returns series showing some amount of predictability.⁸ It will diverge from much of the previous work in two important ways. First, it is assumed that x_t is unobserved. Other papers have considered x_t to be known, or try to connect it to various information variables available at time t . Moving to a framework where agents are trying to estimate x_t using predictive regressions is possible, but would complicate the initial simplicity of the learning dynamics used in this paper. The single hidden state variable maps directly into a very simple state-space representation. Another reason for avoiding predictive regressions is that there is very little agreement as to what the predictors should be, or how good they are for portfolio construction.⁹ Second, I will assume that the disturbances for the x_t process, η_t , and the return noise process, e_t are independent. This is a key difference from much of the work on financial prediction. In that literature this correlation is important, both in estimating, and in interpreting various forecasts. At the weekly frequencies used in this paper, it is not clear if this covariance is significantly different from zero for many predictors. Also, keeping this zero simplifies the learning portfolio construction process allowing for sharper interpretations of many of the results.

Certain aspects of the stochastic structure of r_t will be important for the framework. Both noise shocks, e_t and η_t , will be normally distributed, and are homoskedastic with variances given by σ_e and σ_η . The annualized values of these are given in table 1. Two other important features will be used in choosing parameters. First, the signal to noise ratio in returns series is small. Predictive regressions run at the annual frequency generally yield very small R^2 values, usually between zero and 10 percent. Reflecting this, the parameters are set so that the variance of x_t is 2 percent of the total return variance at the weekly frequency.

⁷See Campbell & Viceira (2002) for many examples.

⁸The predictability literature is extensive. Early examples include Keim & Stambaugh (1986), Campbell & Shiller (1988), and Fama & French (1988). See Ferson, Sarkissian & Simin (2003) and Stambaugh (1999) for further references.

⁹The structure here is the base case for Pastor & Stambaugh (2006) who analyze the problem of “imperfect” predictors, where the true expected return is not observed.

Table 2 reports a monte-carlo simulation of the return process, showing autocorrelations at frequencies of 1 to 4 weeks, and the R^2 for a simulated one year predictive regression. Values in parenthesis are standard deviations across 1000 monte-carlo runs. The relatively short sample lengths are chosen to correspond to those available in many financial time series. The annual prediction experiment assumes the investor knows x_t and regresses the next year's return on the current value. The simulations produce R^2 estimates which are approximately 10 percent with very large dispersion across the simulated cross section. We should expect these numbers to have a slight upward bias due to the fact that in this experiment it is assumed that investors know the value of x_t .

The value of ρ is set to 0.95. This reflects the large persistence believed to characterize many predictor variables. For example, Campbell & Viceira (2002) report a value of 0.957 for an estimate of the quarterly impact of lagged dividend price ratios on current ones. The value of 0.95 is probably slightly too small for weekly persistence, but there are several reasons for choosing this. First, x_t doesn't exactly represent dividend price ratios, but is a stand in for many different predictors. Second, the value of 0.95 is useful in the experiments to see if agents are able to discern between a stationary, and a nonstationary process for x_t . As an initial test, it seems reasonable to move this parameter farther away from 1.

This return calibration is only a loose approximation. The desire was to choose a set of parameters that generally replicated a broad set of features, rather than exactly fitting to a return series from any particular time period. Also, given that the expected return series is unobserved, the degrees of freedom in choosing some of these parameters is obviously large.

2.2 Wealth Evolution

Most of the experiments performed in this paper will concentrate on the evolution of wealth shares across traders. The objective is to find out in a pool of noninteracting strategies with an exogenous returns process, who in the end is left standing through simple compounding of wealth onto successful dynamic portfolios. Before detailing where these strategies will come from in terms of preferences, it is important to state that one can be somewhat agnostic in terms of preferences.¹⁰

Agent i 's strategy each period will be to invest $\alpha_{t,i}$ fraction of wealth in the risky asset, and $1 - \alpha_{t,i}$ fraction in the risk free. The portfolio return from t to $t + 1$ is therefore,

$$R_{t+1,i}^p = \alpha_{t,i}R_{t+1} + (1 - \alpha_{t,i})R_f. \quad (3)$$

¹⁰This follows the general approach taken in Evstigneev et al. (2006) where portfolios are stated as fractions of wealth invested in different assets. Where these strategies come from is not critical, since the authors are only interested in the properties of strategies which would eventually be selected for in terms of wealth.

The wealth share of agent i follows,

$$w_{t+1,i} = \frac{w_{t,i}(1 + R_{t+1,i}^p)}{\sum_{j=1}^N w_{t,j}(1 + R_{t+1,j}^p)}. \quad (4)$$

The dynamics of wealth depends on the realized distribution of returns, wealth shares at period t , and the portfolio strategies at period t , $\alpha_{t,i}$. It is important to add that another case fits easily into this framework. If agents consume a fixed fraction of wealth, λ , each period, then the wealth share dynamics would be given by,

$$w_{t+1,i} = \frac{(1 - \lambda)w_{t,i}(1 + R_{t+1,i}^p)}{\sum_{j=1}^N (1 - \lambda)w_{t,j}(1 + R_{t+1,j}^p)} \quad (5)$$

which is obviously the same.¹¹

2.3 Preferences and portfolio choices

Portfolio choices in the model are determined by a simple myopic power utility function in future wealth. The agent's portfolio problem corresponds to,

$$\max_{\alpha_{t,i}} \frac{E_t^i W_{t+1}^{1-\gamma}}{1-\gamma}, \quad (6)$$

$$st. \quad W_{t+1} = (1 + R_{t+1}^p)W_t. \quad (7)$$

Dropping out constant values known at time t , this becomes,

$$\max \frac{1}{1-\gamma} E_t^i (1 + R_{t,i}^p)^{1-\gamma}. \quad (8)$$

If portfolio returns were log normal, this could be transformed. Unfortunately, portfolio returns are not log normal. Campbell & Viceira (2002) show using a Taylor series approximation that the log portfolio return is approximated by,

$$r_{p,t+1} = r_f + \alpha_t(r_{t+1} - r_f) + (1/2)\alpha_t(1 - \alpha_t)\sigma_t^2 \quad (9)$$

where $r_{p,t} = \log(1 + R_t^p)$, $r_t = \log(1 + R_t)$, and $\sigma_t^2 = \text{var}(r_t)$. Assuming the return on the risky asset is log normal, then the approximate portfolio return is also log normal. This allows the use of the well known fact that for log normal random variables

$$\log(E(Y)) = E \log(Y) + (1/2)\sigma_Y^2. \quad (10)$$

¹¹It is well known in dynamic intertemporal consumption/portfolio choice problems that the c/w ratio is constant when the intertemporal elasticity of substitution is unity, (Giovannini & Weil (1989)).

Returning to the maximization problem in equation 8, taking logs of the expectation, and using 10 we get,

$$\max \frac{1}{1-\gamma} \log(E_t^i(1 + R_{t,i}^p)^{1-\gamma}). \quad (11)$$

$$\max \frac{1}{1-\gamma} E_t^i(1-\gamma) \log(1 + R_{t+1}^p) + (1/2)(1-\gamma)^2 \sigma_{r_p}^2 \quad (12)$$

$$\max E_t^i r_p + (1/2)(1-\gamma) \sigma_{r_p}^2. \quad (13)$$

Using the approximation in 9 gives,

$$\max_{\alpha_t} r_f + \alpha_t(E_t^i r_{t+1} - r_f) + (1/2)\alpha_t(1-\alpha_t)\sigma_t^2 + (1/2)(1-\gamma)\alpha_t^2\sigma_t^2. \quad (14)$$

Dropping the constant, r_f , and solving for the optimal holding, α , gives

$$\alpha_{t,i} = \frac{E_t^i(r_{t+1}) - r_f + \sigma_t^2/2}{\gamma\sigma_t^2}. \quad (15)$$

For most runs the portfolio weights will be restricted to $-0.5 < \alpha_{t,i} < 2$. This allows for some short selling, and some leverage. This range is designed to replicate a moderate risk taking hedge fund more than the average individual investor.

This simplified myopic portfolio strategy will be used throughout the paper. It is important to note that the fixed consumption, myopic strategy approach given here would be optimal in a standard intertemporal model for consumption portfolio choice subject to two key assumptions. First, the intertemporal elasticity of substitution would have to be unity to fix the consumption wealth ratio, and second, the correlation between the innovation to x_t and the return error e_t needs to be zero to eliminate hedging demands.¹²

2.4 Return forecasting

Given returns follow,

$$\begin{aligned} r_{t+1} &= x_{t+1} + e_{t+1} \\ x_{t+1} &= \mu + \rho(x_t - \mu) + \eta_{t+1} \end{aligned}$$

¹²See Campbell & Viceira (1999) for the basic framework. Hedging demands would only impose a constant shift on the optimal portfolio, so it is an interesting question how much of an impact this might have on the results.

and assuming that the state, or expected return component, x_t , is unobserved, the optimal forecasting strategy is a state space model. I will also assume throughout that the unconditional mean, μ , is known to all investors. This is a simple application of the Kalman filter. For completeness the simple learning algorithm is outlined here, but there are many descriptions of this in the economics, engineering, and statistics literatures.¹³

This is a very quick outline of the forecasting problem faced by agents. Let $\hat{x}_{t|t}$ be the forecast of x_t given time t information. The Kalman forecast of next period's state variable is given by

$$\hat{x}_{t+1|t} = \mu + \rho(\hat{x}_{t|t} - \mu). \quad (16)$$

Now define the forecast error conditioned on time t information by

$$p_t = E_t(x_t - \hat{x}_{t|t})^2, \quad (17)$$

and the one step ahead forecast error is given by

$$p_{t+1|t} = E_t(x_{t+1} - \hat{x}_{t+1|t})^2 = \rho^2 p_t + \sigma_\eta^2. \quad (18)$$

The key value in the Kalman filter is the gain level which is given by,

$$k_{t+1} = \frac{p_{t+1|t}}{p_{t+1|t} + \sigma_e^2}. \quad (19)$$

The gain level is used to update the x_{t+1} forecast, and its conditional error, when new information, r_{t+1} , arrives,

$$\begin{aligned} \hat{x}_{t+1|t+1} &= \hat{x}_{t+1|t} + k_{t+1}(r_{t+1} - \hat{x}_{t+1|t}) \\ p_{t+1} &= p_{t+1|t}(1 - k_{t+1}) \end{aligned}$$

The importance of k_t is clear. It tells agents how much weight to give current information as they try to forecast future returns. If the noise level in the system, σ_e^2 is high, then this weight will be low, since new returns are not very informative. As this noise level falls, the signal to noise ratio increases, and new returns are given progressively more weight in forecasts of the hidden expected return level.

¹³For information on state-space models in economic time series contexts see Hamilton (1994) or Harvey (1989). Also, it is easy to show that this model corresponds to an ARMA(1,1) which is the same as this state space representation. See Taylor (1980), Taylor (2005), and LeBaron (1992) for examples.

Kalman filters are implemented recursively, and the parameters are updated over time as in

$$p_t \rightarrow p_{t+1|t} \rightarrow k_{t+1} \rightarrow p_{t+1} \dots \quad (20)$$

It is easy to show in this simple system that this will converge to steady state levels which define a constant gain value, k . This gives a forecast of

$$\hat{x}_{t+1|t} = \mu + \rho(\hat{x}_{t|t} - \mu) \quad (21)$$

$$\hat{x}_{t+1|t} = \mu + \rho(\hat{x}_{t|t-1} - \mu) + \rho k(r_t - \hat{x}_{t|t-1}). \quad (22)$$

Simplifying, but abusing the Kalman notation this can be written as

$$\hat{x}_{t+1|t} = \mu + \rho(\hat{x}_{t|t-1} - \mu) + \omega(r_t - \hat{x}_{t|t-1}). \quad (23)$$

This is comparable to traditional adaptive expectations when $\rho = 1$.¹⁴ The experiments in this paper will consider equation 23 as the structure for forecasts, but will test to see if wealth will select the correct parameter values. Agent i will use forecasts based on

$$\hat{x}_{t+1|t} = \mu + \rho_i(\hat{x}_{t|t-1} - \mu) + \omega_i(r_t - \hat{x}_{t|t-1}), \quad (24)$$

and are indexed by their parameter pair (ω_i, ρ_i) . The dynamics of wealth will be analyzed to explore its properties in terms of finding optimal forecasts. Also, the Kalman filter structure produces optimal values, (ω^*, ρ^*) as a benchmark to compare with other strategies. For the generated time series used in this paper $(\omega^*, \rho^*) = (0.0164, 0.95)$.

3 Results

3.1 Log utility ($\gamma = 1$): no predictability

The first experiments look at the convergence of wealth to the growth optimal strategy in a world with no return predictability. Agents do not attempt to forecast returns, and they are assumed to know the true mean and variance of returns. Returns are drawn from a normal distribution with the same mean and variance that will be used in future simulations, but the predictable expected return component, x_t is fixed

¹⁴See Evans & Honkapohja (2001) or Sargent (1999) for examples of the connections between state-space modeling and older adaptive expectations ideas.

at zero. Figure 1 displays the evolution of wealth fractions for three levels of risk aversion, $\gamma = 0.5, 1, 2$, over time. The results presented here are cross sectional means from 100 different runs.

The steady increase in the solid line represents the convergence onto the $\gamma = 1$ (log) strategy. The other two trader types eventually disappear in terms of wealth. This is just as the theory would predict.¹⁵ The more interesting feature is to notice that this convergence is relatively slow. A significant dispersion between the $\gamma = 1$ strategy and the others does not appear until after 30 to 40 years of simulated data.

Given the cross section of simulated wealth fractions it is also possible to get a picture of the cross sectional dispersion in these time series averages. This is shown in figure 2. This presents results for the $\gamma = 1$ wealth fraction for the same runs from the previous table. However, in this figure the median, 0.95, and 0.05 quantiles taken from the 100 run cross section are used. This shows that the distribution still has a large amount of dispersion even after 500 years of simulated data have gone by.

3.2 Log utility ($\gamma = 1$): Predictability and learning

The early experiments with no return predictability are interesting, but the focus of this paper is on the case where there is some small return predictability. In this section I use returns which follow the predictable process described in the previous section. This section focusses on the case where all agents have $\gamma = 1$ or log preferences. Agents follow adaptive forecasting rules as in equation 23. They will be heterogeneous both in terms of the gain and memory parameters. The objective is to see how well wealth is drawn to the optimal forecasting parameters determined by the Kalman filter.

The two panels of figure 3 show the long run properties for different forecasting parameters. The parameter pairs are organized on a grid by the memory and gain parameters for the adaptive forecasts. Memory parameters vary from 0.9 to 1.0 incremented by 0.01. Gain parameters vary from 0 to 0.1 incremented by 0.01.¹⁶ This gives a 11 by 13 grid for a total of 143 rules. The dark lines mark the optimal forecast parameters, or strategies formed by using the true conditional expectation. The gain levels at zero are an important value to keep track of since this corresponds to ignoring any new return information in the forecast, and using a constant forecast set to the unconditional mean return.

The top panel of figure 3 shows the cross sectional mean of wealth fractions over 100 runs recorded after 500 years. The contour height (displayed on the right legend) is in units of density divided by uniform density ($1/143 = 0.007$). For example, a contour level of 2 indicates that the corresponding rule has twice the wealth density it would have under a uniform wealth distribution. The figure shows a clear long run concentration on the optimal forecasting rule. This indicates that for the $\gamma = 1$ investor wealth will concentrate on the

¹⁵Since the strategies are actually linear approximations, there is still some value in confirming this result.

¹⁶Extra points are added at 0.005, and 0.0164. Both give a finer grid near zero, and the later makes sure that the true optimal forecast parameters are in the grid.

true conditional expectation based trading rule in the long run.

The lower panel of figure 3 shows the estimated expected utility for the different strategies. This is estimated with the time series mean taken over the 500 years and over the 100 cross section. It is reported as an annual certainty equivalent return. This is estimated as,

$$\log(1 + r_p^*) = E(\log(1 + r_{p,t})), \quad (25)$$

where the above expectation is estimated by taking both time and cross sectional means. The time average is over 500 years, and the cross section is the 100 runs. It shows a region around the optimal rule with a annual certainty equivalent of about 10 percent. This drops off as the distance to the optimum increases. It is interesting to note that the drop off is very steep as one moves to the constant forecasting rules on the left side of the panel where the gain is zero. This indicates the economic usefulness of the adaptive strategy for the $\gamma = 1$ investor. There is also somewhat of an asymmetry in the shape in that utility levels are not all that sensitive to reductions in the memory parameter. Dropping the memory parameter to 0.9 has only a small impact on the certainty equivalence level.

Figure 4 shows the time evolution of the wealth density plot from figure 3. At 5 years there is almost no indication of any convergence. Actually, in this early period there is a small indication of wealth concentration on the constant forecast.¹⁷ By year 20 wealth appears to be converging, and by year 50 the convergence to the true parameters is clear. Although still somewhat slow, this convergence process appears to be a little faster than the convergence over different γ values.

Table 3 shows wealth fractions for a small 3×3 parameter grid. This shows evolution of wealth fractions directly. The values reported are the cross sectional averages with snapshots taken at 5, 10, 20, and 200 years. The years correspond to moving clockwise around the four panels. The center box in each panel corresponds to the optimal forecast parameters. Wealth is converging to this value, but as in the larger grid case, convergence appears to be somewhat slow. At 20 years, there is still only 16 percent of wealth at the optimal forecast.

Figure 5 displays the mean portfolio weights for the different forecast parameters taken over time and the 100 runs. The range of forecast parameters yields both portfolios which are leveraged on average, $\alpha > 1$, and other which are holding some fraction of the risk free asset $\alpha < 1$. Moving in the northwest direction (lower gain, higher memory) increases the aggressiveness of the mean portfolio choice. The value at the optimal forecast parameter pair is approximately 1.00.

¹⁷One possibility for this is that return predictability is too weak to show up at this period. The strategy evolution here might be dominated by the increased volatility on the conditional strategies which will reduce the growth rate.

3.3 Risk aversion ($\gamma = 3, 2$): Predictability and learning

The previous section demonstrated that wealth does converge to the optimal forecasting parameters for the $\gamma = 1$ (log) investor. In that case, the objectives of utility maximization and wealth growth maximization are exactly aligned, so this result is not surprising. Unfortunately, it is not at all clear this is a good level of risk aversion for actual investors. It is generally felt that $\gamma = 1$ is somewhat low. Economists and financial advisors often use a wide range of values of $\gamma > 1$. All these higher levels of risk aversion are probably a much better approximation to the actual population than $\gamma = 1$.

Figure 6 repeats the two panel plot of long run wealth and expected utility levels for $\gamma = 3$. It is clear that the upper panel has changed dramatically. Wealth is no longer converging to the optimal forecast parameters. The wealth density is maximized at a (gain, memory) pair of (0.06, 1.00), well off the optimum values which are again marked by the dark lines. The lower panel reports the utility levels as annual certainty equivalents which are given by,

$$(1 + r_p^*)^{1-\gamma} = E((1 + r_{p,t})^{1-\gamma}) \quad (26)$$

This shows clearly that in utility terms, the optimal forecast parameters are utility maximizing for investors. The maximum certainty equivalent return is estimated as 5.25 percent per year at the optimal forecast parameters. It is only 2.91 at the wealth maximizing parameters which is a loss of almost 2.5 percent, or 50 percent of the certainty equivalent return. The utility surface is again relatively flat in the the memory parameter with little change coming from reducing this to 0.90. The surface is steep in the direction of reducing the gain parameter to zero which again corresponds to moving to a constant weight portfolio.

Figure 7 shows the dynamics of the wealth distributions over time. As in the previous four panel plot these are again cross sectional means over 100 different runs. The pattern is interesting in that the movement away from the optimal values begins early. Even at 5 years, it is clear the wealth density is drifting to the northeast corner. By 20 years there is a pronounced large bias in the gain parameter. This continues at 50 and 200 years.

Table 4 again presents the wealth fractions from a small (3x3) grid experiment with 4 time snapshots. In comparison to table 3 the convergence appears to be a little faster. Even at year 10, there is evidence of a drift to the higher gain level. By year 20, over 50 percent of wealth is concentrated on strategies with a gain parameter well over the optimal.

Why is wealth drifting far from the optimum? This initially seems counter intuitive. This question will be explored more in a future section, but the initial answer is to remember that the $\gamma = 3$ investors are maximizing utility not wealth growth. Wealth is selecting, in the evolutionary sense of selection, for a strategy closest to the growth optimal strategy. Figure 8 explores the properties of the different parameter

pairs. It reports the expected growth rate for the different dynamic portfolios across the forecasting pairs. This is given by

$$E(\log(1 + r_{p,t})). \quad (27)$$

It is again estimated as the mean taken over the 500 years of weekly data, and the 100 runs in the cross section. The optimal growth portfolio subject to the $\gamma = 3$ preferences is at a biased set of parameters, and agrees with the long run wealth convergence point.

The mean portfolio weights are given in figure 9. The growth maximizing portfolio would on average be taking a position with $\alpha = 0.55$. The portfolio weight for the utility maximizing (gain, memory) pair is only $\alpha = 0.39$. In contrast to figure 5 the mean portfolio weights are increasing as the gain and memory parameters are increased. This shows that the wealth maximization point is choosing a more aggressive portfolio than the utility maximization point which is closer to the log utility, growth optimal, portfolio.

Finally, figure 10 repeats the wealth and utility plots for $\gamma = 2$. At this more moderate level of risk aversion, the differences are still present, but less dramatic. The wealth density is maximized at a (gain, memory) pair of (0.04, 0.98). The annual certainty equivalent return at this point is 6.25 percent. The corresponding return at the optimal forecast parameters is 7.00. This is a much smaller difference in utility terms than for the $\gamma = 3$ case. However, the difference between the forecast parameter pairs still appears quite large.

4 Robustness Checks and Causes

4.1 Causes and portfolio constraints

The general theoretical reason for wealth selecting distorted forecast parameters is clear. Wealth and utility maximization are not the same, and there is no reason at $\gamma \neq 1$ for wealth to converge to the optimal forecast parameters. However, why is it biasing so far from the optimal forecast parameters, and why does the bias tend to be on the high side? Moving from $\gamma = 1$ to $\gamma = 3$ shifts to a more conservative optimal strategy. The optimal $\gamma = 3$ investor sets α too low relative to the growth optimal investor. Optimal wealth growth will occur at a set of forecast parameters which moves into higher return/risk combinations relative to the optimal portfolio for the $\gamma = 3$ investor.

Moving to higher forecast parameters helps increase expected returns through two mechanisms. First, increasing these parameters will increase the unconditional portfolio fraction, α . How changing these parameters could increase α is not immediately clear, since from the structure of the forecasting problem, the unconditional forecast for all parameters is set to the unconditional expected return. Therefore, $E(\alpha_t)$ will

not change as one moves through the parameter grid.¹⁸ A quick glance at figure 9 shows that this is clearly not the case. The reason for the changes in α comes from the fact that the portfolio weights are constrained. First, short sales were restricted by forcing $\alpha > -0.5$, and extreme leverage was prohibited by requiring $\alpha < 2$. These restrictions cut into the tails of an unrestricted α distribution. If they cut into the left tail more frequently than the right, the unconditional $E(\alpha)$ will increase as α volatility increases.

One way of exploring this is to eliminate all portfolio restrictions. This is done in figure 11. This shows a reduced change in the gain parameter from the optimal forecast parameters. The new wealth concentration is at a (gain, memory) combination of (0.05, 0.98). The utility drop off at this point is still significant as shown in the lower panel. The annual certainty equivalent return at the utility maximizing point is 5.4 percent per year, while at the wealth maximizing point it falls to -4.0 percent per year. This dramatic reduction is coming from the increase in volatility driven by the much more extreme portfolio positions taken in this unconstrained situation.

This also shows that the bounds coming from the constrained α values, and their impact on $E(\alpha)$ are not the only cause for the increase in expected returns. A second effect comes from increasing the covariance of the forecast with returns. This can be easily seen from,

$$E(R_p) = E(\alpha_t(R_{t+1} - R_f) + R_f) \quad (28)$$

$$= cov(\alpha_t, R_{t+1}) + E(\alpha_t)E(R_{t+1} - R_f) + R_f \quad (29)$$

In the unconstrained case changing the memory or gain parameter has no impact on $E(\alpha_t)$, so the change in the expected return of the strategy must come from the covariance. Given that α_t is linear in the expected return, the change in filtering parameters must result in an increase in the covariance of the forecast with future returns. For a linear forecast centered around the unconditional mean, the covariance can be increased by multiplying it by a factor greater than one.

A quick quantitative picture of these values is given in table 5. This table presents strategy summary statistics (averaged over 100 runs, and taken from the 500 year run) for strategies in both the constrained and unconstrained cases. *Umax* refers to strategies at the utility maximizing parameters, and *Wmax* refers to the wealth maximizing parameters. In both cases both the risk and return of the dynamic portfolios increases when one moves to the wealth maximizing parameters. In the constrained case the expected value of the risky asset share, $E(\alpha_t)$, increases, driven by the impact of the short sale constraint. In the unconstrained case the expected risky share is almost constant at 0.37 and 0.39 with the small change coming from sampling variability. In both cases there is an increase in the covariance of α_t with r_{t+1} which drives

¹⁸This can be seen by taking unconditional expectations of equations 15 and 23.

up the expected return to the dynamic strategy.

The last two lines in the table give an indication for how binding the constraints are on the portfolios. For the utility maximizing constrained portfolio, 6 percent are constrained by the short sale restriction, but only 1 percent are constrained by the maximum leverage restriction. These constraints become more binding at the wealth maximizing parameters where 31 percent are bound by the short sale constraint, and 30 percent hit the upper leverage bound. Comparable numbers are presented for the unconstrained case. Here, they correspond to actual portfolio weights outside the constraint bounds.

It seems possible that both of these effects may be important in the real world. The short and leverage constraints can either be thought of as sensible real investment constraints put on by agents following these strategies. They also could be viewed as crude implementation of a Bayesian decision system which may be unsure about the performance of these predictors. A reasonable argument could be made that the original restrictions, allowing α to be as large as 2 might have been too generous. Figure 12 shows the impact of eliminating both short sales and leverage, $0 < \alpha_t < 1$. This pushes the forecast parameters even farther from the optimal target. However, the portfolio constraints reduce the utility loss to a modest change from 4.81 to 4.62 percent per year in certainty equivalent returns.

4.2 Heterogeneous preferences and learning

Up to this point, I have concentrated on a single level of risk aversion in the experiments. Obviously, in the real world we would expect some heterogeneity across individual attitudes toward risk. This section performs some short tests looking at wealth evolution across both agent types, and forecast parameters. To keep the simulations more tractable, the memory parameter will be fixed at 0.95, and the gain parameter is allowed to vary from 0 to 0.06. Risk aversion will vary in the range, $0.5 < \gamma < 2.5$.

The theoretical work on growth optimal strategies suggests that we should see convergence to the log ($\gamma = 1$) type with true probabilities. Figure 13 shows that this is indeed the case with an eventual long run convergence to the (γ, gain) pair of $(1, 0.016)$. These wealth fractions are again, cross sectional means from 100 different series. As in many of the other runs, this convergence appears to be relatively slow. This system actually gives some indication of an early bias toward very low gamma strategies. It is also clear that there is a small wealth identification issue across gain and γ . Moving to higher γ would reduce the variability of the portfolio position. However, this reduction can be partially compensated for by increasing the gain. In other words a set of strategies sitting close to a line going through the optimal point will be similar in terms of their performance, and will be hard to discern in terms of wealth evolution.

These results do not invalidate those from the previous sections. They simply affirm that if the growth

optimal strategy is in the population, then convergence to the true forecasting parameters will occur. In general, it is believed that growth optimal objectives are much less risk averse than the general population.

5 Conclusions

In this paper I have shown that wealth evolution alone can converge to forecasting strategies with distorted beliefs. These results were predicted by theoretical work on agent evolution, but the results here are performed using a familiar forecasting setup calibrated to known patterns in financial time series.

The most important qualitative aspect of the forecast distortion is the selection of gain parameters which are well above the optimal forecast. This translates into wealth concentrating on strategies which put too much weight on current returns in their forecasts. Such strategies could correspond to the presence of momentum and trend following strategies in actual markets.

Several extensions to this model would appear to be important. First, eventually experiments will need to follow the theoretical literature, and much of the computational literature, and endogenize prices. One can always question whether the convergence results given here would be affected by price changes as wealth moves around. This paper deliberately eliminated this effect, but in the future it has to be part of the analysis of evolution and financial markets. A second simplification which may have a large impact was the assumption that the innovations to the expected return process and the return noise process are independent. This deviates from some of the forecasting evidence and many of the prediction models in use. It also has a nontrivial impact on the structure of the Kalman filter forecasting system.¹⁹ If there is a large enough negative correlation then a large recent return could have a negative impact on predicted future expected returns, and therefore would warrant a negative gain parameter in this system. Implementing a richer evolutionary system which better addresses this issue is another important extension.

The basic point of this paper is simple. Wealth evolution, on its own, will not reliably perform the function that it is often assumed to do. This obviously forces some difficult choices for researchers building heterogeneous adaptive models. If optimal forecasts hold any sway as a point markets might be tending toward, then this movement must be coming from the active learning side. Worse, the passive learning side will be slowly, and steadily working against this drift. Unfortunately, modeling active learning is much more difficult than passive learning. These results suggest that there will always remain some amount of wealth concentrated on strategies that would be difficult to explain from the standpoint of optimal forecasting. Understanding exactly what one would expect this wealth distribution to look like as it moves through time, and its impact on prices still remains an interesting question.

¹⁹See Pastor & Stambaugh (2006) for discussion.

References

- Alchian, A. (1950), ‘Uncertainty, evolution, and economic theory’, *Journal of Political Economy* **58**, 211–221.
- Berrada, T. (2006), Bounded rationality and asset pricing, Technical report, Swiss Finance Institute.
- Blume, L. & Easley, D. (1990), ‘Evolution and market behavior’, *Journal of Economic Theory* **58**, 9–40.
- Blume, L. & Easley, D. (2006), ‘If you’re so smart, Why aren’t you rich? Belief selection in complete and incomplete markets’, *Econometrica* **74**, 929–966.
- Breiman, L. (1961), Optimal gambling systems for favorable games, in J. Newyman & E. Scott, eds, ‘Proceedings of the Fourth Berkeley Symp. of Math Statistics, and Probability’, Vol. 1, University of California Berkely Press, Berkely, CA.
- Campbell, J. Y. & Shiller, R. (1988), ‘The dividend-price ratio and expectations of future dividends and discount factors’, *Review of Financial Studies* **1**, 195–227.
- Campbell, J. Y. & Viceira, L. M. (1999), ‘Consumption and portfolio decisions when expected returns are time varying’, *Quarterly Journal of Economics* **114**, 433–495.
- Campbell, J. Y. & Viceira, L. M. (2002), *Strategic Asset Allocation*, Oxford University Press, Oxford, UK.
- Cecchetti, S., Lam, P. & Mark, N. C. (2000), ‘Asset pricing with distorted beliefs: Are equity returns too good to be true?’, *American Economic Review* **90**, 787–805.
- DeLong, J. B., Shleifer, A., Summers, L. H. & Waldmann, R. (1991), ‘The survival of noise traders in financial markets’, *Journal of Business* **64**, 1–19.
- Evans, G. W. & Honkapohja, S. (2001), *Learning and Expectations in Macroeconomics*, Princeton University Press.
- Evstigneev, I. V., Hens, T. & Schenk-Hoppe, K. R. (2006), ‘Evolutionary stable stock markets’, *Economic Theory* **27**, 449–468.
- Fama, E. F. & French, K. R. (1988), ‘Dividend yields and expected stock returns’, *Journal of Financial Economics* **22**, 3–25.
- Ferson, W. E., Sarkissian, S. & Simin, T. T. (2003), ‘Spurious regressions in financial economics?’, *Journal of Finance* **58**, 1393–1413.

- Figlewski, S. (1978), ‘Market “efficiency” in a market with heterogeneous information’, *Journal of Political Economy* **86**(4), 581–597.
- Friedman, M. (1953), *Essays in Positive Economics*, University of Chicago Press, Chicago, IL.
- Giovannini, A. & Weil, P. (1989), Risk aversion and intertemporal substitution in the capital asset pricing model, Technical Report 2824, National Bureau of Economic Research, Cambridge, MA.
- Hakansson, N. H. (1971), ‘Multi-period mean-variance analysis: Toward a general theory of portfolio choice’, *Journal of Finance* **26**.
- Hamilton, J. D. (1994), *Time Series Analysis*, Princeton University Press, Princeton, NJ.
- Harvey, A. C. (1989), *Forecasting, Structural Time Series Models, and the Kalman Filter*, Cambridge University Press.
- Keim, D. & Stambaugh, R. (1986), ‘Predicting returns in the bond and stock markets’, *Journal of Financial Economics* **17**, 357–390. %K %X %F.
- Kelley, J. L. (1956), ‘A new interpretation of information rate’, *Bell Systems Technical Journal* **35**, 917–926.
- Kogan, L., Ross, S. A. & Wang, J. (2006), ‘The price impact and survival of irrational traders’, *Journal of Finance* **61**, 195–229.
- LeBaron, B. (1992), Do moving average trading rule results imply nonlinearities in foreign exchange markets?, Technical report, University of Wisconsin - Madison, Madison, Wisconsin.
- LeBaron, B. (2001), ‘Evolution and time horizons in an agent based stock market’, *Macroeconomic Dynamics* **5**(2), 225–254.
- LeBaron, B. (2006), Agent-based computational finance, in K. L. Judd & L. Tesfatsion, eds, ‘Handbook of Computational Economics’, Elsevier, pp. 1187–1233.
- Markowitz, H. (1976), ‘Investment for the long run: New evidence for an old rule’, *Journal of Finance* **31**, 1273–1286.
- Pastor, L. & Stambaugh, R. F. (2006), Predictive systems: Living with imperfect predictors, Technical report, University of Chicago, Chicago, IL.
- Rubinstein, M. (1991), ‘Continuously rebalanced investment strategies’, *Journal of Portfolio Management* **Fall**, 78–81.

- Samuelson, P. (1971), ‘The “fallacy” of maximizing the geometric mean in long sequences of investing or gambling’, *Proceedings of the National Academy of Science* **68**, 2493–2496.
- Sandroni, A. (2000), ‘Do markets favor agents able to make accurate predictions?’, *Econometrica* **68**, 1303–1341.
- Sargent, T. (1999), *The Conquest of American Inflation*, Princeton University Press, Princeton, NJ.
- Stambaugh, R. (1999), ‘Predictive regressions’, *Journal of Financial Economics* **54**, 375–421.
- Taylor, S. J. (1980), ‘Conjectured models for trends in financial prices, tests and forecasts’, *Journal of the Royal Statistical Society A* **143**, 338–362.
- Taylor, S. J. (2005), *Asset price dynamics, volatility, and prediction*, Princeton University Press, Princeton, NJ.
- Yan, H. (2006), Natural selection in financial markets: Does it work?, Technical report, Yale University.

Table 1: *Return Parameter Values*

Parameter	Value
r_f	0.02
$E(r_t)$	0.07
σ_r	0.20
σ_x^2/σ_r^2	0.02
ρ	0.95

Description: Parameters for return time series. All values are annualized, but simulations are done at the weekly frequency. r_f is the risk free interest rate. $E(r_t)$ is the unconditional expected real return on the risky asset. σ_r is the corresponding annual standard deviation. σ_x^2/σ_r^2 is the signal to noise ratio in the returns series. ρ is the AR(1) persistence parameter for the expected return process.

Table 2: *Return Time Series Simulations*

Sample	ρ_1	ρ_2	ρ_3	ρ_4	R^2
25 Years	0.018 (0.028)	0.013 (0.028)	0.016 (0.029)	0.015 (0.028)	0.110 (0.106)
50 Years	0.019 (0.020)	0.017 (0.020)	0.017 (0.020)	0.015 (0.020)	0.100 (0.075)

Description: Mean values from 1000 monte-carlo return simulations corresponding to 25 and 50 years. ρ_j is the return autocorrelation at j week lag. R^2 is the R^2 of a annual regression of year $t+1$ returns on x_t , the expected return, at the end of year t . Numbers in parenthesis are the standard deviations of these estimated values from the 1000 length cross section.

Table 3: *Mean wealth fractions (small grid), $\gamma = 1$*

Years = 5	Gain			Years = 10	Gain		
Memory	0	0.0164	0.06	Memory	0	0.0164	0.06
0.91	0.097	0.119	0.116	0.91	0.079	0.128	0.120
0.95	0.097	0.121	0.115	0.95	0.079	0.134	0.121
0.99	0.097	0.120	0.120	0.99	0.079	0.130	0.131
Years = 20	Gain			Years = 100	Gain		
Memory	0	0.0164	0.06	Memory	0	0.0164	0.06
0.91	0.047	0.147	0.131	0.91	0.002	0.189	0.118
0.95	0.047	0.159	0.138	0.95	0.002	0.246	0.142
0.99	0.047	0.140	0.143	0.99	0.002	0.149	0.150

Description: This table presents wealth fractions which are cross sectional averages over 100 runs taken at 4 different time periods, 5, 10, 20, and 100 years. The runs estimated wealth fractions from a 3×3 parameter grid, varying both the memory and gain parameter in the forecast filters. The center cell in the grid corresponds to the optimal forecast parameter pair.

Table 4: *Mean wealth fractions (small grid), $\gamma = 3$*

Years = 5	Gain			Years = 10	Gain		
Memory	0	0.0164	0.06	Memory	0	0.0164	0.06
0.91	0.099	0.110	0.118	0.91	0.083	0.103	0.126
0.95	0.099	0.112	0.121	0.95	0.083	0.110	0.136
0.99	0.099	0.116	0.124	0.99	0.083	0.126	0.149
Years = 20	Gain			Years = 100	Gain		
Memory	0	0.0164	0.06	Memory	0	0.0164	0.06
0.91	0.050	0.089	0.144	0.91	0.001	0.018	0.135
0.95	0.050	0.106	0.164	0.95	0.001	0.040	0.250
0.99	0.050	0.143	0.203	0.99	0.001	0.097	0.456

Description: This table presents wealth fractions which are cross sectional averages over 100 runs taken at 4 different time periods, 5, 10, 20, and 100 years. The runs estimated wealth fractions from a 3×3 parameter grid, varying both the memory and gain parameter in the forecast filters. The center cell in the grid corresponds to the optimal forecast parameter pair.

Table 5: *Strategy Summary*

	UMax	WMax	UMax (unconstrained)	WMax(unconstrained)
$E(r_p)x100$	0.16	0.23	0.17	0.43
$\sigma_{r_p}^2 x100$	0.04	0.12	0.04	0.47
$E(\alpha_t)$	0.48	0.73	0.39	0.37
σ_{α}^2	0.35	1.15	0.40	6.04
$cov(\alpha_t, r_{t+1})x100$	0.08	0.14	0.09	0.35
$Prob(\alpha_t \leq -0.5)$	0.06	0.31	0.08	0.36
$Prob(\alpha_t \geq 2)$	0.01	0.30	0.01	0.25

Description: This table presents summary statistics on two dynamic strategies evaluated in both a constrained ($-0.5 < \alpha < 2$) and an unconstrained case. UMax corresponds to the utility maximizing (gain,memory) pair, and WMax corresponds to the wealth maximizing (gain,memory) pair. The table shows the expected return for the portfolio, its variance, the expected value and variance for the portfolio weight α . It also shows the covariance of the weight with future returns, and finally the fraction of portfolio weights hitting the constraints.

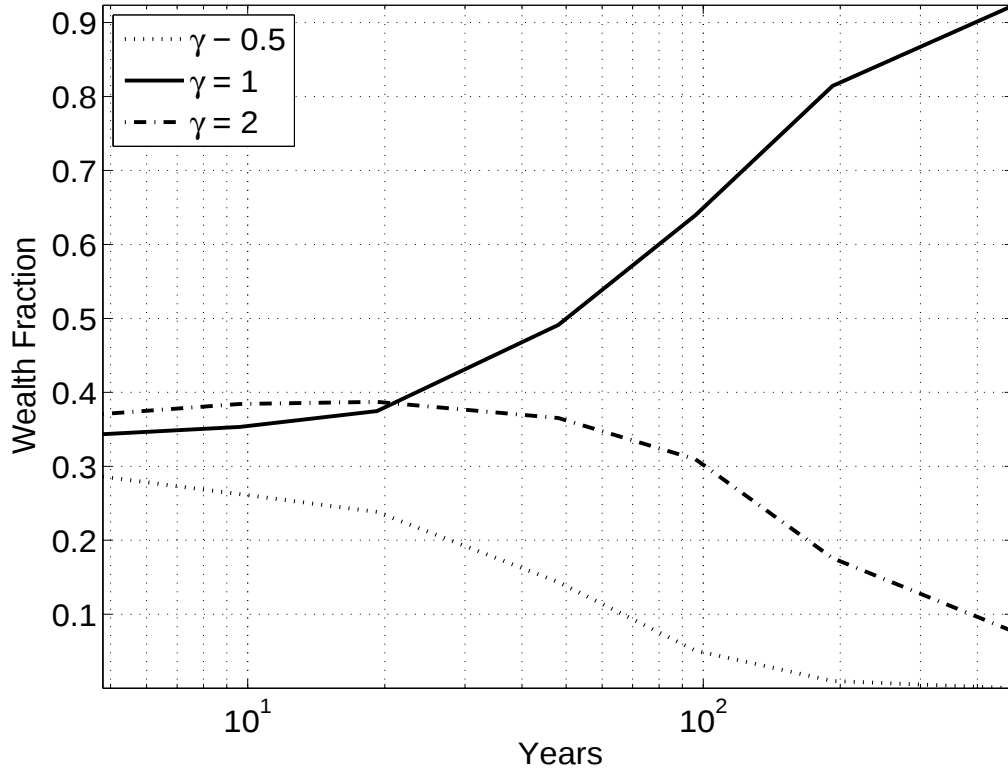


Figure 1: **Mean wealth fractions for $\gamma = 0.5, 1, 2$**

Description : This figure presents mean wealth fractions averaged over a cross section of 100 different runs. Simulated returns are independent, and normally distributed corresponding to the unconditional parameters for the mean and variance given in table 1.

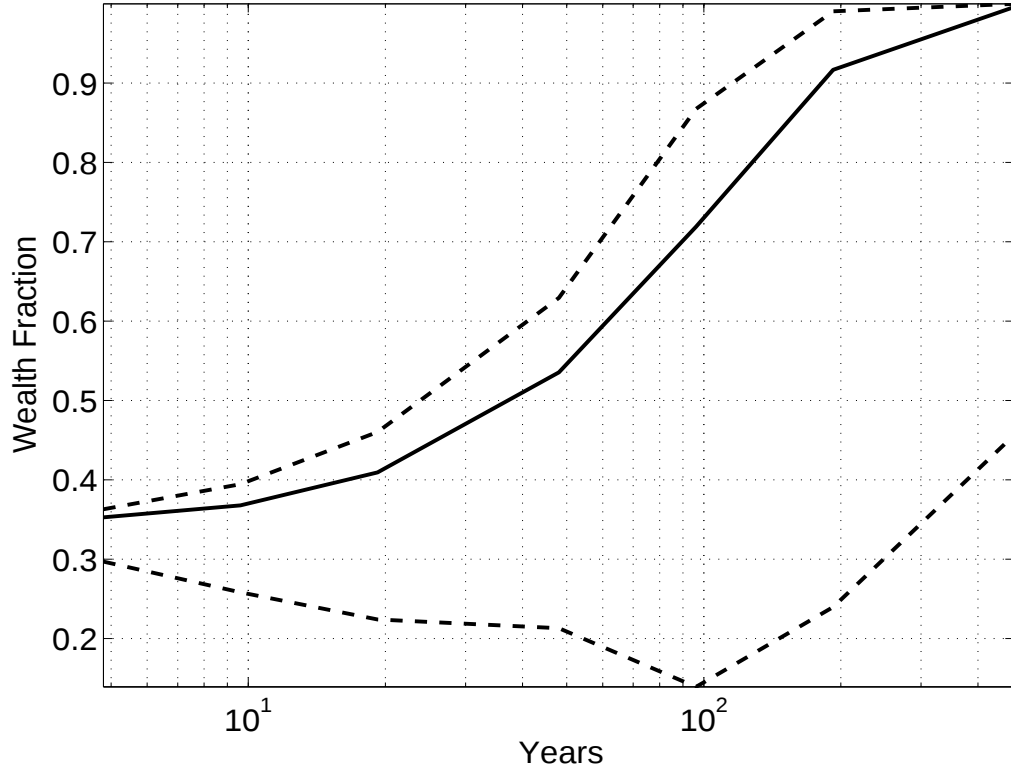


Figure 2: **Wealth fractions for $\gamma = 1$: Median and 5, 95 quantiles**

This figure presents the median, 0.05, and 0.95 quantiles for the wealth fractions of the $\gamma = 1$ strategy. The distribution values are taken from a 100 run cross section sampled at representative time steps.

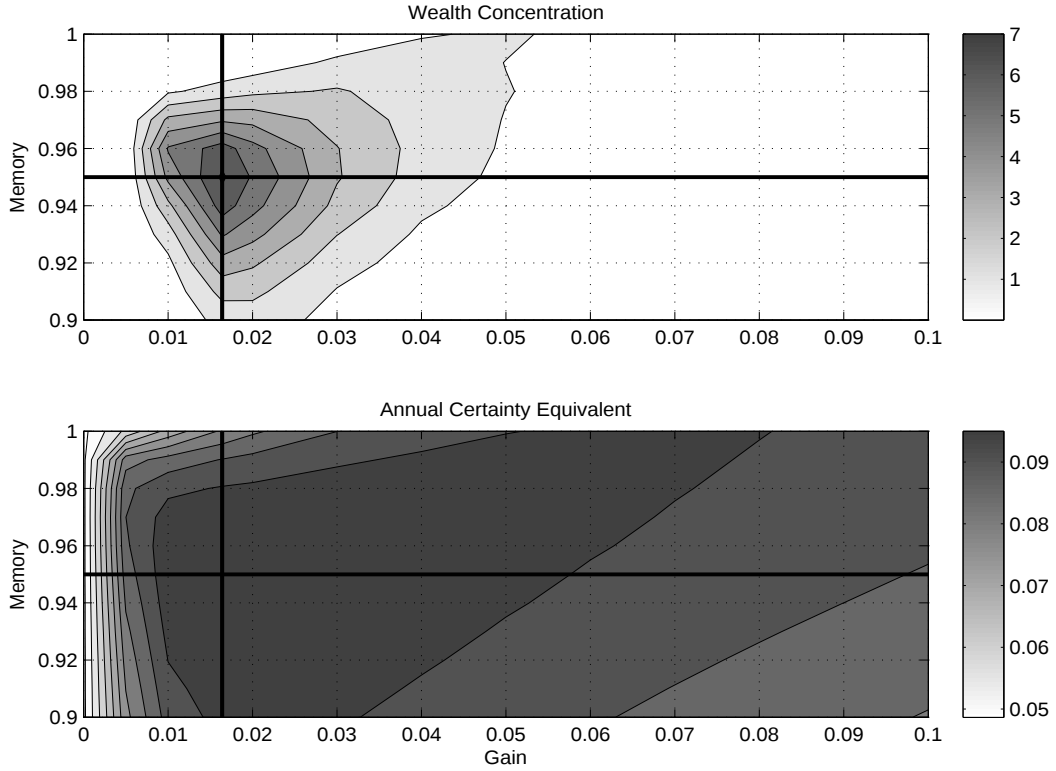


Figure 3: **Wealth and utility surfaces for $\gamma = 1$**

This upper panel in this figure shows the wealth distribution after 500 years estimated as a mean over a 100 run cross section. The figure shows the density over the different strategies indexed by the memory and Kalman gain parameters. The height measures the density at each grid point relative to a uniform density. The lower panel measures the expected utility of each rule reported in units of annual certainty equivalent returns. Both maximums correspond to the optimal Kalman forecast parameters of (0.016, 0.950). The maximum certainty equivalent return is 0.100, or 10 percent per year.

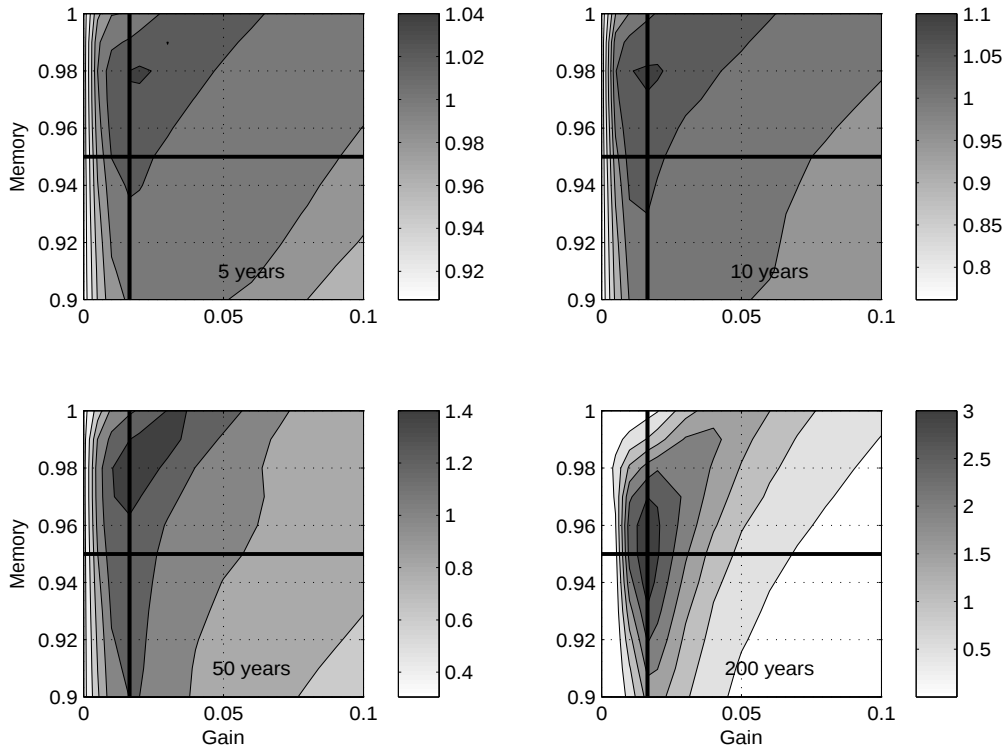


Figure 4: **Wealth surfaces for $\gamma = 1$: Time evolution**

This set of 4 figures displays the time evolution of the cross sectional averages of wealth distributions. Distributions are sampled at the 5, 20, 50, and 200 year time periods.

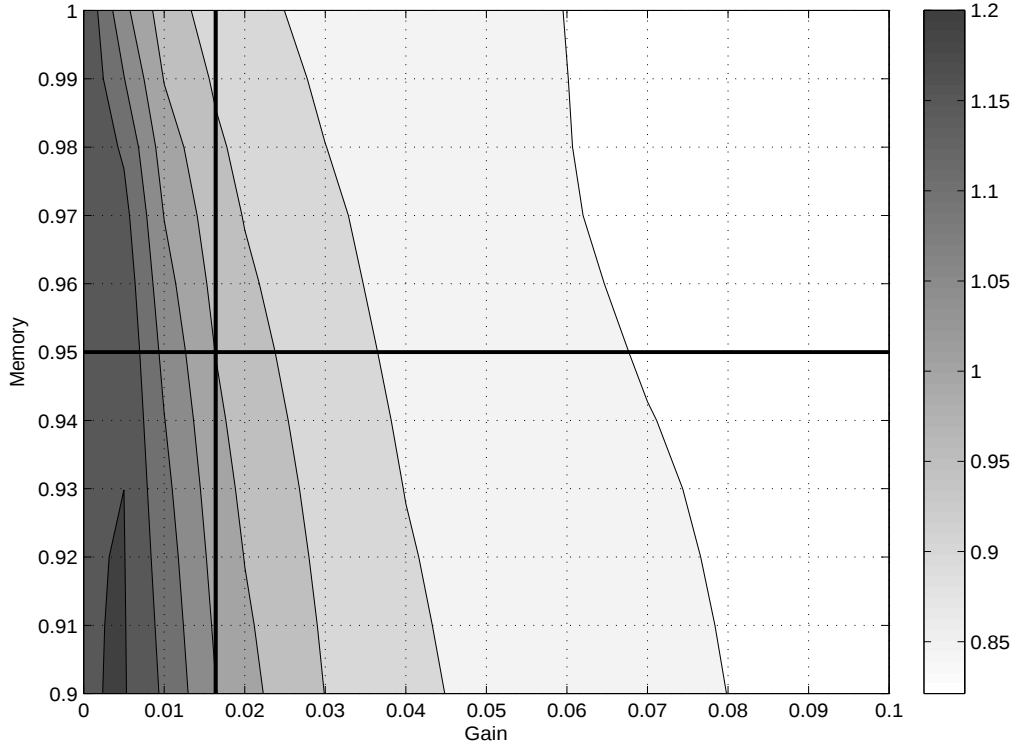


Figure 5: **Mean portfolios: Fraction of wealth in risky asset for $\gamma = 1$**
This figure displays the mean portfolio weights, $E(\alpha_t)$, for each forecast value pair. α_t is the fraction of wealth in the risky asset at time t , and is constrained by $-0.5 < \alpha < 2$. Note that at the optimal forecast parameters the $\gamma = 1$ investor would be fully invested in the risky asset, but not leveraged, $\alpha = 1$.

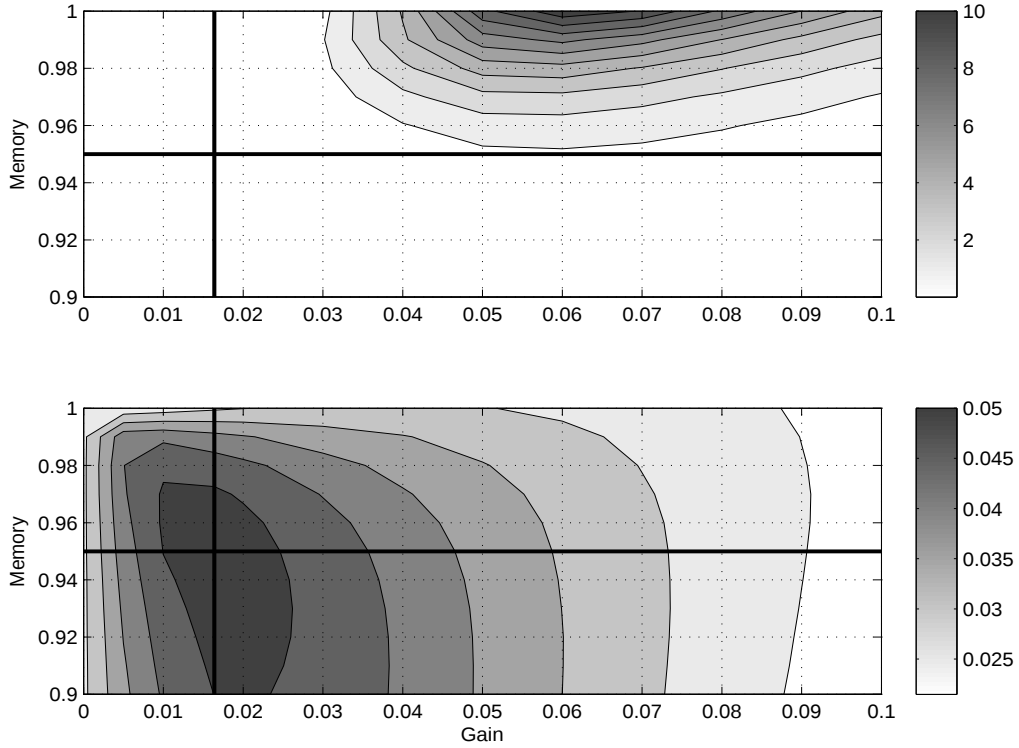


Figure 6: **Wealth and utility surfaces for $\gamma = 3$.**

This upper panel in this figure shows the wealth distribution after 500 years estimated as a mean over a 100 run cross section. The figure shows the density over the different strategies indexed by the memory and Kalman gain parameters. The height measures the density at each grid point relative to a uniform density. The lower panel measures the expected utility of each rule reported in units of annual certainty equivalent returns. The maximum of the wealth density is at the (gain, memory) pair of (0.06, 1.00). The annual certainty equivalent return at this point is 2.91 percent which compares to an annual certainty equivalent return of 5.25 at the optimal forecast parameters.

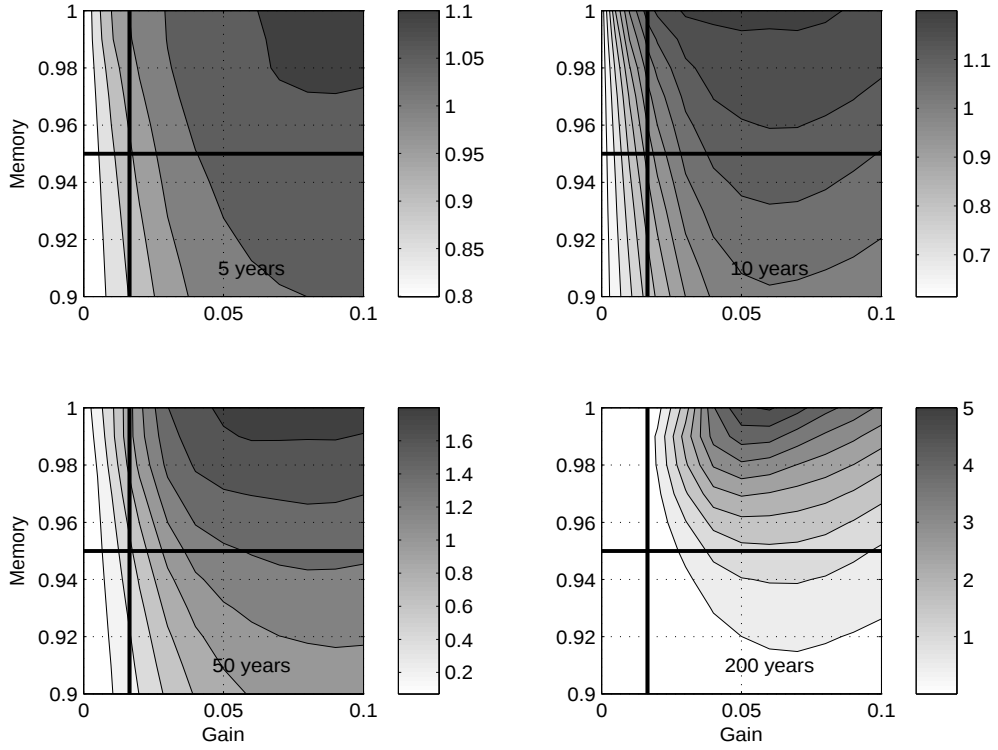


Figure 7: **Wealth surfaces for $\gamma = 3$: Time evolution**

This set of 4 figures displays the time evolution of the cross sectional averages of wealth distributions. Distributions are sampled at the 5, 20, 50, and 200 year time periods.

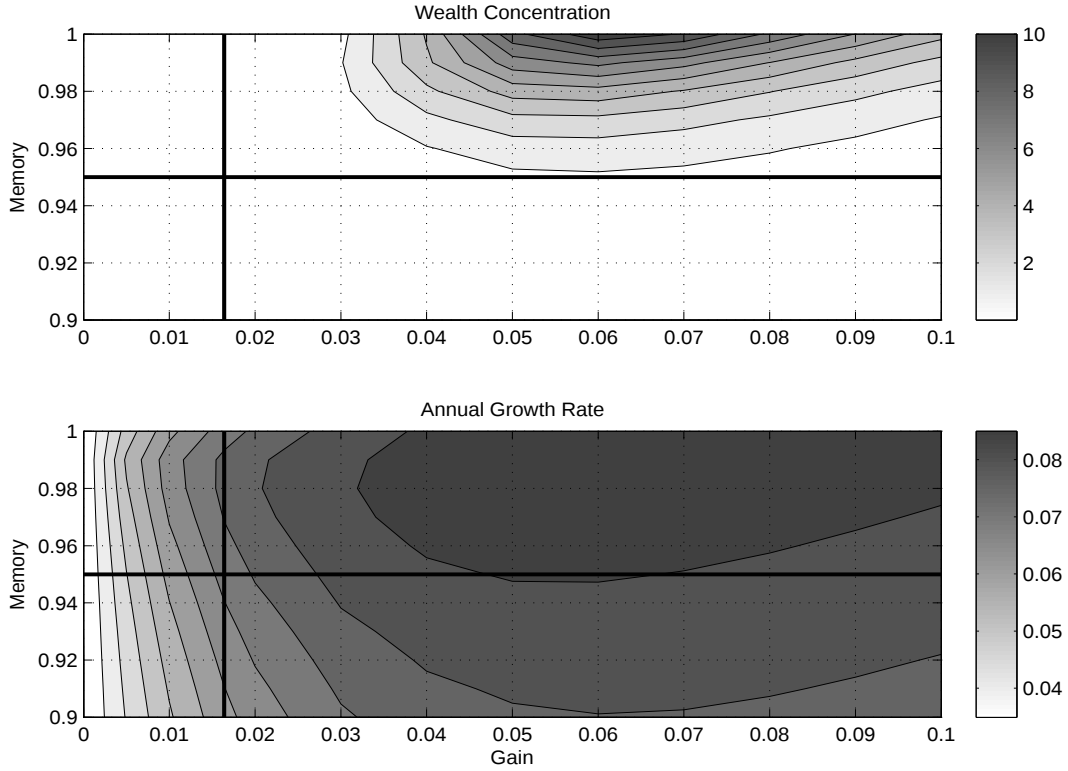


Figure 8: **Wealth expected growth for $\gamma = 3$.**

This figure repeats the 500 year wealth densities in the top panel. The bottom panel now displays the expected growth rate for the dynamic portfolios chosen for the different forecast parameters. The growth maximizing parameters line up with the long run wealth maximizing parameters as it should.

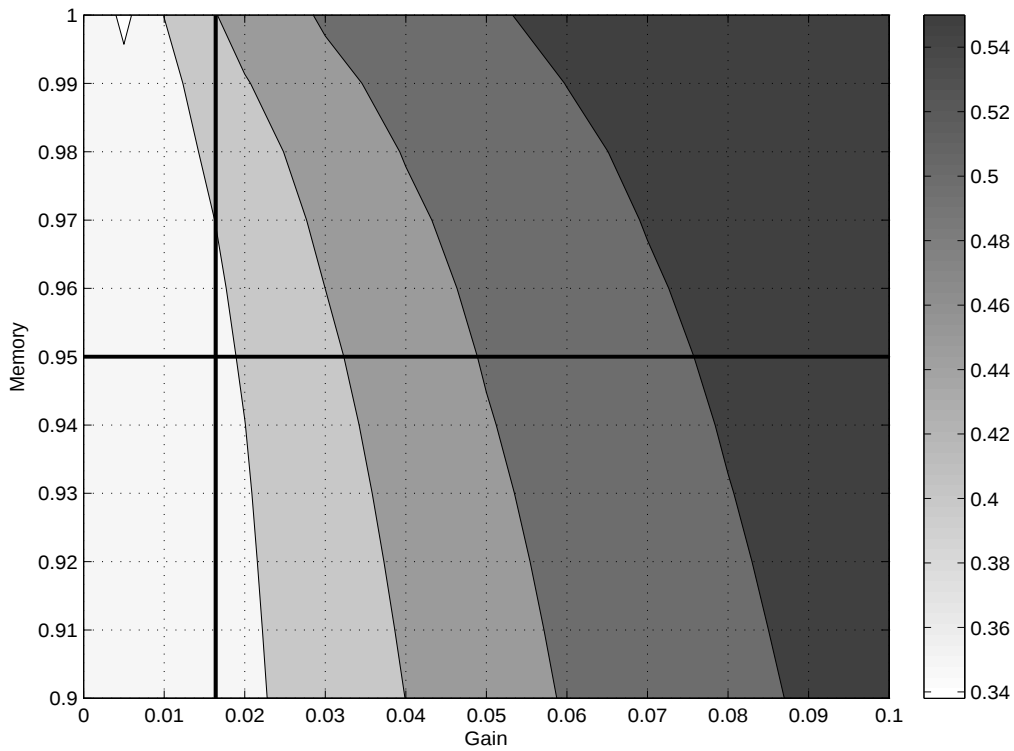


Figure 9: **Mean portfolios: Fraction of wealth in risky asset for $\gamma = 3$**
This figure displays the mean portfolio weights, $E(\alpha_t)$, for each forecast value pair. α_t is the fraction of wealth in the risky asset at time t , and is constrained by $-0.5 < \alpha < 2$.

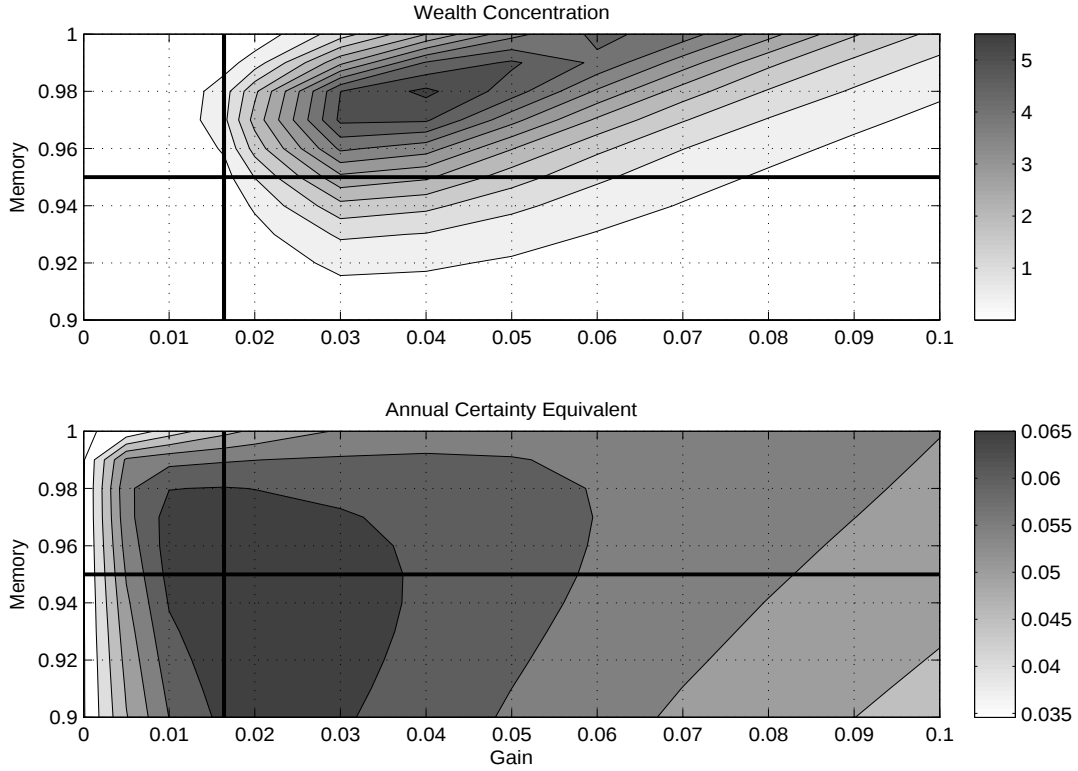


Figure 10: **Wealth and utility:** $\gamma = 2$

This upper panel in this figure shows the wealth distribution after 500 years estimated as a mean over a 100 run cross section. The figure shows the density over the different strategies indexed by the memory and Kalman gain parameters. The height measures the density at each grid point relative to a uniform density. The lower panel measures the expected utility of each rule reported in units of annual certainty equivalent returns. The maximum of the wealth density is at the (gain, memory) pair of (0.04, 0.98). The annual certainty equivalent return at this point is 6.25 percent which compares to an annual certainty equivalent return of 7.00 at the optimal forecast parameters.

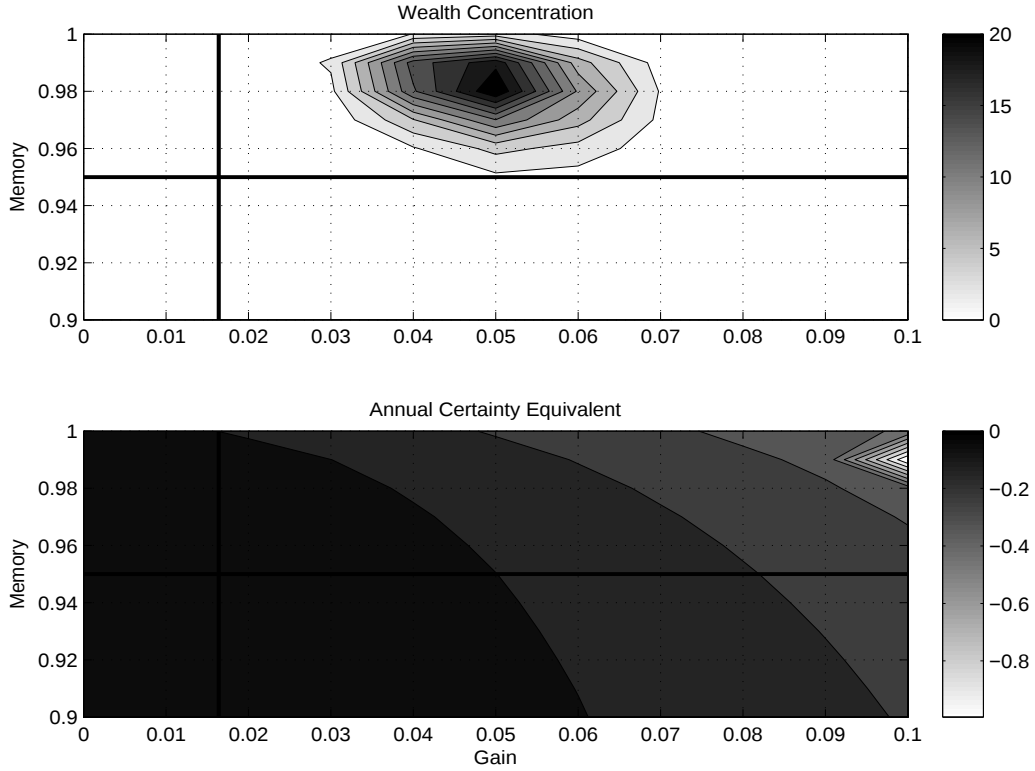


Figure 11: **Wealth and utility:** $\gamma = 3$ α unconstrained

This upper panel in this figure shows the wealth distribution after 500 years estimated as a mean over a 100 run cross section. The figure shows the density over the different strategies indexed by the memory and Kalman gain parameters. The height measures the density at each grid point relative to a uniform density. The lower panel measures the expected utility of each rule reported in units of annual certainty equivalent returns. The maximum of the wealth density is at the (gain, memory) pair of (0.05, 0.98). The annual certainty equivalent return at this point is -4.00 percent which compares to an annual certainty equivalent return of 5.45 at the optimal forecast parameters.

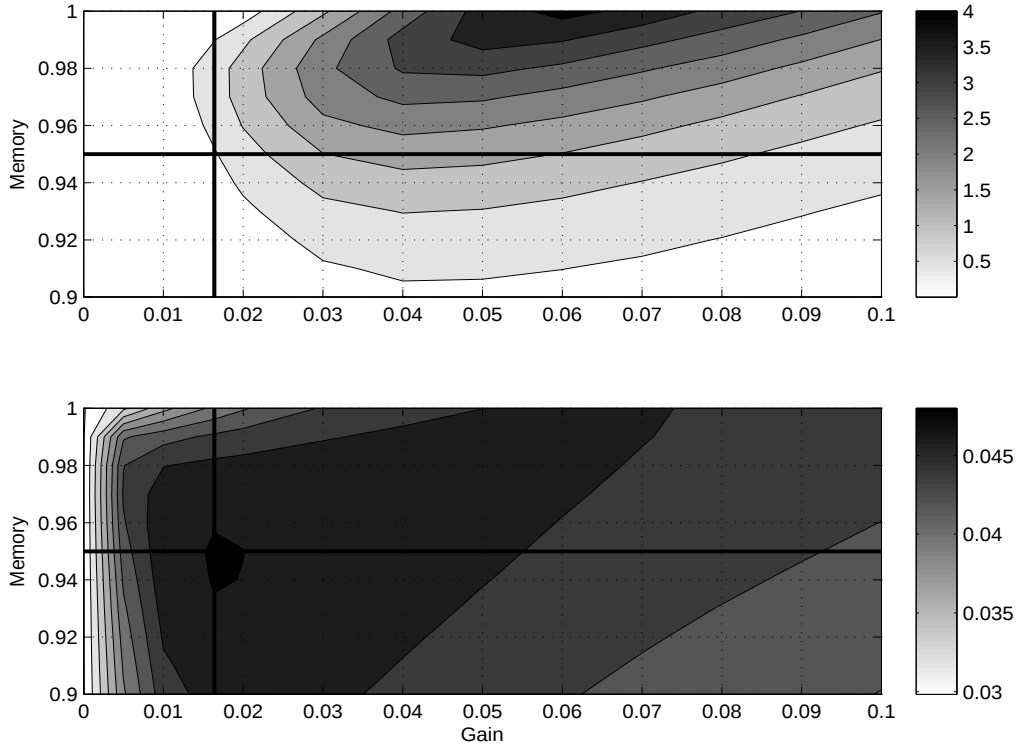


Figure 12: **Wealth and utility:** $\gamma = 3$, $0 < \alpha < 1$

This upper panel in this figure shows the wealth distribution after 500 years estimated as a mean over a 100 run cross section. The figure shows the density over the different strategies indexed by the memory and Kalman gain parameters. The height measures the density at each grid point relative to a uniform density. The lower panel measures the expected utility of each rule reported in units of annual certainty equivalent returns. The maximum of the wealth density is at the (gain, memory) pair of (0.06, 1.00). The annual certainty equivalent return at this point is 4.62 percent which compares to an annual certainty equivalent return of 4.81 at the optimal forecast parameters.

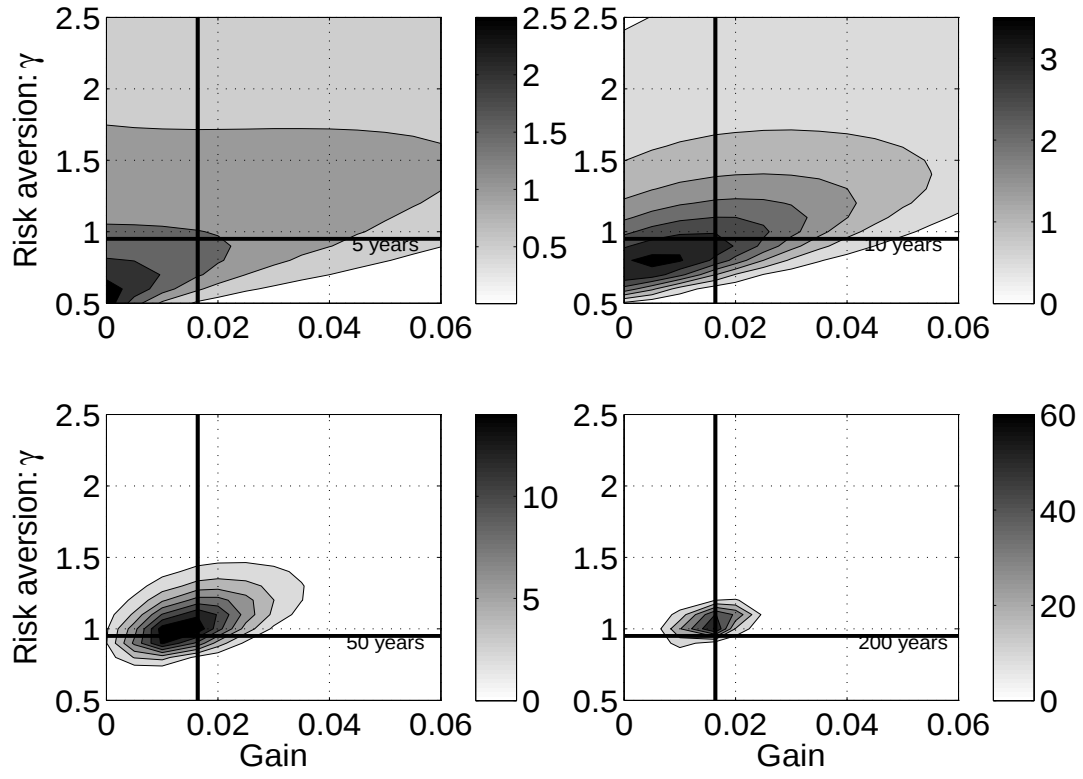


Figure 13: **Wealth and utility:** $\gamma = [0.5, 2.5]$

Description: Evolution over risk aversion (γ) and gain parameter. Means over 100 runs.