

5-4-2009

Why Do Macroeconomic Forecasters Forecast Inaccurately?

Jayant Gokhale

Macalester College, jgokhale@macalester.edu

Follow this and additional works at: http://digitalcommons.macalester.edu/economics_honors_projects



Part of the [Economics Commons](#)

Recommended Citation

Gokhale, Jayant, "Why Do Macroeconomic Forecasters Forecast Inaccurately?" (2009). *Honors Projects*. Paper 17.
http://digitalcommons.macalester.edu/economics_honors_projects/17

This Honors Project is brought to you for free and open access by the Economics Department at DigitalCommons@Macalester College. It has been accepted for inclusion in Honors Projects by an authorized administrator of DigitalCommons@Macalester College. For more information, please contact scholarpub@macalester.edu.

WHY DO MACROECONOMIC FORECASTERS FORECAST INACCURATELY?

An Examination of the Relationship between Herding and Forecast
Accuracy¹

Jayant Gokhale
Honors Thesis
Department of Economics
Macalester College

Faculty Advisor: Professor J. Peter Ferderer

Abstract:

In this paper, I examine why forecasters inaccurately predict the annual growth rate of real GDP in late 1990s (the dot com boom) and early 21st century. I argue that forecasters herd around the lagged consensus (the mean forecast) which, when uninformative, leads them to converge to the wrong prediction. Using data from the *Blue Chip Economic Indicators* newsletter and the Real Time Research Center at the Federal Reserve Bank of Philadelphia for the 1994-2002 period, I econometrically test for the presence of herding and its impact on accuracy. The results suggest that (1) forecasters do herd to “the wisdom of the crowd”, (2) forecaster herding propensities and forecaster accuracy vary from year to year (3) greater forecaster herding leads to greater inaccuracy during the “new economy boom” of the late 1990s.

Keywords: macroeconomic forecasting, business cycles, herding, Blue Chip Survey

¹ Many people played a pivotal role in the completion of this project. I would like to thank my faculty advisor Professor J. Peter Ferderer for his guidance and support. I would like to recognize the contributions of my readers, Professor Gary Krueger and Professor Vittorio Addona. I would also like to thank my peers for their insightful feedback. Lastly, I would like to thank my parents for their constant support and encouragement.

1. Introduction

Good forecasts should be unbiased, efficient, and have serially uncorrelated errors. An unbiased forecast is one that is, on average, neither above nor below the actual value of the forecasted variable. That is, the forecast error has a mean of zero. An efficient forecast is one that incorporates all available information. Finally, the absence of serially correlated errors implies that if, for example, the forecaster underestimated GDP growth this month they tend not to do so in the next month. That is, forecasters learn from their mistakes and do not make systematic errors

To study macroeconomic forecasts, economists have utilized forecasts made by professional forecasters collected in various surveys. For example, Schuh (2001) examines forecasts from three different surveys: the Survey of Professional Forecasters, the Blue Chip survey and a survey of economists conducted by the Wall Street Journal. He finds that participants in the Survey of Professional Forecasters persistently and collectively under-forecast GDP growth during the late 1990s. That is, from 1996 to 1999 every forecaster in the panel provided a GDP growth forecast that was too low relative to actual growth subsequently observed. This is a stunning observation and suggests that the forecasters were using the wrong model to predict the future behavior of the macroeconomic system.

Given that this collective error in the second half of the 1990s occurred during the “new economy” boom suggests that structural changes in the economy may have played a role. The late 1990s was dubbed the dot com boom. Computers and the internet were becoming integrated into production processes and services at a pace far greater than anyone could anticipate. Consequently, the productivity of the individual worker rose

substantially and this led to rapid economic growth. Forecasters had not taken this into account and their estimates deviated significantly from the actual growth rate. Yet, it could not have been the sole reason for their inaccuracy. It took forecasters a long time to incorporate this information into their forecasts. Why?

In this thesis, I utilize forecasts made by the panel of Blue Chip forecasters to study forecasting performance during the 1990s and early part of the 21st century. Specifically, I address two questions of importance. First, does the same pattern observed by Schuh (2001) appear in the Blue Chip panel? Second, if so, how can we explain the tendency of economists to persistently under-forecast economic growth during the boom of 1990s and over-forecast during the bust of the early 21st century?

The paper is organized as follows. Section 2 describes the forecast data from the *Blue Chip Economic Indicators* newsletter and actual Real GDP data from the Federal Reserve Bank of Philadelphia². In this section, I also highlight the controversial debate surrounding the question of an appropriate benchmark (i.e. the actual Real GDP growth estimate) for assessing forecast accuracy. Section 3 provides a review of previous literature on the sources of forecast error. Section 4 engages in discussion of the literature and theory of herding. Section 5 focuses on the econometric analysis of the herding propensities of forecasters and the relationship between herding and forecast accuracy. Section 6 concludes.

² Source: <http://www.philadelphiafed.org/research-and-data/real-time-center/>

2. Forecast and Actual Data

2.1 Blue Chip Data

Since 1976, monthly forecasts made by a set of professional forecasters have been collected and reported in the *Blue Chip Economic Indicators* newsletter. Forecasts for GNP growth, inflation and other major variables are collected by phone during the first three days of each month and disseminated in a newsletter in the same month. In the case of GDP growth, forecasters begin making predictions for the current year in January (a 12-month-ahead forecast) and subsequently revise these forecasts in the following months. This continues until December when the forecasters make their final predictions for the current year³.

The Blue Chip survey data has several advantages over the Survey of Professional Forecasters which is managed by the Federal Reserve Bank of Philadelphia. First, the Blue Chip forecasters are not anonymous and compete with one another for fame and fortune in the forecasting business. Thus, they have an incentive to forecast accurately. Secondly, the Blue Chip forecasts are used by policymakers and receive much attention in the press. Finally, the composition of forecasters in the Blue Chip Panel is diverse and includes forecasts from economists working in investment banks, commercial banks, econometric modeling firms and even the government. Table 1, below, provides a breakdown, by industry⁴, of firms present in the entire Blue Chip data set.

³ Forecasts for the GNP growth rate of the following year are also predicted but I do not focus on these for the sake of brevity.

⁴ Laster et al. (1997) provides the types of classification.

Table 1
Firm Distribution in Blue Chip Survey

Type of firm	Total in each category
Commercial Banks	31
Securities Firms	17
Independents	21
Econometric Modeling Firms	11
Industrial Corporation	16
Other	13
Unknown	15

Note: All the firms in the table provided forecasts for the Blue Chip survey for some interval of time between January 1976 and December 2003.

The more ambiguous industry types are Independents, Econometric Modeling Firms and Other. Econometric Modeling Firms include entities such as Macroeconomic Advisors and UCLA Business Forecast. A lot of econometric modeling firms will be part of research universities. Independents are different because they are private firms. Such firms include Turning Points (Micrometrics) and Econoclast. Finally, the industry type Other includes rating agencies, government agencies and insurance companies.

One disadvantage with the Blue Chip panel is that the composition of forecasters in each year is constantly changing. Though the table above accounts for more than 100 firms in the Blue Chip panel since 1976, not all firms (1) appear simultaneously and (2) forecast consistently from January 1976 to December 2003. In fact, numerous firms are constantly entering and exiting. This makes it slightly difficult to identify a large group of forecasters who forecast for the sample period under observation, 1994-2002. However, I do manage identify 31 forecasters that appear consistently between 1994 and 2002. Table 2 in the appendix summarizes the important characteristics of the 31 forecasters.

2.2 Actual GDP Growth Data

To measure the accuracy of GDP growth forecasts, economists need actual estimates of GDP growth rates. This task is not as easy as it might seem because the U.S. Department of Commerce produces a number of different estimate of GDP over time as more information about the economy is revealed. The *advance estimate* is released a month after the end of the previous quarter and uses incomplete information about economic activity in that quarter. This estimate is revised and released the next month as a *preliminary estimate* and the month after that as a *final estimate*. In addition, each summer the Commerce Department undertakes annual revisions of data for the previous three years. Finally, benchmark revisions (revisions involving changes in macroeconomic definitions or accounting identities) are made every five years.

In 1991 the Federal Reserve Bank of Philadelphia created the Real Time Research Center which collated real time macro-economic data beginning in November 1965⁵. The primary objective was to give economists and policymakers the opportunity to analyze macro-economic policies and their effects given the information set available at the time (Croushore and Stark 2000, pg. 16). Thus researchers no longer had/ have to dig through old publications of the Commerce Department to obtain advance, preliminary and final estimates of GDP; the Real Time Research Center makes this data available on their website to researchers. The Center continues to produce a new vintage every three months, which contains real time data that would have been available to an economic analyst in the middle of the previous quarter. Each new vintage also contains revised actual data for all previous quarters beginning from 1947:Q1.

⁵ They collected and manually entered the data from reports produced by the Bureau of Economic Analysis.

There has been much controversy between economists and policymakers about which actual estimate- – advance, preliminary, final or latter ones that include benchmark revisions - is a better benchmark for measuring accuracy. Zarnowitz (1992) summarizes the debate most aptly:

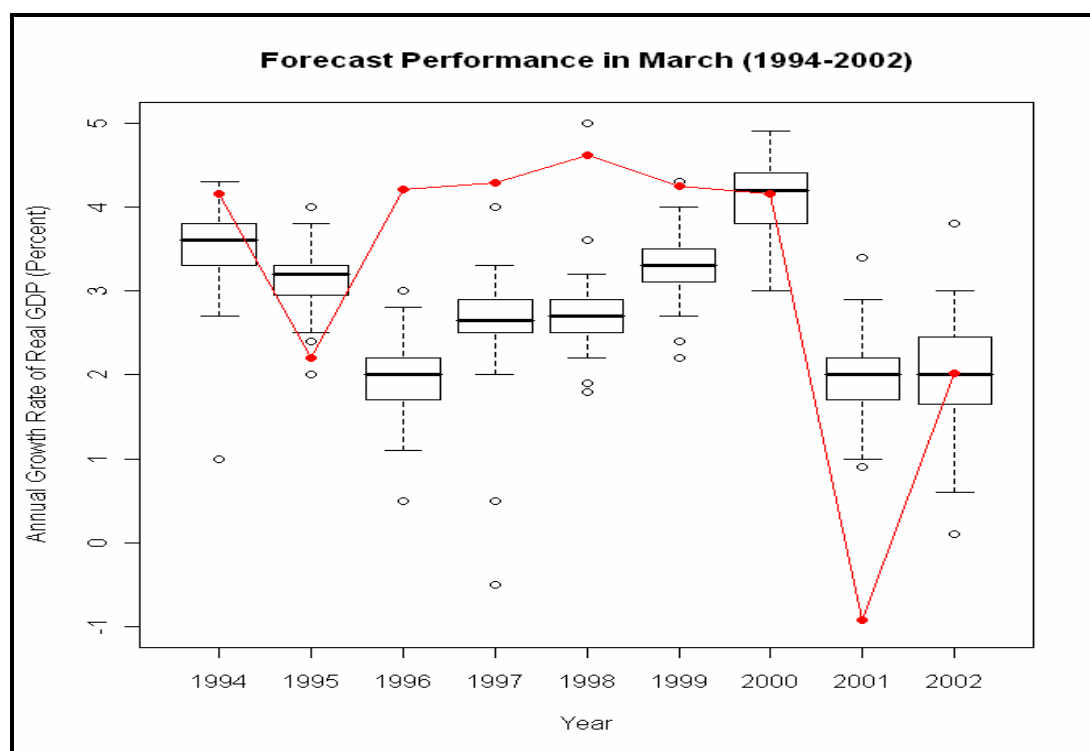
The preliminary figures are most closely related to the latest figures that were available to the forecasters, but they may themselves be partly predictions or “guesstimates” and may seriously deviate from “the truth” as represented by the last revision of the data. On the other hand, the final data may be issued years after the forecast was made and may incorporate major benchmarks. That the forecasters should be responsible for predicting all measurement errors to be corrected by such revisions, is surely questionable.

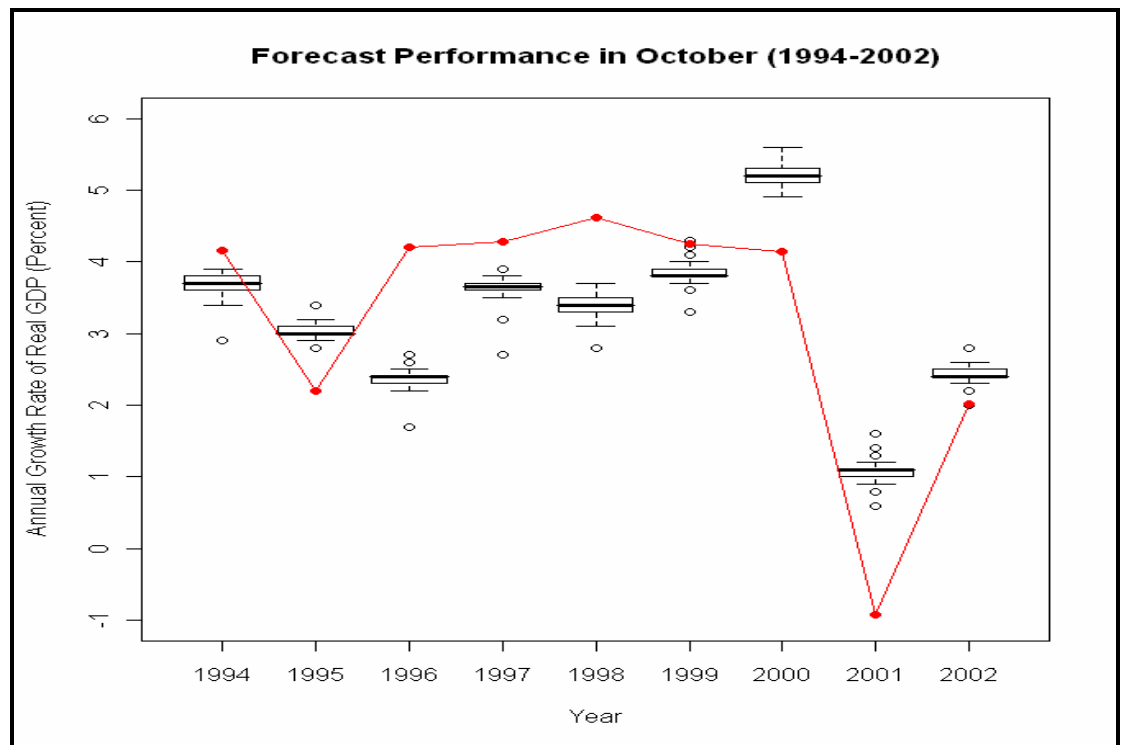
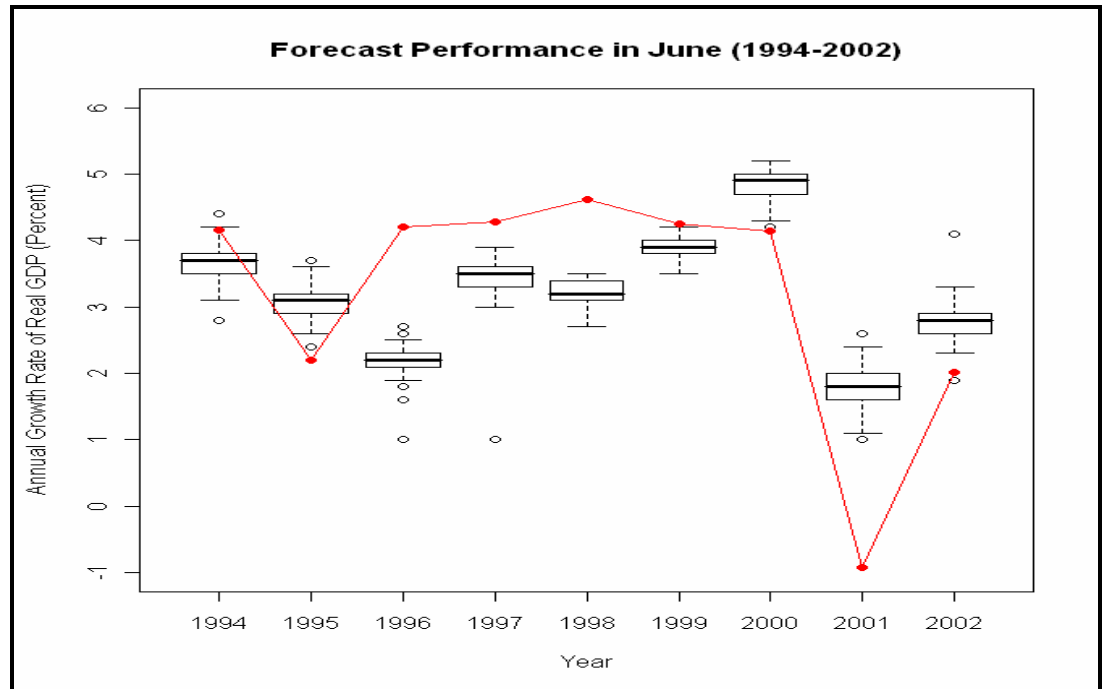
The objective of this paper is to examine the forecasting tendencies of forecasters, which are predicated on the information set available to them. Therefore, it seems most appropriate to judge their performance relative to the advance estimates. Additionally, it seems that forecasts revised several years later distort reality. The latter data undergo (1) revisions using information from income tax records and economic census data collected the year after and (2) benchmark revisions every five years (which, as mentioned earlier, involve changes to macroeconomic definitions or accounting principles). For example, in 1977, the first quarter growth rate as per the May 1977 vintage was 5.2 percent. In August 1979, new information changed the estimate to 8.9 percent. In 1980, a national benchmark revision increased the estimate to 9.6 percent (Croushore and Stark 2000, p19). It would be asking a lot of forecasters for them to foresee the changes in the way the government measures the growth rate of GNP. Therefore, it is best to assess their performance relative to the advance estimates for annual GDP growth.

For the purpose of assessing the accuracy of the Blue Chip panel, Figures 1-3 plot individual forecasts for the annual growth rate of Real GNP for the months of March,

June and October, respectively, from 1994-2002. The graphs also plot the advance estimates for each year as a benchmark to measure forecaster performance. Each figure contains a series of box plots which depict the range of forecasts for the annual growth rate of GNP. The box (minus the tails) provides information about forecasts that fall between the 25th-75th percentile. The dark line inside each box represents the median (50th percentile) forecast. The tip of the top tail represents the highest forecast in the group while the tip of the lower tail represents the lowest forecast in the group. Hence, the greater the length of the box and the further apart the tips of the two tails, the greater the diversity of forecasts. That is, the variance of forecasts is higher.

Figure 1
Blue Chip Forecast Performance by Month for 1994-2002





Two observations about these figures warrant discussion. First, the Blue Chip group persistently under-forecasted GDP growth much like we saw with the Survey of Professional Forecasters as pointed out by Schuh (2001). In particular, each member of

the Blue Chip group under-forecasted GDP growth for four years in a row from 1996 to 1999 when we focus on the 10-month-ahead forecasts made in March. Similarly, each member of the Blue Chip group over-forecasted GDP growth during the 2001 recession. This is a remarkable observation and suggests that it took forecasters several years to figure out that the process driving economic growth during the 1990s had changed.

Second, forecasts made later in the year (October) cluster more around the median and appear to be more accurate. Greater clustering is expected because uncertainty about the state of the economy is reduced as we move through the year and it is expected the less forecaster uncertainty should be associated with less disagreement across forecasters. It is striking that every forecaster still under-forecasted GDP growth in 1996, 1997 and 1998 when forecasting three months before the end of the year. Moreover, the median forecasts for 2000 and 2002 made in October were much less accurate than the median forecasts made in March of that year and the tight clustering around the former suggests that all forecasters over-forecasted 2000 GDP growth by a wide margin in October. Each of these forecasters then proceeded to over-forecast GDP growth in the recession year of 2001.

Likewise, with greater information it is expected that forecasters will be, as a group, more accurate. We see this in the data as the median forecast generally moves closer to the actual GDP growth as we move from Figure 2 to Figure 4. Nevertheless, it

3. Forecast error literature

The process of forecasting is analogous to estimating the amount of a time it takes to drive from A to B. To produce a time prediction, one needs to have a model complete

with parameter estimates which link travel time to its determinants. If the model is imperfect, our predictions are likely to be inaccurate.

For example, when computing traveling time, we must take into consideration factors such as speed, distance, weather, etc. However, it is possible that we might fail to include some important variables or include the wrong variables. These two forms of model misspecification have negative consequences for accuracy.

It might also be the case that a particular factor's relationship with the traveling time changes. Consider the quality of the roads from A to B. In the past, road quality was of little importance because it was guaranteed that the roads were good. However, with the steady deterioration over time, it might take much longer to reach B. This will change the parameter linking the dependent and independent variables (parameter estimation) and lead to greater inaccuracy.

Finally, it is possible that you fail to account for a certain factor which occurs intermittently but has a powerful impact on the traveling time. For example, suppose that on your journey, you encounter a car crash which causes traffic to pileup. In this situation, your traveling time is negatively affected by the unexpected shock.

Similarly, scholars argue that such problems affect forecast accuracy. For a better understanding, assume the model adopts the following form:

$$(1) \quad Y_t = \alpha + \beta Y_{t-i} + \delta X_{t-i} + \varepsilon_t$$

where Y_t represents an $N \times 1$ vector of endogenous variables, X_{t-i} is a $N \times K$ vector of exogenous variables, ε_t is an $N \times 1$ vector of disturbance terms at time t and α , β and δ are parameter matrices. To forecast using this model, the parameter matrices must be

estimated (usually with historical data) and used to project into the future. Hence, we get the following model:

$$(2) \quad \hat{Y}_{t+1} = \alpha + \beta \hat{Y}_{t-i+1} + \delta \hat{X}_{t-i+1}$$

As I mentioned earlier, the forecasting model can be mis-specified in two ways: the inclusion of irrelevant variables and/or the exclusion of relevant variables. Zarnowitz (1997), for example, points out that historically there were two types of business cycle theorists, old and new theorists. The classic (or old) business cycle literature, from the 1890s through the 1960's, focused on "internal dynamics of capitalistic economies: how their component activities interact in successive phases of the process, with what differential timing and intensities and why" (p.4). Contrarily, the more recent models of business cycles rely more on exogenous shocks to explain business cycle fluctuations. If a forecaster uses the new theory to guide his model building and ignores some of the relevant variables or relationships suggested by the older theory, his model will produce forecasting errors that may be systematic.

Another example of model misspecification is the exclusion of The Index of Leading Indicators from X_{t-i} . Its omission could have negative consequences for forecast accuracy because the leading indicators might contain important market information. Batchelor and Dua(1998) examine the effect of including the Consumer Confidence Index (CCI) on forecast accuracy for Real GNP forecasts produced by US forecasters over the previous decade. Using Blue Chip data, the authors find that forecasts would have made smaller errors if they had used the CCI to modify their forecasts during the 1991 recession. However, they conclude that their results cannot be generalized, i.e.

forecasters cannot exploit the relationship between forecast error and the CCI in other forecasting periods.

Another source of forecast error concerns incorrect parameter estimation, namely estimates for α , β and δ matrices. It is evident from the diversity of forecasts that forecasters differ in their parameter estimates. This result might be a function of ideology and technique. Batchelor and Dua (1990) study the effect of different ideologies on the accuracy of forecasts. Using data they gathered through surveys of the Blue Chip forecasters, they find that no one ideology or technique produced consistently better estimates than others⁶. Unfortunately, their ambiguous result might be partially driven by the broadness of their definitions of ideologies and techniques. On this matter of parameter estimation, Zarnowitz (1997) brings to our attention that there is growing evidence in favor of the possibility that relationships between variables are not linear, not constant, but continuously changing during different phases of the business cycle. As the old business cycle theorists tried to argue, “booms generate excesses and imbalances that tend to be reduced in slowdowns and moderate recessions” (p. 4)

Finally, forecast errors could be exacerbated by unpredictable exogenous shocks. According to David Hendry and Neil Ericsson (2001), unanticipated shocks could result in (1) deterministic shifts or (2) stochastic changes. In the case of (1), we would witness a change in the mean of the dependent variable. For example, earthquakes or any such spontaneous deviations from the norm would cause the average annual growth rate of GDP to change drastically. Alternatively, stochastic shifts would change the time varying

⁶ To determine the ideologies and techniques used by Blue Chip forecasters, the authors requested Bob Eggert to include certain questions in his survey. In particular, the authors had Eggert ask forecasters how they would classify themselves (e.g. Keynesian, Monetarist) and which forecasting technique they preferred when forecasting (1) Real GNP, (2) Consumer price index and (3) Treasury bill rate.

error term incrementally but persistently. Both types of shocks are reflected in the error term of equation (1) and cause the predicted and actual outcome to differ even if the model is properly specified and the parameters are accurately estimated

As evident from the discussion above, scholars believe that inaccurate forecasts are driven mainly by three factors, namely model misspecification, incorrect parameter estimation and unpredictable shocks to the error term⁷.

4. Herding

The preceding discussion assumed that forecasters generate forecasts by building a model and using it to project the future independent of what others are forecasting. However, the social learning literature suggests that humans rarely form expectations in a social vacuum and that they often take queues from those around them. In some circumstances, listening to others can help one improve his or her forecast. On the other hand, the more recent literature on information cascades and herding suggests that utilizing the information contained in the forecasts of others might lead a forecaster astray and cause them to produce inaccurate forecasts. This examination of the link between information-based herding and forecast accuracy is exactly what allows this paper to make a unique contribution.

There are two main types of herding- reputation-based herding and information-based herding. Reputation-based herding occurs when forecasters are concerned about more than simply the accuracy of their predictions. According to Lamont (1995), there is a principal-agent problem wherein the agenda of the forecaster is different from that of the consumer of their forecasts. The forecaster is motivated by an incentive structure that

⁷ Some scholars argue that asymmetric loss functions and prior probabilities might lead forecasters to misforecast. See Stekler(1972) and Stekler and Schnaeder(2003).

rewards him on the basis of his reputation. His reputation is based upon two things, namely (1) how accurate he is and (2) how close or far he is with respect to the consensus. The latter is important when the consumers of forecasts cannot judge the quality of a forecaster solely on their short-term performance and they perform in an environment where “bright minds think alike.” Lamont finds that earlier in their careers, forecasters care more about reputation and tend to forecast close to the consensus. But as these young forecasters become older, they care less about reputation and make more radical forecasts. Given that mean forecast of a group tends, over time, to be more accurate than most forecasts of individual group members, one result of this is that older forecasters have relatively greater inaccuracy.

The second type of herding is information-based herding. Unlike reputation herding, an individual is said to herd informationally if he chooses to mimic the actions of others regardless of his private information signal. The motivation for this type of herding is not to fool or mislead the consumer of the forecasts, but simply to increase forecast accuracy by exploiting the information contained in the forecasts of others. Forecasting literature such as Batchelor and Dua(1992) and Ferderer, Pandey and Veletsianos(2005)⁸ examines the information herding tendencies of macro-forecasters, but does not address whether herding affects accuracy.

To illustrate the impact of information-based herding, I rely on Bikhchandani and Sharma (2001). The scholars explain the model by using an example involving investors investing in a certain stock. They set certain rules and conditions to simplify the model.

⁸ Batchelor and Dua(1992) examines how important the lagged individual consensus and the mean consensus forecast are to forecasters. They find that forecasters tend to move toward the consensus because they are conservative and not because they are herding. Ferderer, Pandey and Veletsianos(2005) examine whether forecasters herd and if so, which model, information herding or reputation herding, better explains their herding tendencies.

First, for each of the N investors, the value of the outcome relative to that of the next best investment, V , can be either $+1$ or -1 with equal probability. Second, investors must make their investment decisions in a sequential order, which is determined exogenously. The sequence creates an opportunity for investors to observe and (maybe) be influenced by the actions of their predecessors. Third, aside from being privy to the actions of others, each investor also receives his own private information signal about the stock. The signal can be either good (G) or bad (B). Finally, if $V=+1$, the probability of getting a good signal is p and the probability of getting a bad signal is $1-p$ where $0.5 < p < 1$. As p rises, the signals are more informative.

Now, let Investor 1 act first. If he receives signal G (B) then he will invest (not invest). Since the first investor has no one to observe, he must follow his private signal. Investor 2 moves next. If he receives signal G and he sees that Investor 1 has invested, he will invest. For Investor 2, the actions of Investor 1 will confirm the veracity of his private signal G. However, if he receives signal B and he sees that Investor 1 invested, he will be indifferent since there is an equal probability that either one of them is right. For the sake of simplicity, let us assume he received G and invested accordingly. Now it is Investor 3's move. He sees that both 1 and 2 invested. From this he will infer (maybe incorrectly) that both received signal G. If he has a signal B, he will ignore it because he will assume that the likelihood that two people before him were wrong is less than the likelihood that his signal is wrong (i.e. the superiority of the wisdom of the crowd)⁹. Mathematically, Investor 3 will think of the probability of being wrong despite getting signal G as the probability that $V=-1$ despite an information signal G is $(1-p)$. In some

⁹ If has signal G, he has no reason not to follow the first two investors since the probability of $V=+1$ given a good signal is greater than 0.5.

cases, this will lead him to make the correct decision. However, arriving at the correct outcome is conditional on the fact that the previous investors are guided by informative private signals. If not, Investor 3, along with all the investors who are acting on *his* decision, will converge to the wrong target. Such a cascade of incorrect predictions is known as a negative information cascade.

The information cascade model offers a foundation for the econometric models to follow. In the next section, I discuss the models that I use to test for information-based herding and its effect on forecast accuracy, the basic econometric obstacles that arise when dealing with panel data, ways to overcome the obstacles, and the results.

5. Econometric Analysis

Testing for herding and its effect on accuracy is a two stage process. The first stage involves quantifying the herding tendencies of forecasters. The second stage involves testing the correlation between the quantifiable measure of herding and a measure of forecast accuracy.

5.1 Measuring Herding

To measure herding behavior across the 1994-2002 period, I follow Gallo, Granger and Jeon(2002). Gallo, Granger and Jeon(2002) formulate an econometric model for testing herding behavior. They argue that a forecaster's forecast is built on three factors: (1) a persistence in one's own most recent forecast (the lagged individual forecast), (2) an imitation effect of the average belief expressed in the previous period by the group (the lagged consensus¹⁰ or the "wisdom of the crowd" variable) and (3) an

¹⁰ Since Blue Chip forecast data is monthly, the lagged consensus is the average forecast of all forecasters in the previous month. January is the only month for which the lagged consensus cannot be calculated.

effect due to the desire to move closer together as the time-horizon advances (Granger p. 12). Due to the fact that forecasters forecast simultaneously and not sequentially, they observe the actions of their fellow with a lag of one period.

The following econometric model is derived from Gallo, Granger and Jeon (2002):¹¹

$$(3) \quad y_{T,j}^i = \alpha + w_1^i y_{T,j+1}^i + w_2^i \hat{y}_{T,j+1} + \mu_{T,j}^i$$

for $j = 11, 10, \dots, 1$. The variable $y_{T,j}^i$ is the forecast in year T of forecaster i for j periods ahead, $y_{T,j+1}^i$ is the most recent forecast produced by forecaster i, $\hat{y}_{T,j+1}$ is the lagged consensus forecast, and $\mu_{T,j}^i$ is the disturbance term which captures other information used in the forecast other than the lagged individual forecast and consensus forecast. Lastly, w_1^i and w_2^i represent the weights attached by forecaster i to the most recent forecast and the lagged consensus respectively. The latter is the key parameter to estimate because it provides a measure of the degree to which each forecaster puts weight on the forecasts of others when updating their own individual forecast.

Before discussing the results, it is important to address certain issues that arise when dealing with panel data. One such obstacle in panel data estimation is serially correlated errors. With Blue Chip data, serial correlation is unavoidable due to the fact that forecasters continuously incorporate old information into their estimates. However, model (3) solves the problem of serial correlation by including (1) the individual lagged forecast from one period back which captures all previous information since January and (2) the lagged consensus forecast which also captures past information. Another concern is the presence of heteroskedasticity. To test for heteroskedasticity, I use one cross

¹¹ In my model, I include the most recent forecast and the lagged consensus but do not include the forecast variance.

section (data from any year) of my panel data set. The test statistic for heteroskedasticity is not significant at the 5 percent level.

Now, Table 3 on page 18 displays the herding propensities (w_2^i) and the reliance on the lagged individual forecast (w_1^i) for 31 forecasters over the 1994-2002 sample period. One interesting result is that of the 22 forecasters that herd significantly, the coefficients on lagged consensus range from a maximum of 1.099 to a minimum of 0.191. This suggests that some forecasters derived the forecast at time T almost entirely from the lagged consensus while other forecasters placed very little weight on the lagged consensus. Another curious, though insignificant, result is the negative coefficient on lagged consensus (-0.051) for forecaster 70. This implies that the forecaster actually *deviated* from the consensus rather than herding to it. Overall, the results suggest that forecasters (1) engage in information herding and (2) herd to the lagged consensus.

The results in Table 3 are informative of the general herding behavior of forecasters over the entire sample 1994-2002 but they bring to bear the question of whether forecasters constantly herd with the same intensity. A priori, there is little reason to believe that the tendency to herd will remain the same each year. For example, the level of difficulty involved in forecasting might differ each year and this could alter the incentive to mimic the forecasts of others. Consider 2001. With the occurrence of the recession, it might have been much harder to get an informative private signal to forecast the growth rate (i.e., forecasters were more uncertain). Therefore, forecasters might herd more to exploit the wisdom of the crowd.

Table 3
Forecaster Herding to Lagged Consensus 1994-2002

ID	Lagged Consensus ($\hat{y}_{T,j+1}^i$)	T-stat (pvalue)	Lagged individual forecast ($y_{T,j+1}^i$)	T-stat (pvalue)	R-squared
5	0.068	0.56 (0.57300)	0.930	8.19 (0.00000)***	0.932
6	0.393	2.78 (0.00700)***	0.603	4.26 (0.00000)***	0.933
17	0.347	1.83 (0.07000)*	0.658	3.61 (0.00100)***	0.927
25	0.217	2.30 (0.02400)**	0.774	8.02 (0.00000)***	0.907
26	0.295	4.92 (0.00000)***	0.734	12.27 (0.00000)***	0.953
31	0.746	5.70 (0.00000)***	0.255	2.01 (0.04800)**	0.948
37	0.353	3.01 (0.00300)***	0.672	6.09 (0.00000)***	0.930
39	0.530	6.49 (0.00000)***	0.476	5.63 (0.00000)***	0.955
41	0.107	1.61 (0.11000)	0.939	13.46 (0.00000)***	0.966
43	0.191	2.22 (0.02900)**	0.816	9.45 (0.00000)***	0.921
48	1.099	9.45 (0.00000)***	-0.085	-0.74 (0.46000)	0.957
50	0.715	7.03 (0.00000)***	0.267	2.55 (0.01200)**	0.939
52	0.499	5.87 (0.00000)***	0.525	6.98 (0.00000)***	0.935
57	0.659	6.36 (0.00000)***	0.325	3.20 (0.00200)***	0.918
70	-0.051	-0.30 (0.76400)	1.054	6.17 (0.00000)***	0.945
73	0.103	0.74 (0.46000)	0.901	6.83 (0.00000)***	0.938
79	0.050	0.58 (0.56600)	0.933	12.39 (0.00000)***	0.946
82	0.817	8.71 (0.00000)***	0.254	3.04 (0.00300)***	0.961
84	0.426	2.90 (0.00500)***	0.577	3.82 (0.00000)***	0.932
85	0.880	5.28 (0.00000)***	0.117	0.72 (0.04730)	0.933
89	0.020	0.12 (0.90100)	0.962	6.27 (0.00000)***	0.927
91	0.250	1.13 (0.26200)	0.710	3.60 (0.00100)***	0.907
97	0.202	4.13 (0.00000)***	0.814	14.14 (0.00000)***	0.953
98	0.814	5.31 (0.00000)***	0.220	1.49 (0.14100)	0.948
105	0.583	5.57 (0.00000)***	0.426	3.99 (0.97000)***	0.930
108	0.542	4.58 (0.00000)***	0.503	4.31 (0.00000)***	0.955
111	0.871	7.25 (0.00000)***	0.145	1.26 (0.21200)	0.966
112	0.451	3.90 (0.00000)***	0.598	5.76 (0.00000)***	0.921
114	0.234	2.91 (0.00400)	0.783	11.67 (0.00000)***	0.957
121	0.781	7.03 (0.00000)***	0.193	1.68 (0.09600)*	0.939
122	0.256	1.38 (0.17100)	0.748	4.15 (0.00000)***	0.935

Notes: Dependent variable is the forecast made by the respective forecaster at time T, j months ahead of December. Significance at ten percent, five percent and one percent level is denoted by *, ** and *** respectively.

To test for time-varying (year-specific) coefficients on the lagged consensus, I measure the herding tendency of each forecaster in each year from 1994-2002. I utilize model (3) to conduct the regression but I run it separately for each year. Table 4 on the next page presents the results.

The results suggest that forecasters do not herd with the same intensity over time. In fact, it is interesting to observe the range of parameters for each forecaster. Most of the forecasters have minimum herding parameters below 0. This implies that in certain years they are actually deviating away from the lagged consensus. Similarly, many forecasters have maximum herding parameters above 1. This is not as surprising as it suggests that forecasters place great emphasis on the lagged consensus for information. Finally, with the exception of forecaster 89, all forecasters herd on average to the lagged consensus.

5.2 Measuring the impact of herding on forecast accuracy

The second stage of the analysis focuses on the relationship between the herding propensity and forecast accuracy. The forecasting literature provides a few ways to measure accuracy, namely the absolute error, squared error and mean squared error.¹² Initially, I employ the mean squared error because it penalizes large errors more than the absolute error does.

Figure 2 shows the scatter plot between the herding parameter and the mean squared error which are both estimated over the entire sample (1994-2002). The plot

¹² These three measures are calculated differently. Absolute error is the absolute difference between the actual and the estimated values. Squared error is the square of the difference between the actual and the estimated values. Mean squared error is the average squared error.

focuses on the 10-month horizon in March. I also examine the relationship in June and October but do not include those plots in the text¹³. I focus on these different horizons

¹³ The plots for June and October can be found in the appendix.

Table 4
Time-variant Forecaster Herding to Lagged Consensus

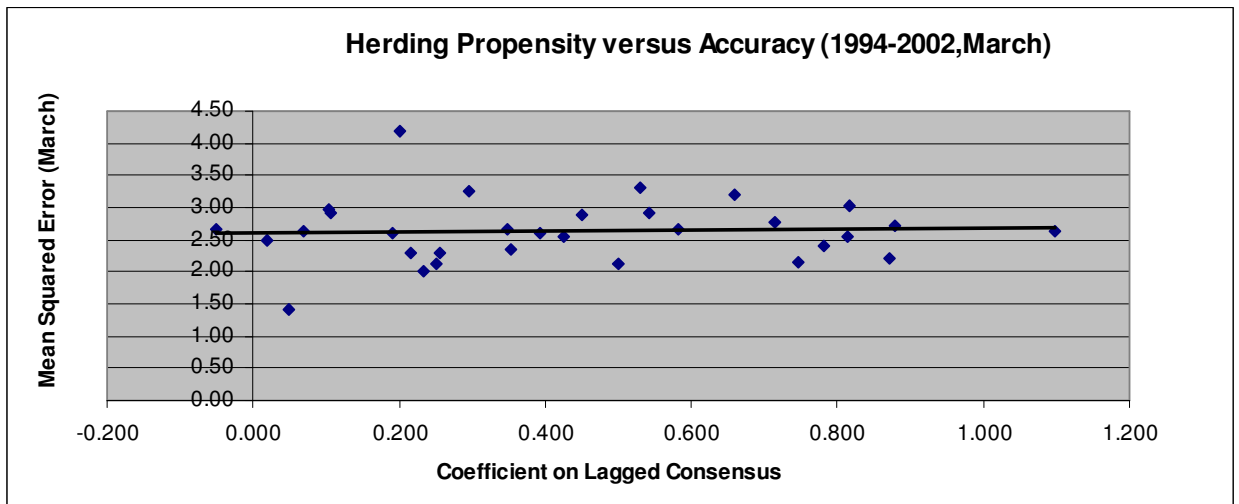
ID	Maximum Coefficient on Lagged Consensus	Minimum Coefficient on Lagged Consensus	Mean Coefficient on Lagged Consensus
5	1.250	-0.934	0.490
6	0.777	-0.042	0.389
17	1.511	-0.063	0.673
25	0.799	-1.151	0.152
26	1.365	-0.434	0.425
31	1.075	0.102	0.489
37	1.239	0.153	0.745
39	1.500	-0.169	0.720
41	0.688	-0.666	0.121
43	0.480	-0.897	0.107
48	2.010	0.008	1.061
50	1.540	-1.169	0.753
52	1.128	-0.249	0.469
57	1.737	-1.148	0.688
70	0.921	-0.836	0.010
73	1.814	-0.883	0.662
79	0.750	-0.045	0.337
82	2.460	-0.027	1.152
84	1.308	-0.385	0.425
85	1.680	-0.536	0.492
89	0.691	-1.797	-0.023
91	1.820	-0.422	0.566
97	1.322	-0.275	0.523
98	1.595	-1.225	0.859
105	1.499	0.209	0.711
108	1.607	0.381	0.928
111	4.165	0.127	1.013
112	0.724	-0.983	0.169
114	1.302	-0.431	0.360
121	1.544	-0.108	0.744
122	1.515	-0.258	0.523

Notes: The table considers herding measures for each year for each forecaster regardless of whether it is significant at the ten, five or one percent level.

because they allow me to consider how systematic differences in uncertainty about the economy at various points in the year affect the result. That is, on average there is more forecast uncertainty in March than in the other months and this has the potential to influence the propensity to herd and the link between herding and forecast errors.

There are two things to note. First, though all three graphs seem to suggest that there is a slight positive relationship between the propensity to herd and the size of the forecast error, we cannot draw any firm conclusions since the best-fit line is insignificant. Secondly, as forecasters get closer to the end of the year, individual errors on average seem to decrease as does within group variance. However, once again the estimated slopes of the best-fit lines are insignificant.

Figure 2
Correlation between Herding and Accuracy in March (1994-2002)



To econometrically test the relationship between herding to the lagged consensus and forecast accuracy for three horizons in the 1994-2002 period, I utilize the following model:

$$(4) \quad MSE^i = \alpha + bw^i + \mu^i$$

where MSE_i is the mean squared error of forecaster i , w^i is the herding parameter for forecaster i calculated in by model (3), and μ^i is the error term for forecaster i . Lastly, b is the coefficient on the herding parameter. In this case, there will be 31 forecasters, each having a mean squared error and a herding parameter.

Table 5 on the next page presents the results for model (4). First, the relationship between herding and accuracy is positive but *insignificant*. This confirms that we could not draw any concrete conclusions from the correlation graphs in Figure 2. Secondly, the steady decline in the coefficient of the constant term from March to October confirms that the forecast error on average decreases as the year passes.

Table 5
Relationship between Herding and Accuracy in March, June and October (1994-2002)

Variables	March	June	October
Herding	0.0832 (0.27)	0.2003 (1.07)	0.0337 (0.39)
Constant	2.5952 (16.27)***	1.7579 (17.92)***	1.2896 (28.47)***
R-squared	0.0026	0.0382	0.0052

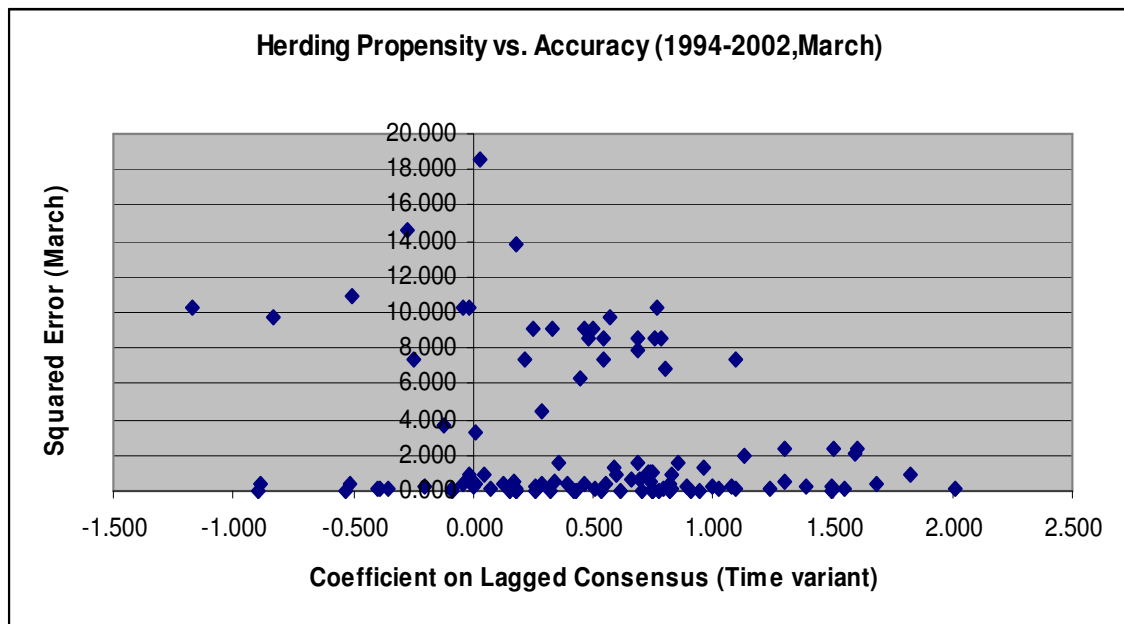
Notes: There are 31 observations for each forecaster. T-statistics are indicated in parentheses. The ten, five and one percent levels of significance are denoted by *, ** and *** respectively.

As we saw above, the herding parameters varied a great deal over time. Given this, it might be the case that the relationship between herding and accuracy might also differ over time. As the literature on information cascades shows, people can get caught up in a positive cascade where they converge to the correct target and this affects accuracy. More importantly, they can also get caught up in a negative target where they mimic their predecessors and converge on the wrong target. To measure forecast

accuracy on a yearly basis, I use the squared error rather than the mean squared error because I only have one error per year.

Figure 3 on the next page plots yet again the correlation between the propensity to herd and forecast error. However, unlike in Figure 2, I utilize the time varying propensities to herd (like those displayed in Table 2) and squared errors for the 1994-2002 period. Though Figure 3 only displays the plot for March, I also plot the relationship for June and October¹⁴. From the plot, it seems that the best-fit line has a negative slope. This suggests that as forecasters herd more, their accuracy increases. This result is at odds with the theory which suggests that increased herding will lead to increased inaccuracy. Once again, we must do further analysis to draw any concrete inferences about the relationship.

Figure 3
Correlation between Herding and Accuracy in March (1994-2002)



¹⁴ The plots for June and October can be found in the appendix.

To examine the impact of herding on forecast accuracy more closely, I examine the relationship between herding propensities and squared errors in various sub-samples such as 1994-2000 and 1994-2002. I utilize the following econometric model:

$$(5) \quad S_T^i = \alpha + bw_T^i + X_T + \mu_T^i$$

where S_T^i represents the squared error forecaster i in year T , w_T^i represents the herding propensity for forecaster i (the time varying coefficient on the lagged consensus) in year T , and X_T represents vector dummies for each year depending on the sub-sample. Finally, μ_T^i represents the error term for forecaster i in year T . This regression is estimated using data from the three months: March, June and October. The motivation for including year dummies is that the difficulty of forecasting differs across years. If the dummies are significant, it would support the decision to investigate the possibility of time varying accuracy.

Tables 6-8 display the results for March, June and October respectively. If you look closely the bolded results for the 1994-2002 sub-sample in Table 5, you will see that the year dummies from 1996-1999 are significant. Each of those coefficients is the difference between the intercept in the specific year and the intercept in the base year (in this case 1994). This suggests that the squared error is differs across years. That is, the difficulty of forecasting differs across years. It is interesting to note that the largest year dummy coefficient is D01. This might be because it was tough to forecast the growth rate in a recession year.

In the 1994-2000 period, these year dummies are significant yet again. This is not surprising since the base year, 1994, has not changed. Once again, we see that 2001 is the year with the largest year dummy coefficient.

The most interesting part of the result is positive and significant coefficient for Herding Propensity (0.191). This implies that when an individual forecaster increases their propensity to herd by 0.01, their accuracy decreases by 0.191. This is in accordance with the theory of herding and negative information cascades.

Finally, none of the coefficients on Herding Propensity in the other sub-samples is significant. The coefficients for Herding in 2001-2002 and 2001 are negative although insignificant. The year dummies continue to remain significant. However, the base year for the remaining samples is either 2001 or 2002.

Table 6
Relationship between Herding and Accuracy in March for various periods in 1994-2002 sample

Variables	1994-2002	1994-2000	2001-2002	2001	2002
Herding Propensity	.0726004 (0.60)	.1909281 (2.20)**	-.4651842 (-0.91)	-1.396736 (-1.25)	.0650464 (0.50)
D94	-	-	-	-	-
D95	.5357633 (1.65)	.5595128 (2.62)***	-	-	-
D96	5.029182 (15.48)***	5.04676 (23.71)***	-	-	-
D97	2.220381 (6.83)***	2.200529 (10.33)***	-	-	-
D98	3.184481 (9.72)***	3.166435 (14.75)***	-	-	-
D99	.6625113 (2.00)**	.6435207 (2.97)***	-	-	-
D00	-.2905294 (-0.87)	-.3004283 (-1.38)	-	-	-
D01	8.533079 (25.70)***	-	-	-	-
D02	-.0572873 (-0.16)	-	-8.463015 (-13.09)***	-	-
Constant	.3917868 (1.64)	.325153 (2.06 **)	9.05894 (20.31)***	9.291184 (14.74)***	.3381716 (2.96)***
R-squared	0.8306	0.8300	0.7853	0.0545	0.0113

Note: Significance at the ten, five and one percent is denoted by *, **, and *** respectively. T-statistics are provided in parentheses below the coefficients. The constant (intercept) represents the intercept for non-independent firms and for the year 1994 for the samples 1994-2002, 1994-2000. For 2001-2002, the constant (intercept) represents the intercept for non-independent firms and the year 2001. In the last two samples, the constant represents only the intercept for the non-independent firms and the specific year.

Table 7 displays the results for June. Once again, the coefficient of interest is the coefficient on Herding Propensity in the sub-sample 1994-2000 (.0803395) which remains positive and significant at the 10 percent level. However, the degree of correlation between herding propensity and accuracy diminishes relative to what we saw in March. In June, if an individual's propensity to herd increases by 1 unit, his accuracy will decrease by .0803395 percent on average.

The year dummies continue to remain significant in the 1994-2002 and the 1994-2000 periods and are still to be interpreted as differences from the base year 1994. Once again, the intercept difference in 2001 is the largest.

Table 7
Relationship between Herding and Accuracy in June for various periods in 1994-2002 sample

Variables	1994-2002	1994-2000	2001-2002	2001	2002
Herding Propensity	.0021878 (0.03)	.0803395 (1.91)*	-.298167 (-1.18)	-.7621155 (-1.31)	-.0576031 (-0.60)
D94	-	-	-	-	-
D95	.4913921 (2.84)***	.5070779 (4.93)***	-	-	-
D96	3.671884 (21.24)***	3.683493 (35.88)***	-	-	-
D97	.4335496 (2.51)**	.4204377 (4.09)***	-	-	-
D98	1.707607 (9.88)***	1.696686 (16.53c)***	-	-	-
D99	-.0756751 (-0.43)	-.0854528 (-0.82)	-	-	-
D00	.2765204 (1.57)	.2728169 (2.62)***	-	-	-
D01	7.491966 (42.82)***	-	-	-	-
D02	.3209745 (1.81)*	-	-7.127628 (-21.76)***	-	-
Constant	.2439543 (1.92)*	.1999448 (2.62)**	7.820006 (32.99)***	7.949891 (23.00)***	.5902998 (7.15)***
R-squared	0.9262	0.9093	0.8984	0.0576	0.0138

Note: Significance at the ten, five and one percent is denoted by *, **, and *** respectively. T-statistics are provided in parentheses below the coefficients. The constant (intercept) represents the intercept for non-independent firms and for the year 1994 for the samples 1994-2002, 1994-2000. For 2001-2002, the constant (intercept) represents the intercept for non-independent firms and the year 2001. In the last two samples, the constant represents only the intercept for the non-independent firms and the specific year.

Lastly, as was the case in March, the coefficient on Herding Propensity in all the other sub-samples remains insignificant. However, unlike the results from March, the coefficients for Herding Propensity in all sub-samples other than 1994-2000 are negative.

Finally, the results in October are displayed in Table 8. The most important change to note here is the lack of significance of the coefficient on Herding Propensity in the 1994-2000 sample. Though the coefficient is positive, it is no longer significant. Therefore, we cannot make any firm inferences about the relationship between herding and accuracy in October.

Table 8
Relationship between Herding and Accuracy in October for various periods in 1994-2002 sample

Variables	1994-2002	1994-2000	2001-2002	2001	2002
Herding Propensity	-.0058398 (-0.19)	-.0247516 (-0.89)	.0766829 (0.69)	.3061751 (1.19)	-.0427426 (-1.42)
D94	-	-	-	-	-
D95	.451086 (5.28)***	.4472902 (6.54)***	-	-	-
D96	3.090455 (36.25)***	3.087646 (45.21)***	-	-	-
D97	.1904959 (2.23)**	.1936688 (2.83)***	-	-	-
D98	1.338459 (15.57)***	1.341344 (19.48)***	-	-	-
D99	-.0521618 (-0.60)	-.0551197 (-0.80)	-	-	-
D00	.8927641 (10.31)***	.8940515 (12.89)***	-	-	-
D01	3.813034 (43.76)***	-	-	-	-
D02	-.0390205 (-0.44)	-	-3.86863 (-27.73)***	-	-
Constant	.2323531 (3.70)***	.2430029 (4.79)***	4.024813 (40.43)***	3.967599 (27.46)***	.2099451 (8.15)***
R-squared	0.9441	0.9393	0.9366	0.0500	0.0743

Note: Significance at the ten, five and one percent is denoted by *, **, and *** respectively. T-statistics are provided in parentheses below the coefficients. The constant (intercept) represents the intercept for non-independent firms and for the year 1994 for the samples 1994-2002, 1994-2000. For 2001-2002, the constant (intercept) represents the intercept for non-independent firms and the year 2001. In the last two samples, the constant represents only the intercept for the non-independent firms and the specific year.

With regard to the year dummies, they continue to remain significant in the 1994-2002 and 1994-2000 sample. Similarly, the coefficients on herding propensity in the remaining sub-samples remain insignificant.

6. Conclusion

In this paper, I wished to address two questions. First, do forecasters in the Blue Chip panel systematically under or over-estimate the growth rate of Real GDP in the period 1994-2002? The motivation for this question arises from the literature, specifically from Schuh (2001). In his paper, Schuh plots forecaster performance from 1968 to 2000 against the actual growth rate and finds that participants in the Survey of Professional Forecasters systematically underestimate GDP growth for three years, 1996-1999. Using data for 31 forecasters from the *Blue Chip Economic Indicators* newsletter and the Real Time Research Center, I observe that the same phenomenon with Blue Chip forecasters. Additionally, I find that the Blue Chip panel overestimates in 2000-2002. This leads to the second question: how can we explain the panel's persistent under-estimation in the late 1990s and the overestimation in the early 21st century?

To answer the second question, I argue that forecasters herd to the “wisdom of the crowd” and this negatively affects their accuracy. I test this hypothesis econometrically by (1) quantifying the herding propensities of forecasters and (2) determining their impact on accuracy. For part (1), I find that the general propensity to herd for 22 of the 31 forecasters is significant at the 10, 5 or 1 percent level. However, it seems unreasonable to assume that forecasters herd with the same intensity over time. Upon further testing, I determine that the propensity to herd varies over time.

With respect to (2), I test the relationship two ways. First, I model the time invariant propensity to herd and mean squared error and find that a positive but insignificant relationship. To examine it more closely, I decide to test year specific propensities to herd with squared errors. In this regression, I also include year dummies to examine whether the difficulty of forecasting varies over time. I find that there is a positive and significant relationship between the herding propensity and squared errors in March and June over the sample 1994-2000. This suggests that if a forecaster herds more, their accuracy will go down. I also find that most of my year dummies are significant for the various sub-samples that I test (e.g. 1994-2002, 1994-2000, etc.).

On a final note, I believe that there are several avenues for further research. One possible course of research would be to examine whether forecast inaccuracy between 1996 and 1999 was caused by a combination of factors, one of which was herding. The late 1990s is known as the boom of the new economy due to the dot-com bubble. Another possible avenue might be to examine whether herding and forecast accuracy by industry. As I mention in my paper, literature such as Laster et al (1996) argue that forecasts by firms in some industries are driven by motivations other than accuracy. This hypothesis has yet to be applied to the problem of information-based herding and forecast accuracy.

Bibliography

- Banerjee, Abhijit V. 1992. A simple model of herd behavior. *The Quarterly Journal of Economics* 507: 797-817.
- Batchelor, Roy and Pami Dua. 1998. Improving macro-economic forecasts: The role of consumer confidence. *International Journal of Forecasting* 14: 71-81.
- Batchelor, Roy and Pami Dua. 1990. Forecaster ideology, forecasting technique, and the accuracy of economic forecasts. *International Journal of Forecasting* 6: 3-10.
- Batchelor, Roy and Pami Dua. 1992. Conservatism and Consensus-seeking among Economic Forecasters. *Journal of Forecasting* 11: 169-181.
- Bikhchandani, Sushil and Sharma, Sunil. 2001. Herd Behavior in Financial Markets. *IMF Staff Papers* 47: 279-310.
- Croushore, Dean and Tom Stark. 2000. A Funny Thing Happened On the Way to the Data Bank: A Real- Time Data Set for Macroeconomists. *Business Review*, Federal Reserve Bank of Philadelphia. 15-27.
- Ferderer, Peter J., Bibek Pandey and George Veletsianos. 2005. Why do the Blue Chip forecasters herd? A test of information and reputation models.
- Gallo, Giampiero M., Clive W.J. Granger and Yongil Jeon. 2002. Copycats and Common Swings: The Impact of the Use of Forecasts in Information Sets. *IMF Staff Papers* 49: 4-21.
- Hendry, David F. and Neil R. Ericsson. 2001. *Understanding Economic Forecasts*. Cambridge: Massachusetts Institute of Technology Press.
- Lamont, Owen. 1995. Macroeconomic Forecasts and Microeconomic Forecasts. Working Paper No. 9617, *NBER Working Papers Series*. 1-33.

- Laster, David, Paul Bennett, and In Sun Geoum. 1996. Rational Bias in Macroeconomic Forecasts. Working Paper No. 9617, *Working Paper*, Federal Reserve Bank of New York.
- Laster, David, Paul Bennett and In Sun Geoum. 1997. Rational Bias in Macroeconomic Forecasts. *Federal Reserve Bank of New York Staff Reports No. 21*, Federal Reserve Bank of New York.
- McNees, Stephen K. 1992. How Large Are Economic Forecast Errors? *New England Economic Review*: 25-42.
- Schuh, Scott. 2001. An Evaluation of Recent Macroeconomic Forecast Error. *New England Economic Review*: 35-56.
- Stekler, Herman O. 2003. Improving our ability to predict the unusual event. *International Journal of Forecasting* 19: 161-163.
- Schander, M. N. and Herman O. Stekler. 1998. Sources of turning point errors. *Applied Economic Letters* 5: 519-521.
- Zarnowitz, Victor. 1992. Twenty-Two Years of The NBER-ASA Quarterly Economic Outlook Surveys: Aspects and Comparisons of Forecasting Performance. Working Paper No. 3965, *NBER Working Paper Series*. 1-52.
- Zarnowitz, Victor. 1997. Business Cycles Observed and Assessed: Why and How They Matter. Working Paper No. 6230, *NBER Working Paper Series*: 1-30.

APPENDIX

Table 2
Information on 31 forecasters in 1994-2002 sample

Forecasting Firm	ID	Industry Classification	Years in Blue Chip Sample for Current year forecast
Bank of America Corp.	5	Commercial Bank	Jan 1977- Jun 2003
Bank One	6	Commercial Bank	Apr 1981- Jun 2003
Chamber of Commerce	17	Other	Feb 1978- Jun 2003
Comerica	25	Commercial Bank	Jan 1990- Jun 2003
Conference Board	26	Other	Jan 1977- Jun 2003
Daimler Chrysler AG	31	Industrial Corporation	Jan 1984- Jun 2003
DuPont	37	Industrial Corporation	Jan 1977- Jun 2003
Econoclast	39	Independent	Jan 1984- Jun 2003
Eggert Economic Enterprises	41	Other	Jan 1977- Jun 2003
Evans, Carrol and Associates	43	Independent	Jan 1980- Nov 2002
Ford Motor Company	48	Industrial Corporation	Aug 1989- Jun 2003
General Motors Corporation	50	Industrial Corporation	Jan, Feb 1977, Jan 1988- Jun 2003
Georgia State	52	Econometric Modeling Firm	Feb 1984- Jun 2003
Inforum- Univ. of Maryland	57	Econometric Modeling Firm	Jan 1986- Jun 2003
Macroeconomic Advisers, LLC	70	Econometric Modeling Firm	Jan 1985- Jun 2003
Merrill Lynch	73	Econometric Modeling Firm	Jan 1983- Jun 2003
Morgan Stanley & Co.	79	Securities Firm	Jan 1982- Jun 2003
Motorola, Inc.	82	Industrial Corporation	Apr 1993- Jun 2003
National City Bank of Cleveland	84	Commercial Bank	Jan 1977- Jun 2003
National Association of Home Builders	85	Other	Mar 1990- Jun 2003
Northern Trust Company	89	Commercial Bank	Apr 1983, Sep 1985-Jun 2003
Perna Associates	91	Unknown	Jan 1991- Jun 2003
Prudential Financial	97	Other	Jan 1977- Oct 2002
Prudential Securities	98	Securities Firm	Apr 1983- Jun 2003
Siff, Oakley, Marks, Inc.	105	Independent	Jan 1979- Jul 2002
Standard and Poors	108	Other	Jan 1994- Jun 2003
Turning Points (Micrometrics)	111	Independent	Mar 1989- Jun 2003
U.S. Trust Co.	112	Commercial Bank	Jan 1977- Jun 2003
UCLA Business Forecast	114	Econometric Modeling Firm	Mar 1977- Jun 2003
Wayne Hummer & Co.	121	Independent	Apr 1978- Jun 2003
Wells Capital Management	122	Commercial Bank	Jun 1991- Jun 2003

Notes: All 31 forecasters appear consistently in the 1994-2002 period. The current year forecast is the forecast for the end of the year in which they are forecasting. The information for each forecaster is collected from the *Blue Chip Economic Indicators* newsletter.

Figure 2 (continued)
Relationship between herding and accuracy in June and October (1994-2002)

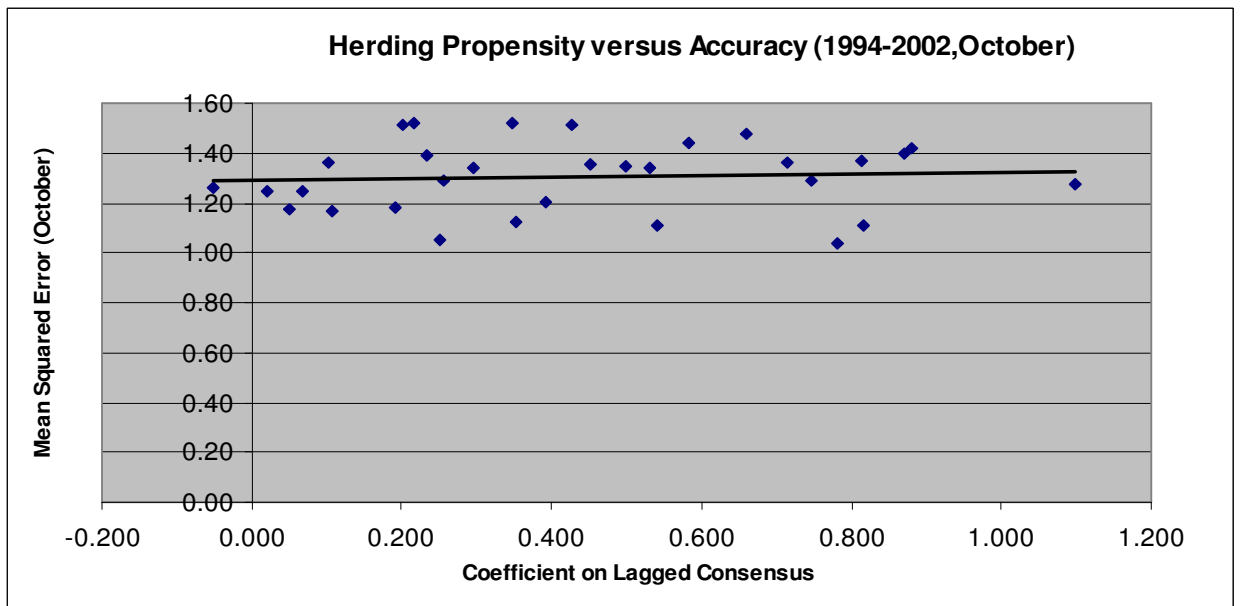
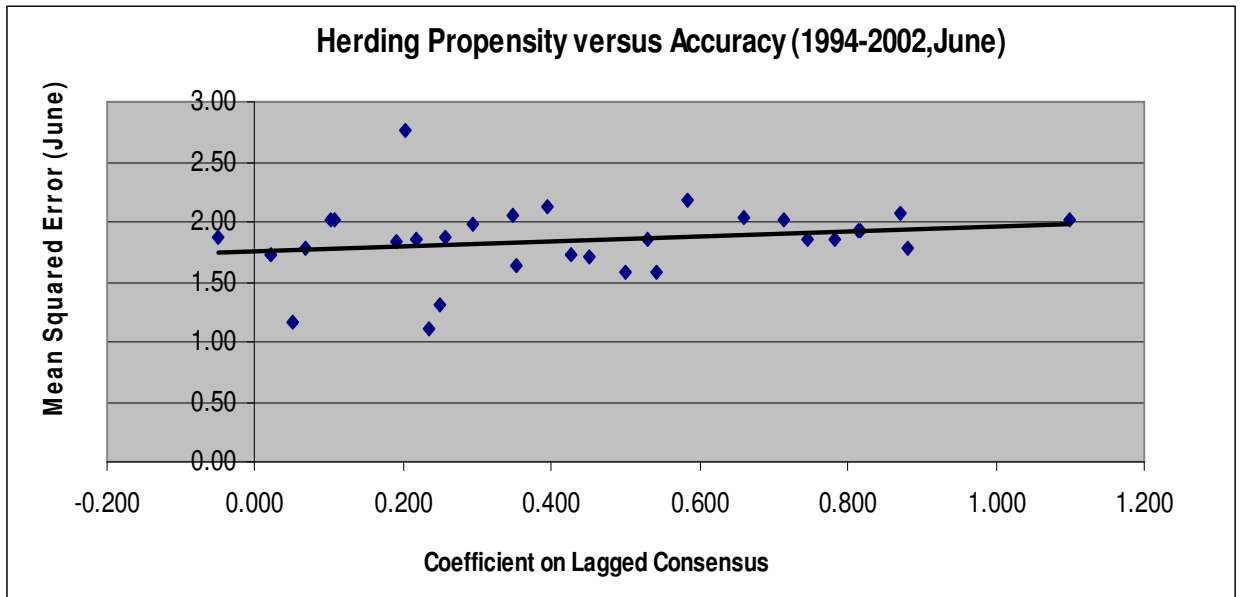


Figure 3 (continued)
Correlation between Herding and Accuracy in June and October (1994-2002)

