



American Finance Association

Judging Fund Managers by the Company They Keep

Author(s): Randolph B. Cohen, Joshua D. Coval, Ľuboš Pástor

Source: *The Journal of Finance*, Vol. 60, No. 3 (Jun., 2005), pp. 1057-1096

Published by: Blackwell Publishing for the American Finance Association

Stable URL: <http://www.jstor.org/stable/3694921>

Accessed: 23/11/2009 12:22

Your use of the JSTOR archive indicates your acceptance of JSTOR's Terms and Conditions of Use, available at <http://www.jstor.org/page/info/about/policies/terms.jsp>. JSTOR's Terms and Conditions of Use provides, in part, that unless you have obtained prior permission, you may not download an entire issue of a journal or multiple copies of articles, and you may use content in the JSTOR archive only for your personal, non-commercial use.

Please contact the publisher regarding any further use of this work. Publisher contact information may be obtained at <http://www.jstor.org/action/showPublisher?publisherCode=black>.

Each copy of any part of a JSTOR transmission must contain the same copyright notice that appears on the screen or printed page of such transmission.

JSTOR is a not-for-profit service that helps scholars, researchers, and students discover, use, and build upon a wide range of content in a trusted digital archive. We use information technology and tools to increase productivity and facilitate new forms of scholarship. For more information about JSTOR, please contact support@jstor.org.



Blackwell Publishing and American Finance Association are collaborating with JSTOR to digitize, preserve and extend access to *The Journal of Finance*.

<http://www.jstor.org>

Judging Fund Managers by the Company They Keep

RANDOLPH B. COHEN, JOSHUA D. COVAL, and ĽUBOŠ PÁSTOR*

ABSTRACT

We develop a performance evaluation approach in which a fund manager's skill is judged by the extent to which the manager's investment decisions resemble the decisions of managers with distinguished performance records. The proposed performance measures use historical returns and holdings of many funds to evaluate the performance of a single fund. Simulations demonstrate that our measures are particularly useful in ranking managers. In an application that relies on such ranking, our measures reveal strong predictability in the returns of U.S. equity funds. Our measures provide information about future fund returns that is not contained in the standard measures.

PEOPLE ARE OFTEN JUDGED BY THE COMPANY they keep. Various human characteristics that are difficult to observe directly are commonly inferred from the characteristics of others who behave in a similar manner. For example, a bored air traveler eager to chat about Madonna's life story is more likely to turn to a fellow passenger reading *The National Enquirer* than to a passenger reading *The Wall Street Journal*, because the readership of the former periodical is generally known to be more interested in celebrity gossip.

In the same spirit, when people compete, their success is often predicted by their techniques, and the quality of a given technique is frequently evaluated based on the track records of the technique's followers. For example, suppose a group of basketball players, some of whom shoot with both hands and some with only one hand, have been taking 10 shots each at the basket. So far, the average score of the two-handers is 8/10, while the one-handers' average is only 4/10.

*Cohen and Coval are from the Harvard Business School. Pástor is from the Graduate School of Business at the University of Chicago, NBER, and CEPR. Research support from the Center for Research in Security Prices and the James S. Kemper Faculty Research Fund at the Graduate School of Business, University of Chicago, is gratefully acknowledged (Pástor). We thank Jonathan Berk, Jeff Busse, George Chacko, Aida Charoenrook, Mike Cooper, Roger Edelen, Gene Fama, Rick Green (editor), Owen Lamont, Rudi Schadt, Erik Sirri, Rob Stambaugh, Pietro Veronesi, Tuomo Vuolteenaho, Russ Wermers, Tong Yao, an anonymous referee, and seminar participants at the 2003 AIM Investment Center Conference on Mutual Funds at the University of Texas at Austin, the 2003 Western Finance Association Meetings, the 2004 Wharton Mutual Fund Conference, the 2004 Chicago Quantitative Alliance Conference, Arizona State University, Berkeley Program in Finance, Dartmouth College, Emory University, Harvard University, Indiana University, Michigan State University, Morningstar, Northwestern University, Purdue University, Thomson Financial, University of Chicago, University of Oregon, University of Pennsylvania, and Vanderbilt University for helpful comments. Huong Trieu provided excellent research assistance.

Two players, one one-hander and one two-hander, have completed only half of their shots so far, and both have scored 4/5. Suppose you are to bet on which of the two players is going to achieve a higher score out of 10. Although the track records of both players are identical, it seems sensible to bet on the two-hander, because the track records of the other players show that two-handed shooters tend to score higher. The one-hander, who is employing what appears to be an inferior technique, is more likely to have been lucky in his first five shots.

Like basketball players, active mutual fund managers rely on a variety of techniques when trying to beat their benchmarks, and many of these techniques are common to groups of managers. For example, managers collect information from different sources and use different valuation methods, but there are clusters of managers who use similar sources and similar methods. Managers using similar techniques are likely to make similar decisions. The premise of this paper is that managers who make similar investment decisions should be expected to deliver similar future performance.

To further support this premise, suppose valuable private information (e.g., a new discovery by a biotech firm) is occasionally received by some managers but not by others. Skill is required to obtain this private information, so managers with more skill receive such information more often. Since they act on the same information, the skilled managers tend to make similar investment decisions (e.g., they all buy the stock of the biotech firm), even if they use different techniques to obtain the private information. As a result, we can tell whether a manager is skilled by comparing his investment decisions with the decisions of other skilled managers.

The premise seems reasonable even if all information is public, but this information is interpreted well by some managers and poorly by others. Managers with more skill are more likely to interpret public information well. In this case, skilled managers again tend to make similar investment decisions because they interpret information similarly. Building on our premise, we propose to judge a fund manager's skill by the extent to which his investment decisions resemble those of other successful managers.

One way to assess the similarity of the managers' investment decisions is to compare the compositions of their portfolios. For example, consider two managers with equally impressive past returns, where one manager currently keeps a big chunk of his portfolio in the stock of Intel, while the other manager holds mostly Microsoft. Suppose also that Intel is currently held especially by managers with good past performance, whereas Microsoft is held mostly by managers with undistinguished records. It seems reasonable to think that the first manager, whose decision to hold Intel is shared by a higher-caliber set of managers, has a superior ability to select stocks, while the second manager, whose decisions coincide with those of subpar managers, has been merely fortunate.

The performance measure proposed in the paper is defined with respect to some simpler reference measure, for which we choose the traditional Jensen's alpha. We show that our measure of a manager's skill is a weighted average of the traditional skill measures across all managers, in which the weights are

essentially the covariances between the manager's current portfolio weights and the current weights of the other managers. Put differently, if two managers have highly similar portfolio weights at the same time, then one manager's skill contributes substantially to our measure of the other manager's skill.

Another way of comparing the managers' investment decisions is to compare their trades. Our trade-based performance measure judges a manager's skill by the extent to which recent changes in his holdings match those of managers with outstanding past performance. This measure is also a weighted average of the traditional skill measures, but now the weights are essentially the covariances between the concurrent changes in the manager's portfolio weights and those of the other managers. According to the trade-based measure, the manager is skilled if he tends to buy stocks that are concurrently purchased by other managers who have performed well, and if he tends to sell stocks that are concurrently purchased by managers who have performed poorly.

Evaluating mutual fund performance is a topic of enormous relevance for the well-being of individual investors and for market efficiency in general. The traditional approach relies solely on historical fund returns to construct measures such as Jensen's alpha (Jensen (1968)) or the Sharpe ratio (Sharpe (1966)). Since the return histories of many funds are short, these traditional measures are often imprecise. To cope with low precision, recent studies propose alternative performance measures that rely not only on fund returns but also on fund holdings.¹ Those measures employ the portfolio holdings of the fund whose performance is being evaluated, but they do not exploit the information contained in the holdings and returns of other funds. Including that additional information and documenting its benefits is the contribution of this paper.

Our performance measures offer substantially higher precision than the traditional return-based measures. Since our measures are weighted averages of the traditional measures, precision is added by pooling information across funds. That is, instead of using just the historical returns of a given manager to estimate his performance, our measures use the return series of all managers whose holdings (or changes in holdings) overlap with those of the given manager. The biggest precision gains from using our measures are obtained for short-history funds. In fact, our measures have reasonably low standard errors even for funds with track records as short as one quarter.

Simulations are conducted to examine the extent to which our skill measures are able to capture true skill. When managers are ranked by true skill and then separately by various performance estimators, our estimators produce higher rank correlations with true skill than standard estimators that do not exploit similarities in holdings or trades across managers. Our measures provide the biggest gains when the number of managers is large and when the fund's return history is short. Our measures exhibit some bias in estimating the

¹ See, for example, Grinblatt and Titman (1993), Daniel et al. (1997), Wermers (2000), and Ferson and Khang (2002). With no reliance on holdings data, Pástor and Stambaugh (2002) and Busse and Irvine (2004) show that additional precision in the traditional performance measures can be achieved by incorporating returns on seemingly unrelated passive assets.

traditional measure, alpha. Due to the weight-averaging of skill across managers, our holdings-based measure is biased toward the average level of skill in the population of managers, similar to a Bayesian shrinkage estimator. The nature of the bias does not impair the rank-ordering of managers, however, and our estimators dominate the usual estimators of alpha even when the objective is to rank managers by their alphas. In sum, the simulations show that our measures are useful especially in applications that involve ranking managers.

One such application is our empirical analysis of return predictability for U.S. equity mutual funds. The literature is ambiguous about whether mutual fund returns are predictable. Grinblatt and Titman (1992), Hendricks, Patel, and Zeckhauser (1993), Goetzmann and Ibbotson (1994), Brown and Goetzmann (1995), Elton, Gruber, and Blake (1996), and Bollen and Busse (2004), among others, all find some persistence in fund performance, especially among the funds lagging their benchmarks. However, much of the persistence has been attributed to the momentum in stock returns (Carhart (1997)), and some also to survivorship bias (Brown et al. (1992)). The mixed nature of the evidence begs for further research.

To analyze fund return predictability, at the beginning of each quarter between April 1982 and July 2002, we sort funds into deciles according to our measures as well as alpha estimated over the past year. We track the deciles' returns over the subsequent quarter, and compare the performance of the decile portfolios over the full sample. We find that fund returns exhibit significant persistence, even after adjusting for momentum in stock returns. The difference between the risk-adjusted returns of the top and bottom deciles sorted by alpha ranges from 3.7% to 5.2% per year across different benchmark models. Our holdings-based measure produces bigger differences, between 5.9% and 7.4% per year, and our trade-based measure yields results similar to alpha.

To see whether our measures contain information about future fund performance that is not contained in alpha, we conduct double sorts, sorting first by alpha and then by our measures. Controlling for alpha, the average difference between the top and bottom quintiles of funds ranked by our holdings-based measure ranges from 2.4% to 4.4% per year, with *t*-statistics ranging from 1.94 to 3.21. The same difference computed using our trade-based measure ranges from 1.2% to 1.4% per year, with *t*-statistics of 2.31 to 2.73. These results indicate that both of our measures contain significant information about future fund returns that is not contained in alpha.

We also examine how much information is contained in alpha but not in our measures, sorting first by our measures and then by alpha. Controlling for our holdings-based measure, the average spread between the top and bottom alpha quintiles is always less than 1% per year and is never significant. Controlling for our trade-based measure, the average spread is 1% to 2% per year, and is significant in half of the cases. Alpha therefore seems to contain some unique information as well, but it seems less informative than our measures.

The highest risk-adjusted returns in our tables are obtained by strategies that combine the information in alpha and in our measures. Consider the (5,5)-(1,1) portfolio, which buys the funds in the top quintiles according to both alpha and

our holdings-based measure, and sells the funds in the bottom quintiles. The Fama–French alpha of this portfolio is 8.52% per year ($t = 3.99$), and the four-factor alpha (after adjusting for momentum) is 5.24% ($t = 2.47$). For the trade-based measure, the two alphas of the (5,5)-(1,1) portfolio are 6.54% ($t = 3.72$) and 4.73% per year ($t = 2.48$). These alphas are higher than those obtained by one-way quintile sorts, suggesting that combining alpha with our measures is useful in predicting future fund performance. Similar results are obtained when portfolios are formed with a one-quarter delay after the portfolio disclosure dates, which allows portfolio holdings to be publicly observable.

We also explore whether investors respond to the information contained in our measures. The investors' response to alphas, in the form of fund flows, is well-documented (e.g., Ippolito (1992), and Sirri and Tufano (1998)). To examine whether investors respond to our measures, we measure net flows into portfolios of funds double-sorted according to alpha and our measures. Fund flows respond strongly to fund alphas, confirming the earlier findings, but they do not respond significantly to our measures. Investors thus seem to be unaware of the information contained in our measures.

The paper proceeds as follows. Section I introduces our performance measures. Section II discusses a simulation exercise that evaluates the usefulness of these measures in capturing true skill. Section III implements the measures empirically to investigate predictability in the returns of U.S. equity mutual funds. Section IV concludes.

I. New Performance Measures

This section introduces two performance measures that judge a fund manager's skill by the extent to which his stock holdings (Section I.A) or stock trades (Section I.B) overlap with those of managers whose other investments have been successful.

A. A Measure Based on Levels of Holdings

Assume that there are M managers, $m = 1, \dots, M$, and N stocks, $n = 1, \dots, N$, each of which is held by at least one manager. Let α_m denote the reference measure of skill for manager m (discussed in the next paragraph), and let $w_{m,n}$ denote the current weight on stock n in manager m 's portfolio. For each stock n , define its quality measure $\bar{\delta}_n$ as

$$\bar{\delta}_n = \sum_{m=1}^M v_{m,n} \alpha_m, \quad (1)$$

where

$$v_{m,n} = \frac{w_{m,n}}{\sum_{m=1}^M w_{m,n}}. \quad (2)$$

The quality of stock n is defined as the average skill of all managers who hold stock n in their portfolios, weighted by how much of the stock they hold. Stocks with high quality are those that are held mostly by highly skilled managers. Managers who hold stocks of high quality are likely to be skilled because their investment decisions are similar to those of other skilled managers (i.e., such managers are in “good company”). Since a larger position in a stock of given quality reveals more about the manager’s ability, the population version of our performance measure is constructed as

$$\delta_m^* = \sum_{n=1}^N w_{m,n} \bar{\delta}_n. \quad (3)$$

This measure of a manager’s performance is the average quality of all stocks in the manager’s portfolio, where each stock contributes according to its portfolio weight.

Our measure δ_m^* is defined in relation to some reference measure of skill, α_m . We choose the population value of Jensen’s alpha, defined as the intercept from the regression of manager m ’s excess returns on the returns of the appropriate benchmark.² Note, however, that alpha is only one of many possible choices for α_m . Other sensible choices include the measures of Grinblatt and Titman (1993) and of Daniel et al. (1997), for example. We opt for the traditional alpha measure solely in the interest of simplicity.

To construct our estimator of managerial skill, we replace α_m in equation (1) with $\hat{\alpha}_m$, the usual OLS estimator of alpha:

$$\hat{\delta}_m^* = \sum_{n=1}^N w_{m,n} \bar{\bar{\delta}}_n, \quad (4)$$

where

$$\bar{\bar{\delta}}_n = \sum_{m=1}^M v_{m,n} \hat{\alpha}_m. \quad (5)$$

Our performance measure has an interesting alternative interpretation. Let W denote the $N \times M$ matrix whose (n, m) element is $w_{m,n}$; let V denote the $N \times M$ matrix whose (n, m) element is $v_{m,n}$; let $\hat{\alpha}$ denote the $M \times 1$ vector of $\{\hat{\alpha}_m\}_{m=1}^M$; let $\bar{\bar{\delta}}$ denote the $N \times 1$ vector of $\{\bar{\bar{\delta}}_n\}_{n=1}^N$; and let $\hat{\delta}^*$ denote the $M \times 1$ vector of $\{\hat{\delta}_m^*\}_{m=1}^M$. Since $\hat{\delta}^* = W' \bar{\bar{\delta}}$ and $\bar{\bar{\delta}} = V \hat{\alpha}$, the vector of our performance measures from equation (4) can be written as

$$\hat{\delta}^* = Z \hat{\alpha}, \quad (6)$$

where $Z = W'V$. The performance of manager m is thus

² The benchmark is assumed to capture any style effects for which the manager should not be rewarded. This paper does not provide any new insights into the choice of the appropriate benchmark.

$$\hat{\delta}_m^* = \sum_{j=1}^M z_{m,j} \hat{\alpha}_j, \quad (7)$$

where $z_{m,j}$ is the (m,j) element of Z , given by

$$z_{m,j} = \sum_{n=1}^N w_{m,n} v_{j,n} = \frac{1}{M} \sum_{n=1}^N w_{m,n} w_{j,n} \frac{1}{h_n}, \quad (8)$$

where

$$h_n = \frac{\sum_{m=1}^M w_{m,n}}{\sum_{n=1}^N \sum_{m=1}^M w_{m,n}} = \frac{\sum_{m=1}^M w_{m,n}}{M} \quad (9)$$

denotes the ratio of the dollar value of stock n held by all M managers to the total dollar value of all stocks held by these managers.

Our measure of manager m 's skill, $\hat{\delta}_m^*$ in equation (7), is a weighted average of the usual skill measures across all managers.³ The weight assigned to the performance of manager j , $z_{m,j}$, is a loose measure of covariance between the weights of managers m and j . This is sensible—if managers m and j own many of the same stocks, they might be using similar techniques, and we want to pay a lot of attention to the performance of manager j when evaluating manager m . The scaling factor $1/h_n$ in equation (8) downweights stocks that receive large weights on average in the portfolios of all managers. For example, if a certain stock occupied 20% of the market capitalization, many managers would have large weights in that stock, and the stock's contribution to the performance measures would be overstated in the absence of the scaling factor. Similarly, if both managers hold a lot of a stock that few others hold, that is valuable information, and the scaling factor emphasizes it.

Our approach adds value only if there is some commonality in the managers' investment decisions. Suppose that manager m holds only stocks that are held by no other manager. In that case, $\hat{\delta}_m^*$ collapses to the traditional measure $\hat{\alpha}_m$.

Due to the symmetry of Z ($z_{i,j} = z_{j,i}$), which follows from equation (8), it is easy to show that the averages of $\hat{\delta}_m^*$ and $\hat{\alpha}_m$ across managers are equal:

$$\frac{1}{M} \sum_{i=1}^M \hat{\delta}_i^* = \frac{1}{M} \sum_{i=1}^M \sum_{j=1}^M z_{i,j} \hat{\alpha}_j = \frac{1}{M} \sum_{j=1}^M \hat{\alpha}_j \sum_{i=1}^M z_{i,j} = \frac{1}{M} \sum_{j=1}^M \hat{\alpha}_j. \quad (10)$$

³ The weights sum to 1:

$$\sum_{j=1}^M z_{m,j} = \sum_{j=1}^M \sum_{n=1}^N w_{m,n} w_{j,n} \frac{1}{M h_n} = \sum_{n=1}^N \frac{1}{M h_n} w_{m,n} \sum_{j=1}^M w_{j,n} = \sum_{n=1}^N w_{m,n} = 1.$$

As a result, our skill measure provides only as much information as the usual measure about the performance of the mutual fund industry as a whole. Nevertheless, our measure can be useful in evaluating the funds' relative performance, as shown later.

Our skill measure can in principle be iterated. The value of $\hat{\delta}_m^*$ computed in equation (4) can be used in place of $\hat{\alpha}_m$ in equation (5) to obtain an iterated estimator of δ_m^* in equation (4). The iterated estimator performs well in the simulations examined in the NBER version of the paper. The gains from repeated iteration are limited, however. Iterating forever is unsatisfactory because it leads to the circular definition $\hat{\delta}^* = Z\hat{\delta}^*$, which implies that the $\hat{\delta}_m^*$'s for all managers are equal to the same arbitrary constant. Intuitively, $\hat{\delta}_m^*$ in equation (7) is a shrinkage estimator, and repeated shrinkage makes all $\hat{\delta}_m^*$'s equal in the limit.

The precision of our estimator can be assessed by its squared standard error, which is displayed along the main diagonal of the $M \times M$ covariance matrix of $\hat{\delta}^*$. This matrix can be computed from equation (6) as

$$\text{Cov}(\hat{\delta}^*, \hat{\delta}^{*'}) = Z\Omega Z', \quad (11)$$

where $\Omega = \text{Cov}(\hat{\alpha}, \hat{\alpha}')$ is the $M \times M$ covariance matrix of $\hat{\alpha}$. If the return histories of all funds spanned the same time period, Ω could in principle be calculated from one big multivariate regression of fund excess returns on benchmark returns. However, funds have histories of unequal lengths spanning different periods. Appendix A describes how Ω can be calculated in such an environment.

To assess the gains in precision provided by our measure, we write the squared standard error of our estimator for manager m as

$$\text{Var}(\hat{\delta}_m^*) = z_m \Omega z_m' = \sum_{i=1}^M \sum_{j=1}^M z_{m,i} z_{m,j} \omega_{i,j}, \quad (12)$$

where z_m is the m^{th} row of Z and $\omega_{i,j}$ is the (i,j) element of Ω . For simplicity, we assume that all m elements of $\hat{\alpha}$ have the same standard error, so that $\omega_{i,i} = \omega^2$ for all $i = 1, \dots, M$. Then $\omega_{i,j} = \omega^2 \rho_{i,j}$, where $\rho_{i,j}$ is the correlation between $\hat{\alpha}_i$ and $\hat{\alpha}_j$. We also assume that all $z_{m,i} > 0$, which is likely to hold in practice for a vast majority of pairs of funds. Then

$$\text{Var}(\hat{\delta}_m^*) = \sum_{i=1}^M z_{m,i}^2 \omega^2 + \sum_{i=1}^M \sum_{j=1, j \neq i}^M z_{m,i} z_{m,j} \omega^2 \rho_{i,j} \quad (13)$$

$$\leq \omega^2 \left(\sum_{i=1}^M z_{m,i}^2 + \sum_{i=1}^M \sum_{j=1, j \neq i}^M z_{m,i} z_{m,j} \right) = \omega^2 \left(\sum_{i=1}^M z_{m,i} \right)^2 = \omega^2, \quad (14)$$

because we already showed that the rows of Z sum to 1. This means that so long as the $\hat{\alpha}_m$'s are not perfectly correlated, $\hat{\delta}_m^*$ has a lower standard error

than $\hat{\alpha}_m$.⁴ Our gains in precision relative to $\hat{\alpha}$ therefore come from the imperfect correlations between the $\hat{\alpha}_m$'s. These correlations tend to be low when the cross-sectional correlation of the managers' residual returns is low. Thus, increasing the number of benchmarks that define $\hat{\alpha}$ can improve the precision of our estimator not only by increasing the precision of $\hat{\alpha}$, but also by reducing the residual correlations among funds.

The fact that $\hat{\delta}^*$'s tend to have substantially lower standard errors than $\hat{\alpha}$'s is confirmed empirically in the NBER version of this paper. In a large sample of U.S. equity funds, we find $\hat{\delta}^*$ to be four to eight times more precise than $\hat{\alpha}$, on average, and to be more precise for 93% to 98% of all funds across three benchmark models. The biggest precision gains are obtained for short-history funds. For example, the correlations between the length of the fund's return history and the ratio of the squared standard error of $\hat{\delta}^*$ to the squared standard error of $\hat{\alpha}$ range from 0.45 to 0.56. The results for the changes estimator $\hat{\delta}^{**}$, which is described in the next section, are similar. We do not report these results in tables to save space and to keep our empirical analysis focused on fund return predictability.

B. A Measure Based on Changes in Holdings

In the previous subsection, we infer that managers are making similar investment decisions if they have similar holdings, regardless of the timing of their trades. In this subsection, we assume that managers make similar decisions if their trades are similar. Since trading involves transaction costs, the decisions to trade are likely to reflect stronger views than the decisions to passively hold. The performance measure developed here exploits similarities between changes in the managers' holdings, rather than their levels.

The return on the portfolio of manager m at time t can be written as

$$R_{m,t} = \sum_{n=1}^N w_{m,n} r_{n,t}, \quad (15)$$

where $r_{n,t}$ denotes the return on stock n . Define the change in the weights as

$$d_{m,n} = w_{m,n,t} - w_{m,n,t-1} \frac{1 + r_{n,t}}{1 + R_{m,t}}, \quad (16)$$

which is the difference between the current weight and the weight obtained if the manager neither bought nor sold any of this stock over the past period.⁵ Let

⁴ If all elements of $\hat{\alpha}$ do not have the same standard error, this relation holds only on average and there may be some m 's for which it does not hold. The calculation in equation (14) is analogous to computing the variance of a portfolio of stocks (with no short positions), which is typically smaller than the variance of any given stock in the portfolio.

⁵ In our empirical analysis, one period equals one quarter. If the manager did not trade at all over the past quarter, so that $d_{m,n} = 0$ for all n , our measure is undefined for this manager. Also note that the time subscripts on d , w , and some related measures below are suppressed to simplify notation.

$\mathcal{N}_m^+ = \{n : d_{m,n} > 0\}$ denote the set of stocks purchased by manager m between $t - 1$ and t , and let $\mathcal{N}_m^- = \{n : d_{m,n} < 0\}$ denote the set of stocks sold by manager m over the same time period. Analogously, let $\mathcal{M}_n^+ = \{m : d_{m,n} > 0\}$ denote the set of managers who made net purchases of stock n between $t - 1$ and t , and let $\mathcal{M}_n^- = \{m : d_{m,n} < 0\}$ denote the set of managers with net sales. We normalize the changes in the weights as

$$x_{m,n}^+ = \frac{d_{m,n}}{\sum_{n \in \mathcal{N}_m^+} d_{m,n}}, \quad x_{m,n}^- = \frac{d_{m,n}}{\sum_{n \in \mathcal{N}_m^-} d_{m,n}}, \quad (17)$$

$$y_{m,n}^+ = \frac{d_{m,n}}{\sum_{m \in \mathcal{M}_n^+} d_{m,n}}, \quad y_{m,n}^- = \frac{d_{m,n}}{\sum_{m \in \mathcal{M}_n^-} d_{m,n}}, \quad (18)$$

so that $x_{m,n}^+$ ($x_{m,n}^-$) captures the fraction of manager m 's purchases (sales) accounted for by stock n , and $y_{m,n}^+$ ($y_{m,n}^-$) captures the fraction of purchases (sales) of stock n accounted for by manager m .

For each stock n , we define its quality measure $\bar{\delta}_n$ as

$$\bar{\delta}_n = \bar{\delta}_n^+ - \bar{\delta}_n^-, \quad (19)$$

where

$$\bar{\delta}_n^+ = \sum_{m \in \mathcal{M}_n^+} y_{m,n}^+ \hat{\alpha}_m, \quad (20)$$

$$\bar{\delta}_n^- = \sum_{m \in \mathcal{M}_n^-} y_{m,n}^- \hat{\alpha}_m, \quad (21)$$

and $\hat{\alpha}_m$ is the usual performance measure, acting as a reference measure again. (Using α_m in place of $\hat{\alpha}_m$ yields the population version of our trade-based skill measure, δ_m^{**} .) The quality of stock n is the difference between the average skill of all managers who bought stock n recently ($\bar{\delta}_n^+$) and the average skill of all managers who sold stock n recently ($\bar{\delta}_n^-$), where the averages are weighted by how much was bought and sold. Stocks of high quality are those that were recently bought mostly by high-skill managers and sold mostly by low-skill managers. Managers who recently bought high-quality stocks and sold low-quality stocks are likely to be skilled, because their decisions are similar to those of other skilled managers. Hence, our trade-based skill measure is

$$\hat{\delta}_m^{**} = \hat{\delta}_m^+ - \hat{\delta}_m^-, \quad (22)$$

where

$$\hat{\delta}_m^+ = \sum_{n \in \mathcal{N}_m^+} x_{m,n}^+ \bar{\delta}_n \quad (23)$$

$$\hat{\delta}_m^- = \sum_{n \in \mathcal{N}_m^-} x_{m,n}^- \bar{\delta}_n. \quad (24)$$

This is the difference between the average quality of the stocks recently bought by manager m and the average quality of the stocks recently sold by this manager. Note that $\hat{\delta}_m^+$ captures the extent to which the manager has been buying high-quality stocks, and $\hat{\delta}_m^-$ captures the extent to which the manager has been selling low-quality stocks. Our measure combines these two aspects of stock-picking skill.

To get more insight into the trade-based measure, we write it out as

$$\begin{aligned}\hat{\delta}_m^{**} &= \sum_{n \in \mathcal{N}_m^+} x_{m,n}^+ \left(\sum_{j \in \mathcal{M}_n^+} y_{j,n}^+ \hat{\alpha}_j - \sum_{j \in \mathcal{M}_n^-} y_{j,n}^- \hat{\alpha}_j \right) \\ &\quad - \sum_{n \in \mathcal{N}_m^-} x_{m,n}^- \left(\sum_{j \in \mathcal{M}_n^+} y_{j,n}^+ \hat{\alpha}_j - \sum_{j \in \mathcal{M}_n^-} y_{j,n}^- \hat{\alpha}_j \right) \\ &= \sum_{j=1}^M c_{m,j} \hat{\alpha}_j,\end{aligned}\tag{25}$$

where

$$\begin{aligned}c_{m,j} &= \sum_{n=1}^N [x_{m,n}^+ y_{j,n}^+ \mathbf{1}_{\{n \in \mathcal{N}_m^+\}} \mathbf{1}_{\{j \in \mathcal{M}_n^+\}} - x_{m,n}^+ y_{j,n}^- \mathbf{1}_{\{n \in \mathcal{N}_m^+\}} \mathbf{1}_{\{j \in \mathcal{M}_n^-\}} \cdots \\ &\quad - x_{m,n}^- y_{j,n}^+ \mathbf{1}_{\{n \in \mathcal{N}_m^-\}} \mathbf{1}_{\{j \in \mathcal{M}_n^+\}} + x_{m,n}^- y_{j,n}^- \mathbf{1}_{\{n \in \mathcal{N}_m^-\}} \mathbf{1}_{\{j \in \mathcal{M}_n^-\}}],\end{aligned}\tag{26}$$

and $\mathbf{1}_{\{\cdot\}}$ denotes an indicator function equal to 1 or 0, depending on whether the associated condition is true. The trade-based measure is a weighted average of the usual measures across all managers, but the weights sum to 0, $\sum_{j=1}^M c_{m,j} = 0$, as shown in Appendix B. The weight on manager j essentially reflects the covariance of the weight changes of managers m and j . This weight, $c_{m,j}$, is a sum of N terms, one for each stock, which capture the products of the managers' weight changes in that stock. If both managers traded that stock in the same direction (i.e., if both bought or sold it), this product is positive and increasing in the extent of the joint action. If one manager bought and the other sold, the product is negative. Hence the loose covariance interpretation. The higher the covariance of the weight changes of managers m and j , the more attention we pay to the skill of manager j when evaluating the skill of manager m .⁶

To summarize the trade-based measure, let C denote the $M \times M$ matrix whose (i, j) element is $c_{i,j}$. Also, let X^+ and X^- denote $M \times N$ matrices whose (i, j) elements are $x_{i,j}^+$ and $x_{i,j}^-$, respectively. Similarly, let Y^+ and Y^- denote $M \times N$ matrices whose (i, j) elements are $y_{i,j}^+$ and $y_{i,j}^-$, respectively. Zeros are substituted for all (i, j) elements of X^+ and Y^+ such that $d_{i,j} < 0$, as well as for all (i, j)

⁶ The average of $\hat{\delta}_m^{**}$ across managers can be shown to equal zero in the special case when managers trade stocks only with each other. Since managers generally trade also with other entities such as individuals, the average $\hat{\delta}_m^{**}$ is likely to be close to zero, but not exactly zero.

elements of X^- and Y^- such that $d_{i,j} > 0$. With this notation, C can be expressed as

$$C = X^+(Y^+)' - X^+(Y^-)' - X^-(Y^+)' + X^-(Y^-)'. \quad (27)$$

Then the $M \times 1$ vector of our performance measures in equation (22) can be written as

$$\hat{\delta}^{**} = C\hat{\alpha}, \quad (28)$$

and the precision of the changes measure can be calculated as

$$\text{Cov}(\hat{\delta}^{**}, \hat{\delta}^{**'}) = C\Omega C', \quad (29)$$

where $\Omega = \text{Cov}(\hat{\alpha}, \hat{\alpha}')$, as before. Henceforth, we often refer to the measure $\hat{\delta}^*$ as the levels measure and to the measure $\hat{\delta}^{**}$ as the changes measure.

Our changes measure in equation (22) weighs stock qualities by the relative magnitudes of the weight changes across stocks. There is an interesting alternative definition, $\hat{\delta}_m^{**A} = \sum_{n=1}^N d_{m,n} \bar{\delta}_n$, which instead relies on the absolute magnitudes of these changes. This alternative changes measure is slightly inferior to $\hat{\delta}^{**}$ in our simulations, but its predictive power for fund returns in one-way sorts is similar to that of $\hat{\delta}^{**}$. We focus on $\hat{\delta}^{**}$, mostly due to its appealing interpretation as the difference between the average qualities of the stocks bought and sold.

C. Some Issues Related to Our Measures

We note that our performance measures are not optimal in the sense of being solutions to an optimization problem. It seems impossible to design the relevant optimization problem without putting some structure on the nature of commonality in skill across managers. Since the true structure of this commonality is unknown, we do not pursue optimality. Instead, we propose measures that capture the underlying intuition in a way that should be reasonably robust to various forms of this commonality.

Our changes measure exploits commonality in trades across managers. The extent to which funds buy and sell stocks at the same time is analyzed by Lakonishok, Shleifer, and Vishny (1992), Grinblatt, Titman, and Wermers (1995), and Wermers (1999), among others. However, this literature on “herding” does not propose using the degree of similarity in the managers’ trades to evaluate their performance. Hong, Kubik, and Stein (2002) show that fund managers from the same city are more likely to hold, buy, or sell the same stock at the same time than are managers from different cities. Given this evidence, managers from the same city are likely to be assigned more similar skill according to our measures than managers from different cities. This seems appropriate. Hong et al. attribute their finding to word-of-mouth information diffusion, and it seems sensible to assign similar skill to managers using similar information.

Given the design of the proposed measures, it seems worthwhile to relate this paper to the literature on “window dressing” (e.g., Haugen and Lakonishok

(1988), Lakonishok et al. (1991), Musto (1997, 1999), and Carhart et al. (2002)). Managers who reshuffle their portfolios around disclosure dates presumably believe that they are judged by the quality of their disclosed portfolios. For example, holdings of unusually risky assets or assets with poor prior performance might lead investors to infer weak managerial ability. In this paper, managers are essentially also judged by the quality of their portfolios, but this quality is assessed differently—by the extent to which these portfolios resemble the portfolios of other well-performing managers. In our approach, a fund manager may be viewed favorably even if he holds stocks that recently performed poorly. This approach seems valuable to rational investors, as discussed in the introduction. In contrast, the literature does not provide a definitive explanation of why window dressing per se should affect the beliefs of rational investors who can observe the funds' track records.

One possible reason behind window dressing is that managers think that they are evaluated by the company they keep. Suppose that the stock of Cisco did well recently. A manager might buy Cisco stock before the disclosure date because she wants her disclosed portfolio to look similar to the portfolios of the successful managers who held Cisco while its stock went up. This tactic may not work, however, because the other managers may have sold Cisco in the meantime. Moreover, this tactic does not confuse our changes measure, because the purchases of Cisco are not concurrent across managers. Also note that the window dressing literature has focused on various asset-pricing issues (e.g., the January effect), but it has not addressed performance evaluation. In short, while there are some similarities between this paper and the window dressing literature, there are also important differences in motivation, focus, and implications for rational investors.

II. Simulations

The purpose of the simulation analysis is to assess the usefulness of our estimators in capturing true managerial skill, which is conveniently known in our simulated environment. Managers receive private signals about each stock's expected return, and a manager's true skill is measured by the probability that the signals he receives are useful. When managers are ranked by their true skill and separately by various performance estimators, useful information about estimator quality is provided by the resulting rank-order correlations. As shown below, our estimators contain important information about true skill that is not captured by standard estimators that make no use of similarities in holdings or trades across managers. The biggest benefits from using our estimators are obtained for funds with relatively short return histories.

A. Simulation Design

The simulations consider a simple setting in which M managers receive signals about expected excess returns of N stocks. Stock n 's excess return at time t is

$$r_{n,t} = \mu_{n,t} + e_{n,t}, \quad n = 1, \dots, N; t = 1, \dots, T, \quad (30)$$

where $\mu_{n,t}$ is the stock's expected excess return and $e_{n,t}$ is an error term. We simulate $\mu_{n,t}$ from $N(0, \sigma_\mu^2)$ and $e_{n,t}$ from $N(0, \sigma_e^2)$, independently across time and stocks.

In every period t , each manager m receives a signal $s_{m,n,t}$ about each stock n . With probability γ_m , this signal is equal to the stock's true expected excess return, and otherwise it is equal to a noise term drawn from an identical distribution:

$$s_{m,n,t} = \begin{cases} \mu_{n,t} & \text{with probability } \gamma_m \\ u_{n,t} & \text{with probability } 1 - \gamma_m, \end{cases} \quad (31)$$

where $u_{n,t} \sim N(0, \sigma_\mu^2)$. Higher γ_m means higher signal quality, so γ_m captures manager m 's true skill. The γ_m 's are drawn independently from the standard uniform distribution, and they are held constant across the T periods.

The managers know the signal structure in equation (31) as well as their own skill γ_m and the error volatility σ_e . They have no information about $\mu_{n,t}$ other than the signal, and they learn about $\mu_{n,t}$ from the signal using the Bayes rule. As shown in Appendix C, the expected excess return on stock n perceived by manager m at time t is

$$E_{m,n,t} = \gamma_m s_{m,n,t}, \quad (32)$$

while the perceived variance of stock n 's returns is

$$V_{m,n,t} = \sigma_e^2 + \sigma_\mu^2 + \gamma_m (s_{m,n,t}^2 - \sigma_\mu^2) - s_{m,n,t}^2 \gamma_m^2, \quad (33)$$

and the perceived return covariance matrix is diagonal. The managers are assumed to maximize the Sharpe ratios of their portfolios, so that manager m 's weight in stock n is

$$w_{m,n,t} \propto V_{m,n,t}^{-1} E_{m,n,t}. \quad (34)$$

As mutual funds are typically not allowed to short stocks, we also require no short sales. This constraint can be imposed simply by putting a zero weight on any stock with a negative signal, thanks to our assumption of uncorrelated stock returns.

In each period t , we record each manager's realized excess return, $r_t^m = w_{m,t}' r_t$, as well as expected excess return, $\alpha_{m,t} = w_{m,t}' \mu_t$, where $w_{m,t}$ is the $N \times 1$ vector of manager m 's portfolio weights, r_t is the vector of $r_{n,t}$'s, and μ_t is the vector of $\mu_{n,t}$'s.⁷ Then

$$\alpha_m = \frac{1}{T} \sum_{t=1}^T \alpha_{m,t} = \frac{1}{T} \sum_{t=1}^T w_{m,t}' \mu_t \quad (35)$$

⁷ For simplicity, we use no benchmark in the simulations, so that excess returns coincide with abnormal returns. Simulating benchmark returns is a complication that would be unlikely to make much difference, other than to make the assumption of uncorrelated e_n more realistic.

denotes manager m 's average expected excess return, which we refer to as the traditional measure of true performance. We calculate four performance estimators. The traditional estimator of α , $\hat{\alpha}$, is simply the manager's average realized excess return:

$$\hat{\alpha}_m = \frac{1}{T} \sum_{t=1}^T r_t^m = \frac{1}{T} \sum_{t=1}^T w'_{m,t} r_t. \quad (36)$$

Our performance measure based on the levels of holdings, $\hat{\delta}^*$, is calculated as

$$\hat{\delta}_m^* = Z_m \hat{\alpha}, \quad (37)$$

where Z_m is defined as row m of matrix Z in equation (6), and the weight matrix W contains the managers' weights in the last (T^{th}) period.⁸ Our performance measure based on changes in holdings, $\hat{\delta}^{**}$, is calculated as

$$\hat{\delta}_m^{**} = C_m \hat{\alpha}, \quad (38)$$

where C_m is defined as row m of matrix C in equation (6), and the weight changes are calculated between the periods $T - 1$ and T . For comparison purposes, we also construct a simple Bayesian estimator that shrinks $\hat{\alpha}$ toward a common mean:

$$\hat{\alpha}_m^B = \frac{1}{2} \hat{\alpha}_m + \frac{1}{2} \tilde{\alpha}, \quad (39)$$

where $\tilde{\alpha}$ is the average of $\hat{\alpha}_m$'s across managers. The four estimators rely only on holdings and return information. We also compute the population versions of δ^* and δ^{**} :

$$\delta_m^* = Z_m \alpha, \quad (40)$$

$$\delta_m^{**} = C_m \alpha, \quad (41)$$

where α is the $N \times 1$ vector of the true α_m 's. These measures possess an unfair advantage over the estimators, since they reflect information that is unknown outside the simulated world, but they help us assess the maximum potential gains from using our estimators.

We conduct 10,000 simulations for each set of parameter values. The number of managers is $M = 30, 100$, and 300 , the number of stocks is $N = 30$ and 100 , and the number of annualized time periods is $T = 1, 5, 10, 20$, and 30 . Throughout, $\sigma_\mu = 0.1$ and $\sigma_e = 0.5$.⁹ In each simulated sample, we calculate α , $\hat{\alpha}$, $\hat{\alpha}^B$, δ^* , δ^{**} , $\hat{\delta}^*$, and $\hat{\delta}^{**}$ for each manager. The managers are ranked according to each

⁸ If there happens to be a stock that is not held by any manager in a given simulated sample, the corresponding row of W is deleted in calculating $\hat{\delta}^*$, to prevent division by zero in equation (2).

⁹ This value of σ_μ assigns 10% cross-sectional dispersion to annual expected excess stock returns. As for σ_e , the median annual return volatility across all 14,149 firms with at least 60 months of contiguous returns on CRSP between January 1926 and December 2002 is 0.4957.

of these measures, as well as according to their true skill γ for the purpose of computing the rank-order correlations.

B. Simulation Results

Table I judges the ability of various performance measures to imitate the ranking of managers by their true skill γ . The table reports the rank-order correlations of each estimator with γ , averaged across 10,000 simulations. For $T = 1$ year, both of our estimators, $\hat{\delta}^*$ and $\hat{\delta}^{**}$, deliver higher rank correlations with γ than the traditional estimator $\hat{\alpha}$. For example, with 300 managers and

Table I
Simulation Evidence about Various Skill Measures Relative to True Skill γ

Random samples of returns on N stocks and signals of M managers over T years are simulated as described in Section II. In each sample, we calculate six different performance measures for each manager. Three of these measures are population versions of the traditional measure α and of our levels and changes estimators, δ^* and δ^{**} . The remaining three measures are sample measures: $\hat{\alpha}$ denotes the average realized excess return on the manager's portfolio, $\hat{\delta}^*$ denotes our levels estimator, which exploits similarities in holdings across managers, and $\hat{\delta}^{**}$ denotes our changes estimator, which exploits similarities in changes in holdings across managers. All managers are ranked according to their performance estimated using the six measures, and rank correlations with the ranking based on true skill γ are reported. All numbers are averaged across 10,000 simulations.

Rank Correlations with True Skill (γ)												
M	$N = 30$						$N = 100$					
	$\hat{\alpha}$	$\hat{\delta}^*$	$\hat{\delta}^{**}$	α	δ^*	δ^{**}	$\hat{\alpha}$	$\hat{\delta}^*$	$\hat{\delta}^{**}$	α	δ^*	δ^{**}
$T = 1$												
30	0.26	0.34	0.35	0.80	0.80	0.82	0.46	0.64	0.65	0.92	0.92	0.93
100	0.27	0.40	0.42	0.81	0.82	0.85	0.47	0.76	0.77	0.93	0.94	0.94
300	0.27	0.44	0.45	0.82	0.83	0.85	0.47	0.80	0.81	0.93	0.94	0.95
$T = 5$												
30	0.53	0.64	0.63	0.94	0.86	0.87	0.77	0.89	0.90	0.98	0.94	0.95
100	0.54	0.72	0.74	0.95	0.85	0.88	0.78	0.93	0.95	0.98	0.95	0.96
300	0.54	0.76	0.78	0.95	0.84	0.88	0.79	0.94	0.96	0.98	0.94	0.96
$T = 10$												
30	0.66	0.75	0.75	0.96	0.86	0.88	0.86	0.93	0.93	0.99	0.95	0.96
100	0.68	0.81	0.83	0.97	0.85	0.89	0.88	0.94	0.96	0.99	0.95	0.96
300	0.68	0.82	0.86	0.97	0.85	0.88	0.88	0.94	0.96	0.99	0.95	0.96
$T = 20$												
30	0.79	0.82	0.83	0.98	0.87	0.89	0.92	0.94	0.95	0.99	0.95	0.96
100	0.80	0.84	0.87	0.98	0.85	0.89	0.93	0.95	0.96	0.99	0.95	0.96
300	0.80	0.84	0.88	0.99	0.85	0.89	0.93	0.95	0.96	1.00	0.95	0.96
$T = 30$												
30	0.84	0.84	0.85	0.98	0.87	0.89	0.94	0.94	0.95	0.99	0.95	0.96
100	0.86	0.84	0.88	0.99	0.85	0.89	0.95	0.95	0.96	1.00	0.95	0.96
300	0.86	0.84	0.88	0.99	0.85	0.89	0.96	0.95	0.96	1.00	0.95	0.96

30 stocks, $\hat{\delta}^*$ and $\hat{\delta}^{**}$ deliver rank correlations of 0.44 and 0.45, respectively, whereas $\hat{\alpha}$ delivers only 0.27. The population measures δ^* and δ^{**} also outperform α . Our measures are particularly effective when the number of managers is large, since there is more cross-sectional information to pool, but they beat $\hat{\alpha}$ even for $M = 30$. As the number of stocks N increases, the managers' portfolios become better diversified and it becomes easier for all measures to detect skill. The measure $\hat{\delta}^{**}$, which uses holdings information not only from period T but also $T - 1$, outperforms $\hat{\delta}^*$.

The length of the sample period matters in judging the estimators' relative success. For $T \in \{5, 10, 20\}$ years, $\hat{\delta}^*$ and $\hat{\delta}^{**}$ still outperform $\hat{\alpha}$, despite the fact that the population measures δ^* and δ^{**} are outperformed by α ; this is due to the higher precision in estimating δ^* and δ^{**} . For $T = 30$, $\hat{\alpha}$ beats $\hat{\delta}^*$, but it still underperforms $\hat{\delta}^{**}$. For $T > 30$ years, $\hat{\alpha}$ beats both $\hat{\delta}^*$ and $\hat{\delta}^{**}$. This is not surprising: As $T \rightarrow \infty$, a manager's skill can be inferred from his average realized return, since true skill is constant over time in the simulations. Therefore, the biggest benefits from using our estimators as opposed to $\hat{\alpha}$ are obtained for funds with relatively short return histories. But even for funds with return histories of up to 30 years or so, the noise in $\hat{\alpha}$ is so large that true skill is better captured by $\hat{\delta}^*$ and $\hat{\delta}^{**}$.

The main goal of this section, accomplished in Table I, is to explore how various skill measures capture true skill γ . In addition, Table II evaluates the ability of these measures to capture α , which is commonly estimated in the literature. Panel A of Table II reports the average rank correlations of each estimator with α . The conclusions are similar to those from Table I. For $T \leq 30$ years or so, $\hat{\delta}^*$ and $\hat{\delta}^{**}$ achieve higher rank correlations with α than the traditional estimator $\hat{\alpha}$, despite the fact that $\hat{\delta}^*$ and $\hat{\delta}^{**}$ are designed to capture their own population quantities, δ^* and δ^{**} , rather than α . For example, with $T = 10$ and $M = N = 30$, both $\hat{\delta}^*$ and $\hat{\delta}^{**}$ deliver rank correlations of 0.77 with α , whereas $\hat{\alpha}$ delivers only 0.68.¹⁰ For $T > 30$, $\hat{\alpha}$ wins, as in Table I. The measure $\hat{\delta}^{**}$ beats $\hat{\delta}^*$ in most cases. The population measures δ^* and δ^{**} attain high rank correlations with α (sometimes as high as 0.99), which present an upper bound on how close we can get to α with our estimators. To sum up, our estimators $\hat{\delta}^*$ and $\hat{\delta}^{**}$ seem successful at imitating the rankings based not only on γ , but also on α .

Panel B of Table II reports the mean squared errors (MSEs) in estimating α for all performance measures.¹¹ The MSEs are calculated as averages across managers and simulations of the squared differences between α and its estimate. As expected, the MSE of the Bayesian estimator $\hat{\alpha}^B$ is lower than the MSE of $\hat{\alpha}$, unless T is large. More interestingly, for $N = 30$ and $T < 30$, $\hat{\delta}^*$ and $\hat{\delta}^{**}$ both have lower MSE than $\hat{\alpha}$, and $\hat{\delta}^{**}$ beats even $\hat{\alpha}^B$, despite the fact that it is not designed to capture α but δ^{**} . As T and N increase, the lowest MSE is ultimately achieved by $\hat{\alpha}$. The MSEs of δ^* and δ^{**} show the extent to which our population measures depart from the true α . Our estimators have high

¹⁰ The Bayesian estimator $\hat{\alpha}^B$ delivers the same rank correlations as $\hat{\alpha}$, by construction.

¹¹ Calculating the MSEs in Table I would not be sensible, as the units of γ and α are different.

precision, which makes them preferred for low T 's, but they also exhibit some bias with respect to α , so the overall effect on the MSE is unclear.

Figure 1 plots the bias in estimating α for four performance estimators, $\hat{\alpha}$, $\hat{\alpha}^B$, $\hat{\delta}^*$, and $\hat{\delta}^{**}$. In each simulated sample, managers are ranked in ascending order by their α . Let α_m denote the true alpha of the m^{th} ranked manager, and $\tilde{\delta}_m$ denote that manager's estimated performance (i.e., $\tilde{\delta}_m = \hat{\alpha}_m, \hat{\alpha}_m^B, \hat{\delta}_m^*$, or $\hat{\delta}_m^{**}$). For each rank m ($m = 1, \dots, M$), the bias is computed by averaging $\tilde{\delta}_m - \alpha_m$ across 10,000 simulated samples. The figure plots the bias against the rank m , with four panels based on different values of M and T . All panels show that the traditional estimator $\hat{\alpha}$ is unbiased, but both the Bayesian estimator

Table II
Simulation Evidence about Various Skill Measures Relative to Traditional Skill α

Random samples of returns on N stocks and signals of M managers over T years are simulated as described in Section II. In each sample, we calculate six different performance measures for each manager. Two of these measures are population versions of our levels and changes estimators, δ^* and δ^{**} . The remaining four measures are sample measures: $\hat{\alpha}$ denotes the average realized excess return on the manager's portfolio, $\hat{\alpha}^B$ denotes the Bayesian estimator that shrinks $\hat{\alpha}$ halfway toward the average of $\hat{\alpha}$ across managers, $\hat{\delta}^*$ denotes our levels estimator, which exploits similarities in holdings across managers, and $\hat{\delta}^{**}$ denotes our changes estimator, which exploits similarities in changes in holdings across managers. All managers are ranked according to their performance estimated using the six measures, and rank correlations with the ranking based on the traditional skill measure α are reported in Panel A. Panel B reports averages across managers of the squared differences between the estimator and the true α (times 100). All numbers are averaged across 10,000 simulations.

Panel A: Rank Correlations with Traditional Skill (α)												
M	$N = 30$						$N = 100$					
	$\hat{\alpha}$	$\hat{\alpha}^B$	$\hat{\delta}^*$	$\hat{\delta}^{**}$	δ^*	δ^{**}	$\hat{\alpha}$	$\hat{\alpha}^B$	$\hat{\delta}^*$	$\hat{\delta}^{**}$	δ^*	δ^{**}
$T = 1$												
30	0.32	0.32	0.41	0.40	0.96	0.95	0.49	0.49	0.67	0.68	0.98	0.98
100	0.33	0.33	0.48	0.47	0.98	0.97	0.50	0.50	0.80	0.80	0.99	0.98
300	0.33	0.33	0.52	0.51	0.99	0.97	0.50	0.50	0.84	0.84	0.99	0.99
$T = 5$												
30	0.55	0.55	0.66	0.65	0.89	0.91	0.78	0.78	0.90	0.91	0.96	0.96
100	0.56	0.56	0.75	0.76	0.88	0.91	0.79	0.79	0.94	0.95	0.96	0.97
300	0.57	0.57	0.78	0.80	0.87	0.91	0.80	0.80	0.95	0.96	0.95	0.97
$T = 10$												
30	0.68	0.68	0.77	0.77	0.88	0.90	0.87	0.87	0.93	0.94	0.95	0.96
100	0.69	0.69	0.82	0.84	0.87	0.90	0.88	0.88	0.95	0.96	0.95	0.96
300	0.70	0.70	0.83	0.87	0.86	0.90	0.89	0.89	0.95	0.96	0.95	0.96
$T = 30$												
30	0.85	0.85	0.85	0.86	0.88	0.90	0.95	0.95	0.95	0.95	0.95	0.96
100	0.86	0.86	0.85	0.88	0.86	0.89	0.95	0.95	0.95	0.96	0.95	0.96
300	0.87	0.87	0.85	0.89	0.85	0.89	0.96	0.96	0.95	0.96	0.95	0.96

(continued)

Table II—Continued

Panel B: Mean Squared Errors												
M	$N = 30$						$N = 100$					
	$\hat{\alpha}$	$\hat{\alpha}^B$	$\hat{\delta}^*$	$\hat{\delta}^{**}$	δ^*	δ^{**}	$\hat{\alpha}$	$\hat{\alpha}^B$	$\hat{\delta}^*$	$\hat{\delta}^{**}$	δ^*	δ^{**}
$T = 1$												
30	2.65	1.71	1.48	0.75	0.11	0.11	0.78	0.53	0.49	0.23	0.09	0.12
100	2.61	1.65	1.40	0.52	0.12	0.10	0.78	0.52	0.49	0.19	0.10	0.13
300	2.62	1.65	1.40	0.47	0.12	0.10	0.78	0.52	0.49	0.19	0.10	0.13
$T = 5$												
30	0.53	0.36	0.36	0.23	0.09	0.09	0.16	0.13	0.17	0.11	0.08	0.08
100	0.53	0.36	0.35	0.17	0.09	0.09	0.16	0.13	0.17	0.10	0.09	0.08
300	0.52	0.35	0.35	0.15	0.10	0.09	0.16	0.13	0.17	0.10	0.10	0.08
$T = 10$												
30	0.26	0.20	0.22	0.16	0.08	0.10	0.08	0.08	0.12	0.10	0.08	0.08
100	0.26	0.19	0.22	0.13	0.09	0.09	0.08	0.08	0.13	0.09	0.09	0.08
300	0.26	0.19	0.22	0.12	0.09	0.09	0.08	0.08	0.13	0.09	0.09	0.08
$T = 30$												
30	0.09	0.09	0.13	0.12	0.08	0.10	0.03	0.05	0.10	0.09	0.08	0.08
100	0.09	0.08	0.13	0.10	0.09	0.09	0.03	0.05	0.10	0.09	0.09	0.09
300	0.09	0.09	0.13	0.10	0.09	0.09	0.03	0.05	0.11	0.09	0.09	0.09

$\hat{\alpha}^B$ and our levels estimator $\hat{\delta}^*$ are biased. For the managers with above-(below-)average skill, both estimators are biased downward (upward), which is entirely expected as a result of the shrinkage of the individual $\hat{\alpha}_m$'s toward their common mean. This shrinkage is explicit in $\hat{\alpha}^B$ and implicit in $\hat{\delta}^*$, which is a weighted average of the $\hat{\alpha}_m$'s across all managers (equation 7). The bias of $\hat{\delta}^*$ is slightly bigger than that of $\hat{\alpha}^B$. The changes estimator $\hat{\delta}^{**}$ is also biased, but unlike $\hat{\delta}^*$, it is biased even on average. This follows by design—while both $\hat{\delta}^*$ and $\hat{\delta}^{**}$ are weighted averages of the $\hat{\alpha}_m$'s, the weights used in $\hat{\delta}^*$ sum to 1, but those in $\hat{\delta}^{**}$ sum to zero, as shown earlier.

The bias with respect to α suggests the need for caution when using our estimators to estimate α . It seems reasonable to use $\hat{\delta}^*$ for such a task, with similar costs and benefits as when using the Bayesian estimator, but $\hat{\delta}^{**}$ is not designed to capture α , even on average. More important, the bias does not compromise the success of our estimators in ranking managers by their α or γ , as we saw in Tables I and II. Our performance measures should thus be particularly useful in applications that involve such ranking, and one such application is presented in the following section.

The simulation results are robust to numerous modifications. First, the results are robust to reasonable changes in σ_μ and σ_e . When the simulations are rerun with $\sigma_\mu = 0.05$ instead of 0.1, the results become even stronger, as both of our estimators outperform $\hat{\alpha}$, even for $T = 30$. With $\sigma_e = 0.3$ instead of the data-implied value of 0.5, the conclusions are again unaffected except that $\hat{\alpha}$ outperforms $\hat{\delta}^*$ and $\hat{\delta}^{**}$ for $T > 10$ years. Less volatile returns improve the relative performance of $\hat{\alpha}$, and more volatile returns hurt $\hat{\alpha}$.

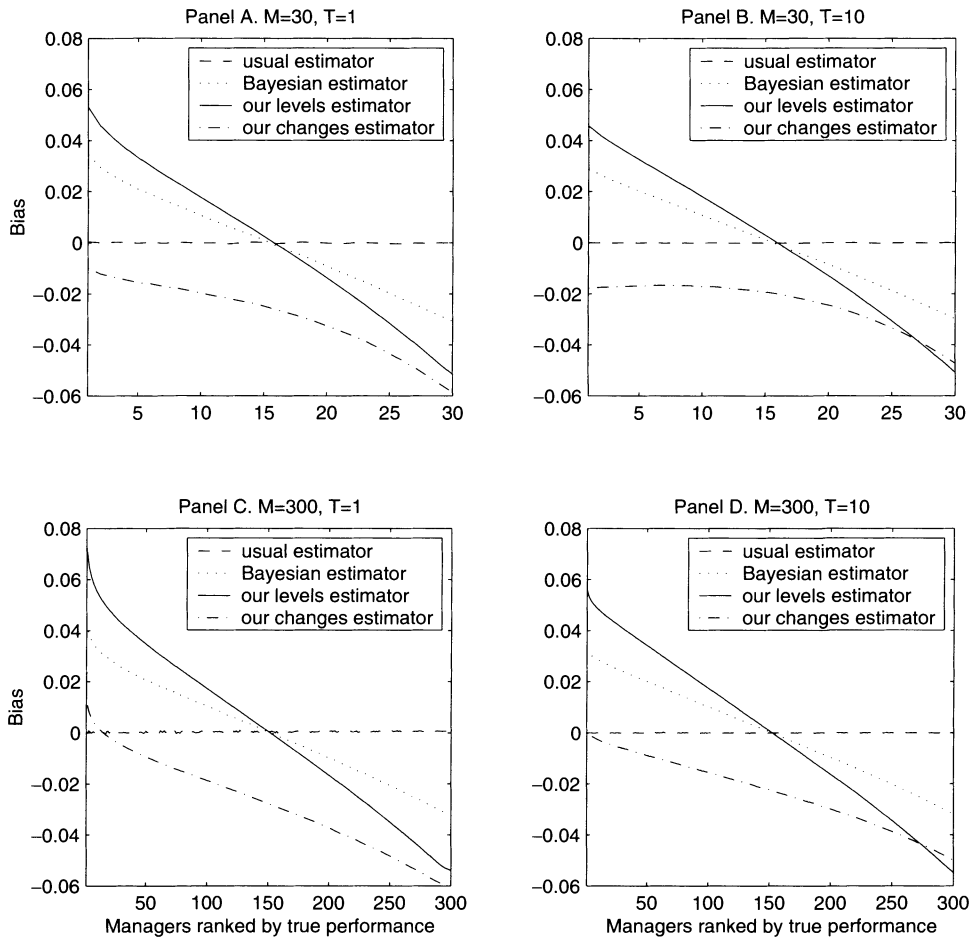


Figure 1. Bias of alternative performance estimators with respect to the traditional performance measure. The figure plots the biases of four performance estimators with respect to the traditional performance measure α , when managers are ranked by the true value of α . In each simulated sample with M (30 or 300) managers, $N = 100$ stocks, and T (1 or 10) years of simulated data, managers are ranked in ascending order by their α . Denote the true performance of the m^{th} ranked manager by α_m . For each estimator $\hat{\delta}$ and each rank m ($m = 1, \dots, M$), the bias is computed by averaging $\hat{\delta}_m - \alpha_m$ across 10,000 simulated samples. Four estimators are considered: the usual estimator $\hat{\alpha}$, the Bayesian estimator $\hat{\alpha}^B$, our levels estimator $\hat{\delta}^*$, and our changes estimator $\hat{\delta}^{**}$.

Allowing stock returns to be correlated leads to the same conclusions as well. To simulate realistic covariance matrices, we first compute return correlations across all pairs of firms with at least 60 months of overlapping contiguous returns on CRSP between January 1926 and December 2002. We sort the resulting 44,418,986 correlations, and record every fifth percentile (i.e., the 5th, 10th, 15th, ..., 95th percentiles).¹² In each simulation, we sample the correlations

¹² The 5th percentile is -0.0839, the 50th percentile is 0.1366, and the 95th percentile is 0.3757.

randomly from the set of 19 percentile correlations, and combine them with σ_e to construct the $N \times N$ return covariance matrix. To rule out short sales with a nondiagonal covariance matrix, we solve the portfolio problem numerically, using techniques of constrained quadratic optimization. To ensure computational feasibility, we limit N to 20, M to 100, and the number of simulations to 5,000. The results, which are not reported to save space, are quite similar to those in Tables I and II. The measure $\hat{\delta}^*$ delivers a higher-rank correlation with true skill than $\hat{\alpha}$ for T up to about 20 years. The measure $\hat{\delta}^{**}$ dominates $\hat{\alpha}$ for $T = 10$ years, but it loses to $\hat{\alpha}$ for $T = 20$. For $T \geq 20$ years, $\hat{\alpha}$ prevails over both of our measures. As argued before, our measures are especially useful for funds with relatively short return histories. A vast majority of real-world mutual funds have return histories shorter than 20 years.

We also let each manager rank stocks by their signals and invest equally in the top 25% of all stocks, instead of using a mean–variance criterion. The results are quite similar to those reported; all rank correlations are slightly lower because managers using this suboptimal strategy do not take full advantage of their skill, but the comparative results are in fact stronger, as $\hat{\delta}^*$ and $\hat{\delta}^{**}$ outperform $\hat{\alpha}$ throughout, even for $T = 30$.

Our measures have one other useful property that is easy to verify by simulation. Imagine a world with a large number of managers, all of whom exhibit no skill ($\gamma = 0$, $\alpha = 0$) and whose holdings overlap sufficiently. The fact that there is no skill in this world would be difficult to detect by $\hat{\alpha}$, because the managers' $\hat{\alpha}$'s would be dispersed around zero due to return noise. However, $\hat{\delta}^*$ and $\hat{\delta}^{**}$ should be zero for all managers, because the manager-specific noise is diversified away by averaging across many $\hat{\alpha}$'s.

III. Empirical Analysis

The simulation results suggest that our skill measures, $\hat{\delta}^*$ and $\hat{\delta}^{**}$, can capture true skill better than the standard measure $\hat{\alpha}$. Therefore, $\hat{\delta}^*$ and $\hat{\delta}^{**}$ should be at least as successful as $\hat{\alpha}$ in predicting fund returns. In this section, the relative predictive ability of all three measures is analyzed empirically in a large sample of equity mutual funds.

A. Data

Mutual fund returns are obtained from CRSP.¹³ These returns are net of fees. Since we are interested in true before-cost managerial ability, we add fees back to obtain gross fund returns. That is, we add the annual expense ratio and

¹³ In the NBER version of this paper, fund returns were computed from the returns on the stocks held by the fund at the most recent quarter-end, assuming that the fund follows a buy-and-hold strategy between the quarter-ends, and that the quarterly rebalancing into the reported holdings is costless. Such hypothetical fund returns (Grinblatt and Titman (1989)) ignore any intraquarter trading by the fund, as well as any ability of the fund to time the market by adjusting its cash holdings. Using the hypothetical fund returns, we found weaker return predictability than reported here.

12(b)1 fees given by CRSP, divide this sum by 12, and add the resulting number to each monthly return in a given year.

Mutual funds are required to file holdings reports with the SEC twice a year, but most funds publicly disclose their portfolio holdings on a quarterly basis. Thomson Financial collects and sells these data. We access the data through WRDS, which currently makes the holdings data available for the period from the end of 1980Q1 through the end of 2002Q2. The data are commonly known as the Spectrum data, since they used to be collected by CDA/Spectrum prior to their purchase by Thomson.¹⁴ The Spectrum mutual fund holdings file contains four columns: date, stock identifier CUSIP, fund identifier, and the number of shares of the given stock held by the given fund on the given date. All dates are quarter-ends (3/31, 6/30, 9/30, or 12/31). Firms with no Spectrum data are recorded as having zero mutual fund ownership. We match each CUSIP to a CRSP PERMNO, the permanent number that CRSP assigns to that security. Holdings associated with CUSIPs for which we found no associated PERMNO are ignored; these account for a very small fraction of holdings.

We merge the Spectrum holdings data with the CRSP mutual fund data via a hand-matching of the funds by name. We often find that separate CRSP listings are merely different share classes of a single fund. In such cases, we select the share class with the greatest number of months of valid data in a given year; if there is no difference by this measure, we use the fund with the lower ICDI identifier. To produce a reliable matching, we also compare total net assets invested in equities as reported in CRSP to the market value of the stock holdings reported in Spectrum, and we exclude funds with substantial discrepancies. The matching procedure yields a sample with 502 funds per quarter, on average. For example, there are 235 funds at the end of 1980 and 1,526 funds in 2002Q2.

One limitation of our data is that holdings are only observed quarterly. If Fund 2 mimics the trades of Fund 1 with a short delay, we are unable to discern such a pattern unless it appears in the funds' quarterly holdings. Both funds may be assigned similar performance by our measures, which seems appropriate only if Fund 1 is unable to profit on its trades before they are mimicked by Fund 2. Useful evidence about the importance of this issue is provided by Chen, Jegadeesh, and Wermers (2000), who find that stocks that experience net purchases by mutual funds have higher subsequent returns than stocks that experience net sales. The authors find that the abnormal performance following the funds' aggregate trades lasts for about a year. Most of the profits from fund trades are therefore not short-lived enough to confuse our performance measures.

Another complication is that the quarter-end dates reported by Thomson Financial do not always correspond to the actual dates when the portfolio snapshots are taken. This leads to a time mismatch for funds with fiscal quarters that differ from calendar quarters. Nonetheless, Wermers (1999, 2000) reports

¹⁴ See Wermers (1999) for detailed information regarding the construction of the database. These data are free of survival bias, as noted by Daniel et al. (1997).

that a vast majority of mutual funds use fiscal years with quarters that coincide with calendar quarters. Like Wermers, we use the approximation that all holdings reported within a given calendar quarter are valid for the end of that calendar quarter. Another problem is that the disclosed portfolios may differ from undisclosed portfolios due to window dressing (Musto (1999)). Since all of these complications add noise to our measures, they should make it more difficult for our measures to predict fund returns. Also note that we do not separate the identity of fund managers from the funds they manage. Baks (2002) constructs a database of fund manager returns and characteristics by tracking the managers' career moves from one fund to another, and uses it to examine fund manager performance. He concludes that the fund typically has a bigger effect on performance than the manager, which helps motivate our focus on funds rather than on managers in the empirical work.

B. Predicting Mutual Fund Returns

At the beginning of each quarter, we compute the traditional measure alpha ($\hat{\alpha}$) for each fund with at least 12 monthly returns by regressing fund returns in excess of the risk-free rate on benchmark returns. If some returns are missing, we run the regression across the months in which returns are available. Using $\hat{\alpha}$ as a reference measure, we then calculate the estimators of our measures, $\hat{\delta}^*$ and $\hat{\delta}^{**}$, from equations (4) and (22).

Nine versions of each measure are computed, using three benchmark models and three lookback periods over which the measures are calculated. The CAPM alpha is calculated with respect to the market benchmark, the Fama–French alpha with respect to the market, size, and value benchmarks of Fama and French (1993), and the four-factor alpha with respect to the market, size, value, and momentum benchmarks, following Carhart (1997).¹⁵ The lookback periods are 12 months, 24 months, and the entire life of the fund. In the interest of space, we only report the results for the Fama–French and four-factor models and only for the 12-month and entire-life lookback periods, but the results for the CAPM alphas and for the 24-month lookback period are similar.

At the beginning of each quarter, funds are sorted into decile portfolios according to each performance measure. The returns on the decile portfolios are calculated over the three months following the portfolio formation, equal-weighting the funds within each decile. The three-month return series are linked across quarters to form a monthly series of returns on each decile portfolio, covering the period from April 1982 through September 2002. (The holdings data begin at the end of 1980Q1, and we require two years of data to obtain the initial estimates.) Table III reports the post-ranking alphas of the decile portfolios as well as of the 10-1 portfolio, constructed by going long the top decile of the best past performers and short the bottom decile. Alpha is defined with respect to the same benchmarks as the reference measure. For example, the table reports the

¹⁵ All benchmark returns, together with the risk-free rate, are obtained from Kenneth French's website.

Table III
Return Predictability for Funds Sorted by Various Measures
of Past Performance

At the end of each quarter, funds are sorted into decile portfolios by various measures of past performance: $\hat{\alpha}$, the OLS estimate of the fund's alpha, $\hat{\delta}^*$, our levels estimator, and $\hat{\delta}^{**}$, our changes estimator. The returns on the decile portfolios are calculated over the three months after portfolio formation, equal-weighting the funds within each decile. The three-month return series are linked across quarters to form a monthly series of returns on each decile portfolio. All performance measures are calculated using the data over the past 12 months (Panel A) and all available past data (Panel B). The table reports the OLS estimates of the deciles' full-period alphas (in percentage per year), as well as their t -statistics (in parentheses). These alphas are defined in the same way as the reference measures $\hat{\alpha}$. The Fama–French alpha is defined with respect to the market, size, and value benchmarks, following Fama and French (1993), and the four-factor alpha with respect to the Fama–French and momentum benchmarks, following Carhart (1997). The sample period is April 1982 through September 2002.

	Decile										
	1	2	3	4	5	6	7	8	9	10	10-1
Panel A: Sorting Funds by Past 12 Months of Performance											
Fama–French Alphas											
$\hat{\alpha}$	−1.62	−0.39	0.00	0.15	0.43	0.75	0.94	1.19	1.62	3.57	5.19
	(−1.62)	(−0.57)	(0.00)	(0.30)	(0.87)	(1.44)	(1.84)	(2.13)	(2.31)	(3.57)	(3.67)
$\hat{\delta}^*$	−1.87	−0.91	−0.75	−0.24	−0.01	−0.01	0.18	2.00	2.72	5.48	7.36
	(−1.30)	(−0.87)	(−1.03)	(−0.42)	(−0.02)	(−0.01)	(0.33)	(2.81)	(2.86)	(4.11)	(3.23)
$\hat{\delta}^{**}$	−1.13	−0.27	−0.12	0.37	0.53	0.07	0.97	0.75	1.51	3.32	4.45
	(−1.23)	(−0.45)	(−0.21)	(0.67)	(1.08)	(0.17)	(1.77)	(1.34)	(2.23)	(3.63)	(4.53)
Four-Factor Alphas											
$\hat{\alpha}$	−1.21	−0.63	0.19	1.13	0.89	0.29	0.65	1.05	1.81	2.48	3.69
	(−1.20)	(−0.80)	(0.31)	(2.13)	(1.81)	(0.54)	(1.29)	(1.68)	(2.63)	(2.60)	(2.64)
$\hat{\delta}^*$	−1.58	−0.89	−0.29	−0.11	0.51	0.72	0.67	1.97	1.33	4.30	5.88
	(−1.14)	(−0.81)	(−0.38)	(−0.17)	(0.91)	(1.32)	(1.25)	(2.56)	(1.37)	(3.46)	(2.73)
$\hat{\delta}^{**}$	−0.60	−0.20	0.30	0.38	0.54	0.76	0.18	0.86	1.15	2.92	3.52
	(−0.62)	(−0.31)	(0.47)	(0.81)	(1.10)	(1.56)	(0.32)	(1.55)	(1.66)	(3.11)	(3.25)
Panel B: Sorting Funds by Entire Past Record											
Fama–French Alphas											
$\hat{\alpha}$	−1.26	−0.05	−0.46	−0.18	0.17	0.24	0.87	1.11	1.77	4.38	5.64
	(−1.42)	(−0.08)	(−0.68)	(−0.35)	(0.39)	(0.48)	(1.82)	(2.04)	(2.79)	(4.41)	(4.49)
$\hat{\delta}^*$	−1.82	−1.36	−0.54	−0.55	0.27	0.21	1.08	1.51	2.47	5.32	7.15
	(−1.67)	(−1.70)	(−0.90)	(−1.12)	(0.52)	(0.40)	(1.95)	(2.34)	(2.66)	(4.12)	(3.84)
$\hat{\delta}^{**}$	−1.08	−0.09	0.31	−0.27	0.53	0.52	1.12	0.26	1.34	3.31	4.39
	(−1.35)	(−0.14)	(0.59)	(−0.49)	(1.05)	(1.05)	(2.05)	(0.46)	(2.05)	(3.61)	(4.67)
Four-Factor Alphas											
$\hat{\alpha}$	−0.74	−0.14	0.36	0.24	0.88	0.66	0.36	1.02	1.47	2.51	3.25
	(−1.03)	(−0.24)	(0.69)	(0.43)	(1.79)	(1.22)	(0.74)	(1.62)	(2.37)	(2.95)	(3.89)
$\hat{\delta}^*$	−1.62	−0.69	−0.27	−0.26	0.40	1.32	1.01	1.54	2.11	3.08	4.70
	(−1.51)	(−0.88)	(−0.43)	(−0.48)	(0.77)	(2.49)	(1.83)	(2.31)	(2.27)	(2.68)	(2.78)
$\hat{\delta}^{**}$	−0.93	0.32	0.15	0.61	0.66	0.47	0.52	0.89	0.39	3.14	4.06
	(−1.14)	(0.52)	(0.28)	(1.07)	(1.23)	(0.98)	(0.95)	(1.40)	(0.56)	(3.43)	(4.90)

post-ranking Fama–French alphas of the portfolios sorted on the past Fama–French $\hat{\alpha}$'s as well as on our measures that use the Fama–French alpha as a reference measure.

All three measures seem capable of predicting fund returns. Consider Panel A of Table III, in which the lookback period is 12 months. The Fama–French alpha of the 10-1 portfolio produced by $\hat{\alpha}$ is 5.2% per year, which is comparable to the 4.5% alpha produced by $\hat{\delta}^{**}$, but is substantially smaller than the 7.4% alpha produced by $\hat{\delta}^*$. Not surprisingly, the persistence in performance weakens when the momentum benchmark is included.¹⁶ Nonetheless, the four-factor alphas remain statistically significant for all three measures. As before, the highest alpha is produced by $\hat{\delta}^*$ (5.9%), while $\hat{\delta}^{**}$ and $\hat{\alpha}$ deliver alphas of 3.5% and 3.7%, respectively. Similar results are shown in Panel B, in which the lookback period is the entire life of the fund. Note that the observed persistence is not restricted to poorly performing funds. In fact, the positive alphas of the top deciles are roughly two to three times larger in absolute value than the negative alphas of the bottom deciles. To summarize, all three measures seem capable of predicting fund returns, and the most predictive power is achieved by $\hat{\delta}^*$.¹⁷

Since our performance measures are correlated with $\hat{\alpha}$, it seems interesting to ask whether our measures contain information not contained in $\hat{\alpha}$ that is useful in forecasting fund returns. To address this question, we perform a series of conditional sorts. Each quarter, we first sort funds into quintiles based on their $\hat{\alpha}$, and then within each $\hat{\alpha}$ quintile, we sort funds into quintiles based on $\hat{\delta}^*$.¹⁸ Panels A of Tables IV and V report the benchmark-adjusted returns for the resulting 25 portfolios, as well as for five 5-1 portfolios that buy funds with high $\hat{\delta}^*$ and short funds with low $\hat{\delta}^*$ within a given $\hat{\alpha}$ quintile. Our cleanest measure of whether $\hat{\delta}^*$ provides information beyond $\hat{\alpha}$ is provided by the portfolio denoted “Avg”, which invests equally in the five 5-1 portfolios.

These results suggest that $\hat{\delta}^*$ contains significant information about future fund returns above and beyond $\hat{\alpha}$. Controlling for $\hat{\alpha}$, the average difference between the top and bottom quintiles of funds ranked by $\hat{\delta}^*$ ranges from 2.4% to 4.4% per year, across both benchmark models and lookback periods. These

¹⁶ Carhart (1997) argues that some apparent persistence in fund performance is due simply to momentum in stock returns. The argument is that, due to momentum, managers who happen to hold mostly stocks that performed well (poorly) over the past year are likely to do well (poorly) also over the following year, even in the absence of any rebalancing on their part. It seems hard to argue that sitting on one's laurels and doing nothing is a managerial skill that should be given credit. Note that Carhart's tables show significant persistence in $\hat{\alpha}$ even after adjusting for momentum.

¹⁷ Berk and Green (2002) argue that net fund returns should be unpredictable if investors compete for returns and if managerial ability exhibits decreasing returns to scale.

¹⁸ We conduct conditional rather than independent sorts because some cells in independent sorts contain few or no funds, due to the correlation between $\hat{\alpha}$ and $\hat{\delta}^*$. In conditional sorts, this correlation leads to the potential concern that the second-step improvement we report may simply be obtained from a finer sort on $\hat{\alpha}$ (i.e., sorting each $\hat{\alpha}$ quintile again by $\hat{\alpha}$ into five subquintiles). However, the spreads in the returns seem too large to be attributed to within-quintile spreads in $\hat{\alpha}$. Moreover, this concern is dismissed in the reverse sorts reported below.

Table IV
Double Sorts Comparing δ^* to Alpha, Using Past 12 Months of Performance

At the end of each quarter, funds are sorted into quintile portfolios according to $\hat{\alpha}$, the OLS estimate of the fund's alpha, and then sorted within the quintiles according to $\hat{\delta}^*$, our levels estimator. Both estimators are constructed using the past 12 months' performance record of each fund. Fund returns are then averaged within each of the 25 portfolios over months 1, 2, and 3 following portfolio formation. The three-month return series are linked across quarters to form a monthly series of returns on each portfolio. The alphas of the resulting 25 return series are reported in Panel A. These alphas are defined in the same way as the reference measures $\hat{\alpha}$. The final two rows report the difference between each of the 5th and 1st portfolios and the associated t -statistic. Panel B reverses the order, sorting first by $\hat{\delta}^*$, and then, within each quintile, by $\hat{\alpha}$. The sample period is April 1982 through September 2002.

Panel A: Sorting Funds by $\hat{\alpha}$ and Then by $\hat{\delta}^*$												
Quintile of $\hat{\delta}^*$	Quintile of $\hat{\alpha}$					Avg.	Quintile of $\hat{\alpha}$					Avg.
	1	2	3	4	5		1	2	3	4	5	
1	-2.89	-0.58	-0.12	0.12	-0.18	-0.73	-1.55	-0.95	0.05	-0.76	0.40	-0.56
2	-1.61	-1.79	-0.53	0.46	0.54	-0.59	-1.28	-0.23	-1.08	-0.01	1.18	-0.28
3	-1.73	0.21	0.14	-0.55	2.18	0.05	-2.05	1.21	0.58	0.81	1.28	0.37
4	-1.05	0.22	0.61	0.77	3.77	0.86	-0.85	1.20	0.95	1.81	2.29	1.08
5	2.34	2.40	2.60	4.65	6.58	3.71	1.22	2.27	1.91	2.78	5.41	2.72
5-1	5.22	2.98	2.72	4.53	6.76	4.44	2.77	3.22	1.86	3.54	5.01	3.28
t -stat	(2.68)	(1.66)	(1.74)	(2.58)	(3.38)	(2.77)	(1.57)	(1.66)	(1.09)	(2.00)	(2.66)	(2.06)

Panel B: Sorting Funds by $\hat{\delta}^*$ and Then by $\hat{\alpha}$												
Quintile of $\hat{\alpha}$	Quintile of $\hat{\delta}^*$					Avg.	Quintile of $\hat{\delta}^*$					Avg.
	1	2	3	4	5		1	2	3	4	5	
1	-2.24	-0.05	0.24	2.00	3.83	0.76	-1.37	-0.90	1.42	1.01	2.36	0.51
2	-2.51	-0.21	0.21	1.52	3.64	0.53	-1.94	0.08	1.05	1.47	2.16	0.56
3	-0.87	-1.26	0.44	0.80	3.66	0.56	-1.19	0.67	0.70	1.63	2.96	0.95
4	-0.94	-0.82	-0.35	0.08	3.22	0.24	-1.12	-0.01	-0.08	1.37	2.83	0.60
5	-0.41	-0.08	-0.36	0.95	6.28	1.28	-0.54	-0.93	0.06	1.09	3.87	0.71
5-1	1.84	-0.04	-0.60	-1.04	2.45	0.52	0.83	-0.04	-1.36	0.08	1.51	0.20
t -stat	(1.49)	(-0.04)	(-0.67)	(-1.07)	(2.61)	(0.82)	(0.68)	(-0.05)	(-1.53)	(0.08)	(1.51)	(0.33)

Table V
Double Sorts Comparing δ^* to Alpha, Using Entire Fund Performance Record

At the end of each quarter, funds are sorted into quintile portfolios according to $\hat{\alpha}$, the OLS estimate of the fund's alpha, and then sorted within the quintiles according to δ^* , our levels estimator. Both estimators are constructed using the entire past performance record of each fund. Fund returns are then averaged within each of the 25 portfolios over months 1, 2, and 3 following portfolio formation. The three-month return series are linked across quarters to form a monthly series of returns on each portfolio. The alphas of the resulting 25 return series are reported in Panel A. These alphas are defined in the same way as the reference measures $\hat{\alpha}$. The final two rows report the difference between each of the 5th and 1st portfolios and the associated t -statistic. Panel B reverses the order, sorting first by δ^* , and then, within each quintile, by $\hat{\alpha}$. The sample period is April 1982 through September 2002.

Panel A: Sorting Funds by $\hat{\alpha}$ and Then by δ^*										
Quintile of δ^*	Quintile of $\hat{\alpha}$					Quintile of $\hat{\alpha}$				
	1	2	3	4	5	Avg.	1	2	3	Avg.
1	-1.66	-1.96	-0.84	0.19	0.42	-0.77	-1.92	-0.65	0.24	0.51
2	-1.41	-0.84	0.06	0.33	2.29	0.09	-1.49	-0.78	-0.65	1.46
3	-1.46	-0.78	-0.17	0.65	2.41	0.13	-0.58	0.62	0.69	2.35
4	-0.14	0.31	-0.07	0.84	3.18	0.83	0.12	1.27	1.79	1.50
5	1.56	1.51	2.18	2.76	7.12	3.03	1.74	1.13	1.73	4.15
5-1	3.22	3.47	3.02	2.57	6.71	3.80	3.65	1.78	1.49	3.63
t -stat	(2.09)	(2.52)	(2.49)	(2.02)	(4.02)	(3.21)	(2.16)	(1.23)	(1.10)	(2.36)
Panel B: Sorting Funds by δ^* and Then by $\hat{\alpha}$										
Quintile of $\hat{\alpha}$	Quintile of δ^*					Quintile of δ^*				
	1	2	3	4	5	Avg.	1	2	3	Avg.
1	-2.18	-1.30	0.31	1.55	2.82	0.24	-1.41	-1.06	0.69	2.07
2	-1.89	-0.02	-0.07	1.42	2.76	0.44	-2.19	-0.62	1.12	1.25
3	-1.39	-0.36	0.27	0.79	3.62	0.59	-0.81	0.97	1.13	2.95
4	-0.94	-0.19	0.48	1.48	4.85	1.14	-0.17	-0.29	-0.09	1.64
5	-1.62	-0.78	0.32	1.35	5.40	0.93	-1.04	-0.42	1.41	4.23
5-1	0.56	0.51	0.01	-0.20	2.58	0.69	0.37	0.65	0.72	2.16
t -stat	(0.62)	(0.90)	(0.01)	(-0.31)	(3.18)	(1.61)	(0.42)	(1.02)	(0.98)	(2.69)

quantities are significant not only economically but also statistically, with t -statistics ranging from 1.94 to 3.21.

We also examine how much information about future fund returns is contained in $\hat{\alpha}$ but not in $\hat{\delta}^*$. To do that, we sort funds into quintiles in reverse, first by $\hat{\delta}^*$ and then by $\hat{\alpha}$, and report the results in Panels B of Tables IV and V. Controlling for $\hat{\delta}^*$, the average 5-1 spread produced by $\hat{\alpha}$ is always less than 1% per year, and is never significant. This suggests that most of the information contained in $\hat{\alpha}$ is already in $\hat{\delta}^*$.

Tables VI and VII report the results for $\hat{\delta}^{**}$. These results are only slightly weaker than those for $\hat{\delta}^*$. Controlling for $\hat{\alpha}$, the average difference between the top and bottom quintiles of funds ranked by $\hat{\delta}^{**}$ ranges from 1.2% to 1.4% per year, with t -statistics ranging from 2.31 to 2.73. This means that $\hat{\delta}^{**}$ adds incremental information about future fund returns over and above $\hat{\alpha}$. Controlling for $\hat{\delta}^{**}$, the average 5-1 spread produced by $\hat{\alpha}$ ranges from 1.3% to 2.1% per year, which is significant in half of the cases. Thus, $\hat{\alpha}$ seems to contain some incremental information beyond $\hat{\delta}^{**}$.

Since both our measures and $\hat{\alpha}$ seem to contain some incremental information relative to each other, it appears that mutual fund portfolio strategies would benefit from combining the information in these measures. Indeed, the highest risk-adjusted returns in our tables are typically offered by the (5,5)-(1,1) portfolio, which buys the funds in the top quintiles according to both $\hat{\alpha}$ and one of our measures, and sells the funds in the bottom quintiles. For example, such a portfolio for $\hat{\delta}^*$ in Panel B of Table IV has a Fama–French alpha of 8.52% per year ($t = 3.99$), and a four-factor alpha of 5.24% ($t = 2.47$). For $\hat{\delta}^{**}$, the two alphas of the (5,5)-(1,1) portfolio in Panel B of Table VI are 6.54% ($t = 3.72$) and 4.73% per year ($t = 2.48$). All these alphas are higher than those obtained by using one-way quintile sorts, which suggests that mutual fund investors would benefit from combining the information in our measures with the information in $\hat{\alpha}$.

To assess the benefits of our measures to investors, we must examine only portfolio strategies that are feasible. The holdings data used to compute $\hat{\delta}^*$ and $\hat{\delta}^{**}$ become publicly available with a lag, since mutual funds can report their holdings to the SEC with a lag of up to two months (e.g., Myers et al. (2001)). To design realistic investment strategies, we assume that fund portfolios that rely on holdings data are formed with a one-quarter delay after the portfolio disclosure date. Specifically, when relying on $\hat{\delta}^*$ or $\hat{\delta}^{**}$, we use returns and holdings through month t to predict returns in months $t + 4$ through $t + 6$ (as opposed to $t + 1$ through $t + 3$). The sample period over which returns are calculated is shifted accordingly by one quarter to July 1982 through December 2002. Because the returns data are available with no lag, sorts that rely on $\hat{\alpha}$ are conducted with no delay (i.e., returns through month $t + 3$ are used to predict returns in months $t + 4$ through $t + 6$). We conduct double sorts as in Tables IV through VII, dropping the reverse sorts to save space, and report the results in Tables VIII and IX.

These results show that our measures help predict fund returns even after allowing for the reporting delays associated with fund holdings. Across all

Table VI
Double Sorts Comparing δ^{**} to Alpha, Using Past 12 Months of Performance

At the end of each quarter, funds are sorted into quintile portfolios according to $\hat{\alpha}$, the OLS estimate of the fund's alpha, and then sorted within the quintiles according to δ^{**} , our changes estimator. Both estimators are constructed using the past 12 months performance record of each fund. Fund returns are then averaged within each of the 25 portfolios over months 1, 2, and 3 following portfolio formation. The three-month return series are linked across quarters to form a monthly series of returns on each portfolio. The alphas of the resulting 25 return series are reported in Panel A. These alphas are defined in the same way as the reference measures $\hat{\alpha}$. The final two rows report the difference between each of the 5th and 1st portfolios and the associated t -statistic. Panel B reverses the order, sorting first by δ^{**} , and then, within each quintile, by $\hat{\alpha}$. The sample period is April 1982 through September 2002.

Panel A: Sorting Funds by $\hat{\alpha}$ and Then by δ^{**}												
Quintile of δ^{**}	Quintile of $\hat{\alpha}$					Avg.	Quintile of $\hat{\alpha}$					Avg.
	1	2	3	4	5		1	2	3	4	5	

Table VII
Double Sorts Comparing δ^{**} to Alpha, Using Entire Fund Performance Record

At the end of each quarter, funds are sorted into quintile portfolios according to $\hat{\alpha}$, the OLS estimate of the fund's alpha, and then sorted within the quintiles according to δ^{**} , our changes estimator. Both estimators are constructed using the entire past performance record of each fund. Fund returns are then averaged within each of the 25 portfolios over months 1, 2, and 3 following portfolio formation. The three-month return series are linked across quarters to form a monthly series of returns on each portfolio. The alphas of the resulting 25 return series are reported in Panel A. These alphas are defined in the same way as the reference measures $\hat{\alpha}$. The final two rows report the difference between each of the 5th and 1st portfolios and the associated t -statistic. Panel B reverses the order, sorting first by δ^{**} , and then, within each quintile, by $\hat{\alpha}$. The sample period is April 1982 through September 2002.

Panel A: Sorting Funds by $\hat{\alpha}$ and Then by $\hat{\delta}^{**}$												
Quintile of $\hat{\delta}^{**}$	Quintile of $\hat{\alpha}$						Quintile of $\hat{\alpha}$					
	1	2	3	4	5	Avg.	1	2	3	4	5	Avg.
1	-1.54	0.23	Fama-French Alphas					Four-Factor Alphas				
2	-1.26	-0.64	-0.55	1.12	2.50	0.35	-2.04	0.10	1.34	0.69	1.85	0.39
3	-0.07	-1.19	0.26	0.52	1.29	0.04	-0.10	-0.38	0.79	0.92	0.49	0.34
4	-0.81	0.44	-0.99	2.08	2.89	0.54	-0.66	1.38	-0.86	0.21	1.11	0.24
5	0.72	0.62	0.98	-0.39	2.84	0.61	-0.05	0.27	0.16	0.08	2.03	0.50
5-1	2.26	0.40	0.32	0.57	5.29	1.50	0.77	1.16	1.59	0.30	4.53	1.67
t -stat	(2.12)	(0.51)	0.87	-0.55	2.79	1.15	2.82	1.06	0.25	-0.39	2.68	1.28
			(1.11)	(-0.67)	(2.49)	(2.31)	(2.47)	(1.19)	(0.30)	(-0.41)	(2.34)	(2.34)
Panel B: Sorting Funds by $\hat{\delta}^{**}$ and Then by $\hat{\alpha}$												
Quintile of $\hat{\alpha}$	Quintile of $\hat{\delta}^{**}$						Quintile of $\hat{\delta}^{**}$					
	1	2	3	4	5	Avg.	1	2	3	4	5	Avg.
1	-1.17	-0.88	Fama-French Alphas					Four-Factor Alphas				
2	-1.91	0.46	0.09	-0.48	1.98	-0.09	-0.86	0.13	0.51	0.71	1.46	0.39
3	-1.03	-0.49	0.76	0.45	1.32	0.22	-1.05	0.26	0.46	1.02	0.24	0.18
4	0.26	-0.12	-0.21	-0.17	1.09	-0.17	-0.49	1.11	-0.12	-0.33	1.40	0.31
5	-0.32	1.46	0.25	0.93	4.06	1.08	-0.48	0.65	-0.35	0.54	2.01	0.47
5-1	0.85	2.34	1.44	2.25	5.02	1.97	0.78	0.80	1.40	1.74	4.22	1.79
t -stat	(0.70)	(2.61)	1.35	2.72	3.04	2.06	1.65	0.67	0.89	1.03	2.76	1.40
			(1.53)	(2.64)	(2.46)	(2.73)	(1.69)	(0.70)	(1.23)	(1.19)	(2.91)	(2.81)

Table VIII
Double Sorts Using Past 12 Months of Performance with One-Quarter Delay in Estimating δ^* and δ^{**}

At the end of each quarter, funds are sorted into quintile portfolios according to $\hat{\alpha}$, the OLS estimate of the fund's alpha, and then sorted within the quintiles according to either $\hat{\delta}^*$ (Panel A) or $\hat{\delta}^{**}$ (Panel B), computed at the end of the previous quarter. All estimators are constructed using the past 12 months performance record of each fund. Fund returns are then averaged within each of the 25 portfolios over months 1, 2, and 3 following portfolio formation. The three-month return series are linked across quarters to form a monthly series of returns on each portfolio, and the alphas of the resulting 25 return series are reported. These alphas are defined in the same way as the reference measures $\hat{\alpha}$. The final two rows report the difference between each of the 5th and 1st portfolios and the associated t -statistic. The sample period is July 1982 through December 2002.

Panel A: Sorting Funds by $\hat{\alpha}$ and Then by $\hat{\delta}^*$										
Quintile of $\hat{\delta}^*$	Quintile of $\hat{\alpha}$					Quintile of $\hat{\alpha}$				
	1	2	3	4	Avg.	1	2	3	4	Avg.
1	-3.31	-2.42	-1.40	-0.56	-0.69	-3.79	-1.81	-0.19	-1.22	0.43
2	-2.17	-0.04	0.33	0.20	1.96	-1.47	0.51	0.35	0.39	1.18
3	-1.62	0.02	1.00	0.24	2.97	-0.55	0.12	-0.02	0.10	2.74
4	-1.07	0.63	0.53	1.45	3.49	-0.33	0.11	0.90	2.25	2.32
5	0.88	2.27	1.38	2.73	5.05	0.46	1.03	1.33	2.06	4.61
5-1	4.19	4.69	2.78	3.29	5.74	4.25	2.84	1.52	3.27	4.18
t -stat	(2.44)	(2.76)	(1.70)	(1.92)	(2.99)	(2.57)	(1.70)	(0.92)	(1.90)	(2.24)
Panel B: Sorting Funds by $\hat{\alpha}$ and Then by $\hat{\delta}^{**}$										
Quintile of $\hat{\delta}^{**}$	Quintile of $\hat{\alpha}$					Quintile of $\hat{\alpha}$				
	1	2	3	4	Avg.	1	2	3	4	Avg.
1	-3.19	-0.19	-0.19	-0.78	2.04	-2.90	-0.71	-0.30	-0.38	1.89
2	-0.87	0.17	0.55	0.68	2.53	-0.41	-0.73	0.94	0.04	1.42
3	-1.02	0.26	0.68	1.38	0.62	0.20	1.85	0.89	0.97	1.31
4	-0.42	1.13	0.58	1.63	3.25	-0.84	0.93	-0.45	0.97	1.60
5	-0.19	-0.21	-0.33	0.79	3.57	-0.13	-0.23	1.16	1.03	4.04
5-1	3.00	-0.01	-0.15	1.57	1.54	2.78	0.48	1.46	1.41	2.15
t -stat	(2.65)	(-0.01)	(-0.17)	(1.60)	(1.44)	(2.39)	(0.57)	(1.50)	(1.64)	(1.95)

Table IX
Double Sorts Using Entire Fund Performance Record with One-Quarter Delay in Estimating δ^* and δ^{}**
At the end of each quarter, funds are sorted into quintile portfolios according to $\hat{\alpha}$, the OLS estimate of the fund's alpha, and then sorted within the quintiles according to either $\hat{\delta}^*$ (Panel A) or $\hat{\delta}^{**}$ (Panel B), computed at the end of the previous quarter. All estimators are constructed using the entire past performance record of each fund. Fund returns are then averaged within each of the 25 portfolios over months 1, 2, and 3 following portfolio formation. The three-month return series are linked across quarters to form a monthly series of returns on each portfolio, and the alphas of the resulting 25 return series are reported. These alphas are defined in the same way as the reference measures $\hat{\alpha}$. The final two rows report the difference between each of the 5th and 1st portfolios and the associated t -statistic. The sample period is July 1982 through December 2002.

Panel A: Sorting Funds by $\hat{\alpha}$ and Then by $\hat{\delta}^*$												
Quintile of $\hat{\delta}^*$	Quintile of $\hat{\alpha}$						Quintile of $\hat{\alpha}$					
	1	2	3	4	5	Avg.	1	2	3	4	5	Avg.
1	-3.61	-1.03	-1.16	0.21	0.73	-0.97	-3.03	-0.67	-0.20	-0.67	-0.21	-0.95
2	-0.85	-0.64	-0.07	0.82	1.01	0.06	-1.08	0.22	-0.29	0.30	0.83	0.00
3	-0.33	0.60	0.34	-0.03	2.42	0.60	0.15	0.62	0.28	0.44	1.83	0.66
4	-0.59	0.19	-0.30	0.62	3.37	0.66	-0.12	1.29	0.86	0.01	1.41	0.69
5	1.18	0.71	1.43	1.47	5.60	2.08	0.49	0.45	1.74	1.56	5.05	1.86
5-1	4.79	1.74	2.59	1.27	4.87	3.05	3.51	1.11	1.94	2.23	5.26	2.81
t -stat	(3.07)	(1.35)	(2.11)	(0.98)	(2.76)	(2.56)	(2.09)	(0.79)	(1.48)	(1.68)	(3.20)	(2.20)
Panel B: Sorting Funds by $\hat{\alpha}$ and Then by $\hat{\delta}^{**}$												
Quintile of $\hat{\delta}^{**}$	Quintile of $\hat{\alpha}$						Quintile of $\hat{\alpha}$					
	1	2	3	4	5	Avg.	1	2	3	4	5	Avg.
1	-1.43	-0.38	-0.82	0.89	1.55	-0.04	-1.82	0.04	0.21	-0.06	1.40	-0.05
2	-0.47	1.12	0.08	0.42	0.97	0.43	-0.55	0.76	0.31	-0.45	1.03	0.22
3	-0.59	-0.40	0.88	1.36	3.40	0.93	0.39	0.34	0.58	0.56	1.79	0.73
4	-0.09	0.15	0.75	0.76	2.31	0.78	0.83	1.51	0.30	1.57	0.59	0.96
5	0.06	-0.23	-0.73	-0.62	3.51	0.40	-1.05	0.61	0.33	-0.71	3.68	0.57
5-1	1.49	0.15	0.08	-1.51	1.95	0.43	0.78	0.57	0.12	-0.65	2.28	0.62
t -stat	(1.31)	(0.19)	(0.09)	(-2.06)	(1.62)	(0.82)	(0.70)	(0.61)	(0.14)	(-0.80)	(2.11)	(1.17)

quintiles of $\hat{\alpha}$, sorting on $\hat{\delta}^*$ creates a positive 5-1 spread in benchmark-adjusted future fund returns. The average spread between the top and bottom $\hat{\delta}^*$ portfolios is between 2.8% and 4.1% per year, with t -statistics ranging from 2.14 to 2.67 across four-model-lookback combinations. Combining $\hat{\delta}^*$ with $\hat{\alpha}$ helps further increase the spread—the alphas of the (5,5)-(1,1) portfolio range from 8.08% to 9.21% per year! The results for $\hat{\delta}^{**}$ are more affected by the reporting delays, since the average 5-1 portfolio alphas range from 0.4% to 1.7%, with t -statistics ranging from 0.82 to 3.03. These results indicate that mutual fund investors can benefit from the information in our measures, especially in $\hat{\delta}^*$.

A natural question raised by these results is whether mutual fund investors are aware of the valuable information contained in our measures, and whether they respond to it by adjusting their fund allocations. Mutual fund flows are well known to respond strongly to $\hat{\alpha}$ (e.g., Ippolito (1992), and Sirri and Tufano (1998)). To see whether they also respond to our measures over and above their response to $\hat{\alpha}$, we conduct similar conditional sorts as before and report the results in Table X. At each year-end, we sort funds into quintiles by $\hat{\alpha}$ and then within each quintile by $\hat{\delta}^*$ (Panel A) or $\hat{\delta}^{**}$ (Panel B) that are lagged by one quarter for reasons explained earlier. For each of the resulting 25 portfolios, we compute the net fund inflow over the following year, and report the average annual net fund inflow across the whole sample.¹⁹ Only the results using a 12-month lookback period are reported, but the results using the funds' full history are very similar.

Looking across the columns confirms that net fund flows increase reliably with $\hat{\alpha}$. The differences between the net inflows into the fifth and first portfolios range from 13% to 34% per year and are almost always statistically significant. Looking down the rows, however, reveals no reliable relation between net flows and our performance measures. These results suggest that while investors pay close attention to $\hat{\alpha}$ in allocating their capital across funds, they pay little if any attention to $\hat{\delta}^*$ and $\hat{\delta}^{**}$.

This ignorance is costly, given the evidence in Tables IV through IX. For example, suppose that investor A sorts funds into quintiles by $\hat{\alpha}$ using a 12-month lookback period and invests in the 5-1 portfolio.²⁰ This investor earns a four-factor alpha of 3.1% per year ($t = 2.8$) over the whole sample period. Investor B performs a conditional sort into quintiles, first by $\hat{\alpha}$ and then by one-quarter-lagged values of $\hat{\delta}^*$, and invests in the (5,5)-(1,1) portfolio. Investor B earns a four-factor alpha of 8.4% per year ($t = 4.0$). The difference between the alphas of investors A and B, 5.3% per year ($t = 3.6$) over a 20.5-year period, is highly economically significant, and is likely to exceed reasonable rebalancing costs. We conclude that mutual fund investors would benefit from using our measures to predict future fund performance.

¹⁹ Net fund inflow is computed as the percentage change in the fund's total net assets minus the fund's return.

²⁰ For simplicity, we assume that it is possible to short funds, but our conclusions do not depend on the ability to short, because a large part of the profits of our long-short strategies comes from the long position.

Table X
Average Net Fund Flows into Funds Sorted by Past Performance

At the end of each year, funds are sorted into quintile portfolios by $\hat{\alpha}$, the OLS estimate of the fund's alpha, and then within the quintiles by values of either $\hat{\delta}^*$ (Panel A) or $\hat{\delta}^{**}$ (Panel B) that are lagged by one quarter. The measure $\hat{\alpha}$ is computed using data from January through December of the current year, whereas the estimators of $\hat{\delta}^*$ and $\hat{\delta}^{**}$ use data from October of the previous year through September of the current year. Average percentage net fund inflows are then calculated for each portfolio over the next year, and the time-series averages of the average annual net inflows are reported. The final two rows and columns report the differences between the net inflows into the 5th and 1st portfolios and their associated t -statistics (in parentheses). The sample period over which the inflows are calculated is January 1983 through December 2002.

Panel A: Sorting Funds by $\hat{\alpha}$ and Then by $\hat{\delta}^*$														
Quintile of $\hat{\delta}^*$	Quintile of $\hat{\alpha}$					Quintile of $\hat{\alpha}$					t-stat	t-stat		
	1	2	3	4	5	1	2	3	4	5				
	Four-Factor Alpha as reference measure													
1	8.34	11.88	23.29	21.96	37.41	29.08	(3.03)	5.03	13.07	22.19	22.57	31.94	26.90	(2.99)
2	7.64	8.88	15.34	19.91	33.26	25.62	(3.10)	13.62	13.64	9.68	15.53	26.69	13.07	(1.63)
3	6.79	16.02	11.76	20.01	38.70	31.91	(3.78)	9.56	10.66	11.84	26.56	33.65	24.09	(3.02)
4	11.53	16.25	18.20	18.67	29.99	18.45	(2.14)	16.35	13.85	15.70	23.57	33.03	16.68	(1.86)
5	19.05	12.15	19.34	18.40	40.90	21.85	(2.21)	19.03	10.84	20.85	26.49	39.60	20.57	(2.04)
5-1	10.71	0.27	-3.95	-3.56	3.49			13.99	-2.23	-1.34	3.92	7.66		
t-stat	(1.41)	(0.03)	(-0.41)	(-0.44)	(0.30)			(1.84)	(-0.30)	(-0.13)	(0.44)	(0.68)		
Panel B: Sorting Funds by $\hat{\alpha}$ and Then by $\hat{\delta}^{**}$														
Quintile of $\hat{\delta}^{**}$	Quintile of $\hat{\alpha}$					Quintile of $\hat{\alpha}$					t-stat	t-stat		
	1	2	3	4	5	1	2	3	4	5				
	Fama-French Alpha as reference measure													
1	5.25	11.29	15.06	16.21	28.14	22.89	(2.71)	6.75	9.66	13.04	13.15	28.26	21.51	(2.48)
2	6.78	11.11	19.50	23.71	32.49	25.71	(3.13)	6.36	17.07	23.25	28.37	33.71	27.35	(3.35)
3	10.50	8.93	15.24	16.85	35.44	24.94	(2.54)	6.86	4.59	9.89	15.31	32.36	25.50	(3.05)
4	8.95	12.49	15.88	15.61	28.72	19.78	(2.34)	10.27	12.98	12.59	11.24	31.29	21.03	(2.20)
5	9.56	15.31	15.58	22.72	43.60	34.04	(3.36)	10.98	14.61	22.90	26.96	42.68	31.70	(3.05)
5-1	4.31	4.02	0.52	6.51	15.46			4.23	4.95	9.86	13.81	14.42		
t-stat	(0.64)	(0.43)	(0.06)	(0.80)	(1.36)			(0.53)	(0.57)	(1.12)	(1.57)	(1.31)		

IV. Conclusion

This paper proposes new performance measures that exploit the information contained in the similarity of a manager's holdings (or changes in holdings) to those of managers who have performed well, and in their distinctiveness from those of managers who have performed poorly. These performance measures use historical returns and holdings of many funds to evaluate the performance of a single fund.²¹ As a result, these measures are typically more precise than the traditional return-based measures. In simulations, our measures are found to be well-suited for empirical applications that involve ranking managers. This suitability is confirmed in our empirical analysis. Our measures are shown to contain a significant amount of information about future fund returns, which is not contained in the standard alpha measure. Our evidence suggests that mutual fund investors could significantly benefit from investing in funds selected by combining the information contained in alpha and in our measures, at least before costs and fees. Despite this potential benefit, mutual fund flows do not respond to our measures nearly as strongly as they respond to alpha.

The basic idea in the paper is that managers who make similar investment decisions have similar skill. The proposed skill measures capture this idea in a simple manner, but future research can extend these measures in various dimensions. For example, the measure of stock quality can be modified to give more weight to the $\hat{\alpha}$'s with lower standard errors. For another example, useful information about similarities in investment decisions may also be contained in the funds' residual return correlations.²² Finally, our measures of a manager's skill rely on the manager's most recent holdings or trades, without considering her historical holdings. The idea is that a manager's current decisions should be more informative than his past decisions about his future performance. It is not clear, without making further assumptions, how exactly historical holdings should be used, but it is clear that they could contain useful information about managerial skill. It would be interesting to design performance measures that exploit similarities in historical holdings or trades across managers, and perhaps also the correlations between historical holdings and subsequent holding returns (Grinblatt and Titman (1993)). Since such measures use yet more information, they might be able to predict fund returns even better than the simple measures proposed in this paper.

This paper uses the proposed performance measures to study the predictability of fund returns, but a variety of other applications can be pursued in future

²¹ Some information pooling across funds takes place also in the Bayesian frameworks developed by Jones and Shanken (2004) and by Stambaugh (2004), in which a fund's alpha is related to the alphas of other funds through a link in the prior. The techniques as well as the objectives of these papers are quite different from ours. Neither study considers similarities in fund managers' decisions and their relation to performance evaluation.

²² Since such correlations are computed across time, any related similarities in investment decisions must be assumed to be stable over time under this approach. No such assumption is needed in our approach, because similarity in holdings or trades can be evaluated across stocks at any given point in time.

work. For example, our evidence about fund flows suggests that the information contained in our measures is not used by many retail investors. However, this information may perhaps be used within fund families to decide whether a manager should be promoted. Evans (2003) finds that manager promotions and demotions are significantly related to alpha, and it seems worth investigating whether they are also related to our measures.

Appendix A: Computing $\Omega = \text{Cov}(\hat{\alpha}, \hat{\alpha}')$ for Funds Whose Return Histories Are Not Aligned in Time

Let $S = \{1, \dots, T\}$ denote the set of dates in the whole sample period in which fund returns may be available. Suppose that returns on Fund 1 are available in the subset S_1 of the whole sample period, $S_1 \subset S$. These returns are stacked in the vector R_1 , whose dimension is $N_1 \times 1$, where N_1 is the number of observations for Fund 1. Analogously, suppose that returns on Fund 2 are available in $S_2 \subset S$, and that they are stacked in the $N_2 \times 1$ vector R_2 . Let R_B denote the vector of returns on K benchmark portfolios, available for the whole period S . All returns are in excess of the risk-free rate. Define

$$R_{1,t} = \alpha_1 + R_{B,t}\beta_1 + \epsilon_{1,t}, \quad t \in S_1 \quad (\text{A1})$$

$$R_{2,t} = \alpha_2 + R_{B,t}\beta_2 + \epsilon_{2,t}, \quad t \in S_2. \quad (\text{A2})$$

The estimated intercepts $\hat{\alpha}_1$ and $\hat{\alpha}_2$ are obtained by running separate OLS regressions, one using the data from S_1 and the other using the data from S_2 .²³ The standard errors of $\hat{\alpha}_1$ and $\hat{\alpha}_2$ are computed accordingly. To calculate Ω , we also need $\text{Cov}(\hat{\alpha}_1, \hat{\alpha}_2)$ for each pair of funds. Rewrite equations (A1) and (A2) as

$$R_1 = X_1\theta_1 + \epsilon_1 \quad (\text{A3})$$

$$R_2 = X_2\theta_2 + \epsilon_2, \quad (\text{A4})$$

where $\theta_j = (\alpha_j \beta_j')'$, X_j is the subset of $X = [\iota_T R_B]$ corresponding to S_j (i.e., we take only the rows that correspond to the dates in S_j), $j = 1, 2$, and ι_T is a T -vector of 1's. Then

$$\hat{\theta}_1 = (X_1' X_1)^{-1} X_1' R_1 = (X_1' X_1)^{-1} X_1' (X_1 \theta_1 + \epsilon_1) = \theta_1 + (X_1' X_1)^{-1} X_1' \epsilon_1$$

$$\hat{\theta}_2 = (X_2' X_2)^{-1} X_2' R_2 = (X_2' X_2)^{-1} X_2' (X_2 \theta_2 + \epsilon_2) = \theta_2 + (X_2' X_2)^{-1} X_2' \epsilon_2.$$

²³ Using the techniques of Pástor and Stambaugh (2002), the intercepts of short-history funds can be estimated more precisely by incorporating the returns on longer-history funds. The most efficient use of all information involves regressing fund returns on the returns of all longer-history funds and the benchmarks. To ensure feasibility, we would have to choose a subset of longer-history funds, and the estimates would depend on that subset, which seems undesirable. (Pástor and Stambaugh use a small set of carefully selected longer-history passive assets instead of longer-history funds.) We opt for the simplicity of estimating the regression intercepts fund by fund.

Hence,

$$\text{Cov}(\hat{\theta}_1, \hat{\theta}_2) = \text{E}[(\hat{\theta}_1 - \theta_1)(\hat{\theta}_2 - \theta_2)'] = (X_1' X_1)^{-1} X_1' \text{E}(\epsilon_1 \epsilon_2') X_2 (X_2' X_2)^{-1}. \quad (\text{A5})$$

Note that ϵ_1 is $N_1 \times 1$ and ϵ_2 is $N_2 \times 1$, so that $\text{E}(\epsilon_1 \epsilon_2')$ is $N_1 \times N_2$. Let σ_{12} denote the contemporaneous covariance between ϵ_1 and ϵ_2 . Then $\text{E}(\epsilon_1 \epsilon_2')$ is a matrix whose (i, j) element is σ_{12} if $S_1(i) = S_2(j)$ and is zero otherwise, since the epsilons are assumed to be uncorrelated over time. Let $\mathcal{O}$ denote the overlap of the funds' sample periods, $\mathcal{O} = S_1 \cap S_2$, let $X_{\mathcal{O}}$ denote the row subset of X corresponding to $\mathcal{O}$, let $N_{\mathcal{O}}$ denote the number of elements in $\mathcal{O}$, and let $I_{N_{\mathcal{O}}}$ denote the identity matrix of dimension $N_{\mathcal{O}}$. Equation (A5) can be rewritten as

$$\text{Cov}(\hat{\theta}_1, \hat{\theta}_2) = (X_1' X_1)^{-1} X_1' (\sigma_{12} I_{N_{\mathcal{O}}}) X_{\mathcal{O}} (X_2' X_2)^{-1} \quad (\text{A6})$$

$$= \sigma_{12} (X_1' X_1)^{-1} (X_1' X_{\mathcal{O}}) (X_2' X_2)^{-1}. \quad (\text{A7})$$

Our estimate of $\text{Cov}(\hat{\theta}_1, \hat{\theta}_2)$ is given by the $(1, 1)$ element of (A7). (For example, if the history of Fund 2 is subsumed by the history of Fund 1, $S_2 \subset S_1$, so that $\mathcal{O} = S_2$, then $\text{Cov}(\hat{\theta}_1, \hat{\theta}_2) = \sigma_{12} (X_1' X_1)^{-1}$.) To estimate σ_{12} , one can run the regressions (A1) and (A2) on the overlapping data $\mathcal{O}$ and take the sample covariance of the resulting residuals.²⁴

Appendix B: Proof of the Statement Immediately Following Equation (26)

$$\begin{aligned} \sum_{j=1}^M c_{m,j} &= \sum_{j=1}^M \sum_{n=1}^N [x_{m,n}^+ y_{j,n}^+ \mathbf{1}_{\{n \in \mathcal{N}_m^+\}} \mathbf{1}_{\{j \in \mathcal{M}_n^+\}} - x_{m,n}^+ y_{j,n}^- \mathbf{1}_{\{n \in \mathcal{N}_m^+\}} \mathbf{1}_{\{j \in \mathcal{M}_n^-\}} \cdots \\ &\quad - x_{m,n}^- y_{j,n}^+ \mathbf{1}_{\{n \in \mathcal{N}_m^-\}} \mathbf{1}_{\{j \in \mathcal{M}_n^+\}} + x_{m,n}^- y_{j,n}^- \mathbf{1}_{\{n \in \mathcal{N}_m^-\}} \mathbf{1}_{\{j \in \mathcal{M}_n^-\}}] \\ &= \sum_{n=1}^N x_{m,n}^+ \mathbf{1}_{\{n \in \mathcal{N}_m^+\}} \sum_{j=1}^M y_{j,n}^+ \mathbf{1}_{\{j \in \mathcal{M}_n^+\}} - \sum_{n=1}^N x_{m,n}^+ \mathbf{1}_{\{n \in \mathcal{N}_m^+\}} \sum_{j=1}^M y_{j,n}^- \mathbf{1}_{\{j \in \mathcal{M}_n^-\}} \cdots \\ &\quad - \sum_{n=1}^N x_{m,n}^- \mathbf{1}_{\{n \in \mathcal{N}_m^-\}} \sum_{j=1}^M y_{j,n}^+ \mathbf{1}_{\{j \in \mathcal{M}_n^+\}} + \sum_{n=1}^N x_{m,n}^- \mathbf{1}_{\{n \in \mathcal{N}_m^-\}} \sum_{j=1}^M y_{j,n}^- \mathbf{1}_{\{j \in \mathcal{M}_n^-\}} \\ &= \sum_{n=1}^N x_{m,n}^+ \mathbf{1}_{\{n \in \mathcal{N}_m^+\}} - \sum_{n=1}^N x_{m,n}^+ \mathbf{1}_{\{n \in \mathcal{N}_m^+\}} - \sum_{n=1}^N x_{m,n}^- \mathbf{1}_{\{n \in \mathcal{N}_m^-\}} + \sum_{n=1}^N x_{m,n}^- \mathbf{1}_{\{n \in \mathcal{N}_m^-\}} \\ &= 0. \end{aligned} \quad (\text{B1})$$

²⁴ There is a more sophisticated approach to estimating σ_{12} when $S_2 \subset S_1$. The results in Stambaugh (1997) can be used to estimate σ_{12} using also the data in S_1 that is not in S_2 . However, since $S_2 \subset S_1$ is unlikely to happen for all pairs of funds, and since we want the same procedure for all pairs, we simply take the estimate of σ_{12} from the overlapping data.

Appendix C: Inference of Our Simulated Investors

Here we calculate the expected return and the covariance matrix of returns as perceived by the simulated mean–variance investors considered in Section II. The subscripts m and t , which denote investor and time, are dropped throughout for convenience.

First, note that with a diffuse prior on μ_n , the Bayes rule implies that

$$\mu_n | s_n, \gamma = \begin{cases} s_n & \text{with probability } \gamma \\ u_n \sim N(0, \sigma_\mu^2) & \text{with probability } 1 - \gamma, \end{cases} \quad (\text{C1})$$

so that

$$E(\mu_n | s_n, \gamma) = \gamma s_n$$

$$\text{Var}(\mu_n | s_n, \gamma) = E(\mu_n^2 | s_n, \gamma) - [E(\mu_n | s_n, \gamma)]^2 = (\gamma - \gamma^2)s_n^2 + (1 - \gamma)\sigma_\mu^2.$$

Let $E(\gamma)$ denote the investor's expected perception of his own skill, and let $\text{Var}(\gamma)$ denote the variance associated with this perception. Using the law of iterated expectations, the expected return on stock n after observing the signals is equal to

$$E(r_n | S) = E[E(\mu_n | s_n, \gamma)] = E(\gamma)s_n. \quad (\text{C2})$$

Using the variance decomposition rule, the perceived variance of stock n 's return is

$$\begin{aligned} \text{Var}(r_n | S) &= \sigma_e^2 + E[\text{Var}(\mu_n | s_n, \gamma)] + \text{Var}[E(\mu_n | s_n, \gamma)] \\ &= \sigma_e^2 + \sigma_\mu^2 + E(\gamma)(s_n^2 - \sigma_\mu^2) - s_n^2[E(\gamma)]^2 + s_n \text{Var}(\gamma). \end{aligned} \quad (\text{C3})$$

The perceived covariance between the returns on stocks i and j is

$$\begin{aligned} \text{Cov}(r_i, r_j | S) &= E[\text{Cov}(\mu_i, \mu_j | S, \gamma)] + \text{Cov}[E(\mu_i | S, \gamma), E(\mu_j | S, \gamma)] \\ &\quad + \text{Cov}(e_i, e_j | S) \\ &= E[E((\mu_i - \gamma s_i)(\mu_j - \gamma s_j) | S, \gamma)] + \text{Cov}[\gamma s_i, \gamma s_j] + \rho_e \sigma_e^2 \\ &= E[\gamma^2 s_i s_j - \gamma^2 s_i s_j - \gamma^2 s_i s_j + \gamma^2 s_i s_j] + s_i s_j \text{Var}(\gamma) + \rho_e \sigma_e^2 \\ &= s_i s_j \text{Var}(\gamma) + \rho_e \sigma_e^2, \end{aligned} \quad (\text{C4})$$

where ρ_e denotes the correlation across stocks, which is assumed to be zero in the basic simulation. For simplicity, we also assume that each manager knows his own γ , so that $E(\gamma) = \gamma$ and $\text{Var}(\gamma) = 0$.

REFERENCES

- Baks, Klaas P., 2002, On the performance of mutual fund managers, Working paper, Emory University.
- Berk, Jonathan B., and Richard C. Green, 2002, Mutual fund flows and performance in rational markets, Working paper, University of California, Berkeley.
- Bollen, Nicolas P. B., and Jeffrey A. Busse, 2004, Short-term persistence in mutual fund performance, *Review of Financial Studies*, forthcoming.
- Brown, Stephen J., and William N. Goetzmann, 1995, Performance persistence, *Journal of Finance* 50, 679–698.
- Brown, Stephen J., William N. Goetzmann, Roger G. Ibbotson, and Stephen A. Ross, 1992, Survivorship bias in performance studies, *Review of Financial Studies* 5, 553–580.
- Busse, Jeffrey A., and Paul J. Irvine, 2004, Bayesian alphas and mutual fund persistence, Working paper, Emory University.
- Carhart, Mark M., 1997, On persistence in mutual fund performance, *Journal of Finance* 52, 57–82.
- Carhart, Mark M., Ron Kaniel, David K. Musto, and Adam V. Reed, 2002, Learning for the tape: Evidence of gaming behavior in equity mutual funds, *Journal of Finance* 57, 661–693.
- Chen, Hsiu-Lang, Narasimhan Jegadeesh, and Russ Wermers, 2000, The value of active mutual fund management: An examination of the stockholdings and trades of fund managers, *Journal of Financial and Quantitative Analysis* 35, 343–368.
- Daniel, Kent D., Mark Grinblatt, Sheridan Titman, and Russ Wermers, 1997, Measuring mutual fund performance with characteristic-based benchmarks, *Journal of Finance* 52, 1035–1058.
- Elton, Edwin J., Martin J. Gruber, and Christopher R. Blake, 1996, The persistence of risk-adjusted mutual fund performance, *Journal of Business* 69, 133–157.
- Evans, Richard B., 2003, Does alpha really matter? Evidence from mutual fund incubation, termination and manager change, Working paper, Wharton.
- Fama, Eugene F., and Kenneth R. French, 1993, Common risk factors in the returns on stocks and bonds, *Journal of Financial Economics* 33, 3–56.
- Ferson, Wayne E., and Kenneth Khang, 2002, Conditional performance measurement using portfolio weights: Evidence for pension funds, *Journal of Financial Economics* 65, 249–282.
- Goetzmann, William N., and Roger G. Ibbotson, 1994, Do winners repeat? Patterns in mutual fund performance, *Journal of Portfolio Management* 20, 9–18.
- Grinblatt, Mark, and Sheridan Titman, 1989, Mutual fund performance: An analysis of quarterly portfolio holdings, *Journal of Business* 62, 394–415.
- Grinblatt, Mark, and Sheridan Titman, 1992, The persistence of mutual fund performance, *Journal of Finance* 47, 1977–1984.
- Grinblatt, Mark, and Sheridan Titman, 1993, Performance measurement without benchmarks: An examination of mutual fund returns, *Journal of Business* 66, 47–68.
- Grinblatt, Mark, Sheridan Titman, and Russ Wermers, 1995, Momentum investment strategies, portfolio performance, and herding: A study of mutual fund behavior, *American Economic Review* 85, 1088–1105.
- Haugen, Robert A., and Josef Lakonishok, 1988, *The Incredible January Effect: The Stock Market's Unsolved Mystery* (Dow Jones-Irwin, Homewood, IL).
- Hendricks, Darryl, Jayendu Patel, and Richard Zeckhauser, 1993, Hot hands in mutual funds: Short-run persistence of relative performance, 1974–1988, *Journal of Finance* 48, 93–130.
- Hong, Harrison, Jeffrey D. Kubik, and Jeremy C. Stein, 2002, Thy neighbor's portfolio: Word-of-mouth effects in the holdings and trades of money managers, Working paper, Stanford University.
- Ippolito, Richard A., 1992, Consumer reaction to measures of poor quality: Evidence from the mutual fund industry, *Journal of Law and Economics* 35, 45–70.
- Jensen, Michael, 1968, The performance of mutual funds in the period 1945–1964, *Journal of Finance* 23, 389–416.
- Jones, Christopher S., and Jay Shanken, 2004, Mutual fund performance with learning across funds, Working paper, University of Southern California.

- Lakonishok, Josef, Andrei Shleifer, Richard Thaler, and Robert W. Vishny, 1991, Window dressing by pension fund managers, *American Economic Review* 81, 227–231.
- Lakonishok, Josef, Andrei Shleifer, and Robert W. Vishny, 1992, The impact of institutional trading on stock prices, *Journal of Financial Economics* 32, 23–44.
- Musto, David K., 1997, Portfolio disclosures and year-end price shifts, *Journal of Finance* 52, 1563–1588.
- Musto, David K., 1999, Investment decisions depend on portfolio disclosures, *Journal of Finance* 54, 935–952.
- Myers, Mary M., James M. Poterba, John B. Shoven, and Douglas A. Shackelford, 2001, Copycat funds: Information disclosure regulation and the returns to active management in the mutual fund industry, Working paper, University of Chicago.
- Pástor, Ľuboš, and Robert Stambaugh, 2002, Mutual fund performance and seemingly unrelated assets, *Journal of Financial Economics* 63, 315–349.
- Sharpe, William F., 1966, Mutual fund performance, *Journal of Business* 39, 119–138.
- Sirri, Erik R., and Peter Tufano, 1998, Costly search and mutual fund flows, *Journal of Finance* 53, 1589–1622.
- Stambaugh, Robert F., 1997, Analyzing investments whose histories differ in length, *Journal of Financial Economics* 45, 285–331.
- Stambaugh, Robert F., 2004, Inference about survivors, Working paper, Wharton.
- Wermers, Russ, 1999, Mutual fund herding and the impact on stock prices, *Journal of Finance* 54, 581–622.
- Wermers, Russ, 2000, Mutual fund performance: An empirical decomposition into stock-picking talent, style, transaction costs, and expenses, *Journal of Finance* 55, 1655–1695.