

An exploration of the persistence of UK unit trust performance

Jonathan Fletcher^{a,*}, David Forbes^b

^a*Department of Accounting and Finance, University of Strathclyde, Curran Building,
100 Cathedral Street, Glasgow G4 0LN, UK*

^b*Glasgow Caledonian University, Glasgow, UK*

Abstract

We examine the persistence in UK unit trust performance between January 1982 and December 1996. We find significant persistence in the relative rankings of trusts using different performance measures. We also find significant persistence in the performance of portfolios of trusts, formed on the basis of prior year excess returns, when performance is evaluated relative to models based on the capital asset pricing model (CAPM) or arbitrage pricing theory (APT). However this persistence is eliminated when performance is evaluated relative to a model similar to Carhart [Journal of Finance 52 (1997) 57]. Using a conditional performance measure leads to significant reversals in performance with this model.

© 2002 Elsevier Science B.V. All rights reserved.

JEL classification: G12

Keywords: Performance persistence; Benchmark portfolios

1. Introduction

Over the past 10 years there has been a substantial growth in empirical studies examining whether fund performance is predictable over time. This has largely focused on US mutual funds and includes studies by Grinblatt and Titman (1992), Hendricks et al. (1993), Goetzmann and Ibbotson (1994), Malkiel (1995), Brown and Goetzmann (1995), Elton et al. (1996) and Carhart (1997) amongst others. The evidence in these studies suggests that past performance provide information about future performance. More recent studies have focused on other types of financial institutions. Christopherson et al. (1998)

* Corresponding author. Tel.: +44-141-548-3892; fax: +44-141-552-3547.

E-mail address: j.fletcher@strath.ac.uk (J. Fletcher).

and Christopherson et al. (1999) examine persistence in US pension funds and Brown et al. (1999) examine hedge funds.

There are at least two reasons why performance persistence is an important issue of study. The first is to examine whether fund managers have superior investment skill. Gruber (1996) argues that since US mutual funds sell at net asset value, then management skill is not priced. Provided funds with good performance do not subsequently increase their fees¹ then performance should be predictable. Evidence in Carhart (1997) suggests that the persistence in US mutual funds is not due to superior investment skill when a four factor model is used to evaluate the out of sample performance of portfolio strategies formed on the basis of prior performance. A second reason for studying persistence is to consider whether past performance provides information about future performance for investors. Performance league tables are reported in financial magazines and performance consultancy services. Evidence in Sirri and Tufano (1998) and Gruber (1996) shows that funds with good past performance receive the largest inflows of new cash² flows. This has the advantage for funds because the size of these funds will tend to grow and will receive higher fees since fees are usually set as a percentage of net asset value.

Empirical evidence of persistence has been, until fairly recently, limited outside that of the USA. A number of recent studies have examined the predictability in UK fund performance. Brown et al. (1997) and Allen and Tan (1999) find significant persistence in the relative rankings of UK pension funds and unit trusts³, respectively. The persistence of the relative rankings of trusts has been questioned by a recent study by Rhodes (2000) who documents that the persistence in relative rankings has been considerably weaker in the 1990s compared to the early 1980s. Blake and Timmermann (1998) find that the portfolio of the top quartile of trusts formed on the basis of prior performance generates significantly positive abnormal returns relative to a three-factor model. In addition the bottom quartile of trusts earns significantly negative abnormal returns relative to a three-factor model.^{4,5} Lunde et al. (1999) examine the impact of the attrition of unit trusts on performance persistence tests.

This paper uses the methodologies of Brown and Goetzmann (1995) and Carhart (1997) to examine a number of issues that have not received a great deal of attention in the prior literature of UK fund performance. The paper begins by re-examining the persistence in the relative rankings of trusts over consecutive periods of time and whether this is robust to alternative performance measures. This is explored further by examining the persistence in trust performance when trusts are compared to an absolute benchmark, such as a stock market index, and again how robust this is to alternative performance measures. There is little prior research of this latter issue in UK fund performance. The main concern of the paper is to explore possible explanations of the persistence in performance of portfolios of

¹ Elton et al. (1996) and Gruber (1996) find that funds tend not to increase their fees more than other funds after periods of good performance.

² Gruber (1996) and Zheng (1999) examine the fund selection ability of mutual fund investors as to whether movements in cash flows are able to predict future performance.

³ UK unit trusts are equivalent to open-ended US mutual funds.

⁴ This is a similar model to Elton et al. (1993).

⁵ This is particularly strong in UK smaller companies sector.

trusts, formed on the basis of prior year excess returns, that has been found in studies such as [Blake and Timmermann \(1998\)](#). Using a methodology similar to [Carhart \(1997\)](#), we examine whether the persistence can be explained by evaluating performance relative to different factor models or the use of the conditional performance measure developed by [Ferson and Schadt \(1996\)](#). In particular, we address whether the evidence suggests that the persistence is a reflection of superior stock selection ability.

A number of findings emerge from the study many of, which are consistent with the prior research of UK and US funds. However there are some interesting differences with US fund research. The first main finding is that there is significant persistence in the relative rankings of trusts over consecutive 1-year and 2-year intervals, which is fairly robust to the performance measure used. However when persistence is examined by comparing trust performance to an absolute benchmark, the persistence is largely driven by repeat underperformance. The second main finding is that there is significant persistence in the performance of portfolios formed on the basis of prior year excess returns which cannot be explained by evaluating performance on the basis of factor models based on the CAPM or APT. The third main finding is that the persistence in performance is eliminated when performance is estimated relative to the [Carhart \(1997\)](#) model. When the conditional performance measure is used, this leads to significant reversals in performance. The evidence in the paper suggests that the persistence in performance of UK trusts is not a manifestation of superior stock selection strategy and can be explained by factors that are known to capture cross-sectional differences in stock returns.

The paper is organised as follows. Section 2 describes the methodology. Section 3 discusses the sample of unit trusts and other data. Section 4 reports empirical results. The final section contains concluding comments.

2. Methodology

This study employs the methodologies of [Brown and Goetzmann \(1995\)](#) and [Carhart \(1997\)](#) to evaluate the persistence in trust performance. [Brown and Goetzmann \(1995\)](#) use 2*2 contingency tables to evaluate performance persistence. Within each period, the trusts are defined as either Winners or Losers depending on their performance. Using two consecutive periods, a trust can be allocated to one of four categories. These are Winner in both periods (WW), a Winner in the first period and a Loser in the second period (WL), a Loser in the first period and a Winner in the second period (LW) and a Loser in both periods (LL). If there is evidence of positive persistence, then we would expect to observe more trusts in either the WW or LL categories. If there is reversal in performance, we would expect more trusts in the WL or LW categories.

[Brown and Goetzmann \(1995\)](#) propose the log–odds ratio to test for significant persistence. This is defined as

$$\log - \text{odds ratio} = \ln(WW*LL)/(WL*LW) \quad (1)$$

Under the null hypothesis of no persistence, the log–odds ratio will equal zero. The null hypothesis can be tested by the *z* test, which is simply the log–odds ratio divided by the

standard error. The standard error equals $\sqrt{(1/WW)+(1/WL)+(1/LW)+(1/LL)}$. This has a standard normal distribution. A significantly positive log-odds ratio is evidence of persistence in performance and a significantly negative log-odds ratio is evidence of reversal in performance.

The performance of the trusts over different subperiods is measured in different ways. The first is based on the cumulative excess return of the trusts. A second approach is to use market-adjusted returns of the trust. This is the cumulative annual excess return on the trust less the annual excess return on the market index.⁶ The third approach is the unconditional Jensen (1968) measure. This can be estimated from the following regression:

$$r_{it} = \alpha_i + \sum_{k=1}^K \beta_{ik} r_{kt} + \varepsilon_{it} \quad (2)$$

where r_{it} is the excess return on trust i in period t , r_{kt} is the excess return on the k th portfolio in the benchmark for $k=1, \dots, K$, K is the number of portfolios in the benchmark, ε_{it} is a random error term with $E(\varepsilon_{it})=0$ and $E(\varepsilon_{it}r_{kt})=0$ for $k=1, \dots, K$.

The β_{ik} coefficients are the betas of trust i to each of the portfolios in the benchmark. The intercept α_i is the Jensen (1968) performance measure.

The final method is the conditional Jensen measure of Ferson and Schadt (1996). This allows the betas and risk premiums to vary through time. Ferson and Schadt (1996) assume that the fund beta is a linear function of the information variables used by investors to set prices. Ferson and Schadt (1996) show that the conditional Jensen measures can be estimated within a CAPM framework from the following regression:

$$r_{it} = \alpha_i + \beta_i r_{mt} + \sum_{l=1}^L \delta_{il} r_{mt} z_{lt-1} + e_{it} \quad (3)$$

where z_{lt-1} is the de-meaned l th information variable for $l=1, \dots, L$, β_i is the average conditional beta and L is the number of information variables. The additional terms $r_{mt} z_{lt-1}$ captures the covariance between the conditional beta of the trust and risk premiums. The intercept α_i is the conditional Jensen measure which will equal zero if the trust exhibits no abnormal performance. For multi-factor models, the conditional Jensen measure is the intercept in the excess return regression of the trust return on a constant, the benchmark excess returns and the products of the information variables with each benchmark portfolio.

The paper also uses the approach of Carhart (1997)⁷ to examine the persistence in performance of quartile (decile) portfolios of trusts formed on the basis of prior year excess returns. This approach is adopted to examine in more detail possible explanations of the persistence in performance of such portfolios. Focusing on these portfolios

⁶ The trusts are also compared to a notional index fund where the annual fee of 1/2% or 1% is subtracted from the annual return on the market index.

⁷ Elton et al. (1996) propose an alternative methodology to evaluate performance persistence.

can be motivated by a recent study of [Del Guercio and Tkac \(2000\)](#) who document that the past returns of a fund is a more significant determinant in the allocation of new cash flows into funds by mutual fund investors compared to other performance measures.⁸ At the start of each year, trusts are ranked on the basis of their annual excess return over the prior year (ranking period) and grouped either into decile or quartile portfolios. Equally weighted monthly excess returns are then estimated on each portfolio over the next 12 months (evaluation period). If a trust disappears during this period then the trust will remain in the portfolio up to that point. This process is repeated each year and generates a time-series on monthly portfolio excess returns for each decile or quartile.⁹ The performance of the decile or quartile portfolios is then evaluated using both Eqs. (2) and (3).

One issue that can affect inferences about performance persistence is the question of survivorship bias. [Brown et al. \(1992\)](#) argue that using a survivorship only sample of funds can create the illusion of positive persistence. This arises whenever funds disappear because of poor performance over a single period. [Brown et al. \(1992\)](#) show that high volatility funds, which survive, will tend to have higher returns and tend to be winners in both periods. [Grinblatt and Titman \(1992\)](#) point out that repeat losers are most likely to disappear from a survivorship biased sample of funds, which implies we are more likely to find reversals in performance. This argument posits that funds disappear due to poor performance over more than one period.

Recent papers by [Carpenter and Lynch \(1999\)](#) and [Carhart et al. \(2000\)](#) provides more extensive evidence on the issues of various biases in performance persistence tests. These papers distinguish between two related biases of survivorship bias and look-ahead bias. Survivorship bias is defined as only including funds in the sample that are alive at the end of the sample period. Look-ahead bias is defined as requiring funds to survive a minimum period of time. [Carhart et al. \(2000\)](#) point out that methodologies, which test for performance persistence, can be subject to look-ahead bias. [Carhart et al. \(2000\)](#) examine the impact of these biases in the performance persistence tests of [Carhart \(1997\)](#) for US mutual funds. They find that using either a survivorship biased or look-ahead biased sample of funds tends to weaken the evidence of persistence but the impact is greater for the survivorship biased sample of funds. They attribute this to US mutual funds disappearing primarily to poor performance over more than one period.¹⁰

The methodology of [Brown and Goetzmann \(1995\)](#) does suffer from look-ahead bias as trusts are required to exist over the consecutive periods of time to be included in the tests. However, [Carhart et al. \(2000\)](#) argue that the methodology of [Carhart \(1997\)](#) minimises the impact of look-ahead bias as funds are still included in the decile portfolios even if they disappear over the evaluation period.

⁸ This differs from pension funds. [Del Guercio and Tkac \(2000\)](#) attribute this to mutual fund investors being less financially sophisticated compared to pension fund clients.

⁹ We also experimented with ranking trusts on the basis of alternative performance measures, such as the Jensen performance of the trusts over the prior 3 years relative to the FTA model. This tended to yield similar inferences.

¹⁰ [Lunde et al. \(1999\)](#) find evidence that prolonged underperformance by a trust is a major factor in the disappearance of UK funds.

3. Data and benchmarks

3.1. Unit trust data

The performance persistence tests of the unit trusts is examined between January 1982 and December 1996. The sample of trusts was selected to include all trusts with UK equity objectives recorded in the annual Unit Trust Yearbooks. Using the 1982 Yearbook, all trusts with UK General, UK Income and UK Growth objectives were identified. Using the subsequent Yearbooks, new trusts were added to the sample if they had UK equity objectives. The history of each trust was tracked throughout the sample period. If the trusts were taken over or wound up or changed objective to a non-UK equity objective, then monthly returns were collected on these trusts until these events. Name changes and transfers of unit trusts were treated as a continuation of the original trust. No survivorship requirements were imposed on the trusts. There were 724 trusts in the final sample. There were 139 trusts with continuous return data throughout the period. Additional information was also collected from the Yearbooks on the characteristics of the trusts. This includes the investment objective, size, annual charge and load charge of the trust at the start of each year.

The monthly unit trust returns are continuously compounded returns and are calculated from the offer prices at the end of the month and dividends paid by the trusts in the month the dividend is declared ex-dividend. The offer price of the trust includes the load charge,¹¹ brokerage fees and stamp duty. Adjustments are also made for the accrued annual management charge not yet paid. The returns on the trusts can be viewed approximately as gross of the load charge and trading costs but net of the management charge. Excess returns on the trusts were calculated using the monthly return on a 30-day UK Treasury Bill as the risk-free asset. This is obtained from the London Business School Share Price Database (LBS).

The monthly returns of the trusts were collected from a variety of sources. If trusts were in existence at August 1993, then monthly returns could be calculated up to this period from the Finstat managed fund database provided by the Financial Times Information Service. In addition a second managed fund dataset was purchased from Finstat which allowed the monthly returns of trusts that were in existence at the end of 1996 to be calculated from September 1993 to December 1996. Where the information to compute returns was not on Finstat, the monthly offer prices were collected from Money Management and dividends from the Extel UK Dividend and Fixed Interest Record.

3.2. Benchmark portfolios

A number of models are used to estimate the unconditional and conditional Jensen measures. The first model is the single factor model, which uses the excess return on the Financial Times All Share (FTA) index¹² obtained from LBS. The second model is based

¹¹ Nearly all unit trusts have a load charge.

¹² The FTA index is a value-weighted index of the largest companies on the London stock exchange.

on Carhart (1997). This is an extension of the Fama and French (1993) three-factor model and constructs a fourth factor which captures the momentum anomaly of Jegadeesh and Titman (1993). The three additional factors are used because they are known to explain the cross-sectional predictability in US stock returns.¹³ Another reason for using a similar model to Carhart (1997) is that the additional factors control for the possibility of generating abnormal returns through following mechanical trading strategies, e.g. momentum or value.¹⁴

The first factor is the excess return on the FTA index. The size and book to market factors are constructed from six portfolios as in Fama and French (1993). To be included in the six portfolios, companies require a non-zero market value at the start of the year from LBS and a book value of common equity from Datastream in the previous calendar year. Only non-financial companies are included in the construction of the six portfolios.

Beginning in July 1981, stocks were ranked in ascending order on the basis of their market value at the start of the year and grouped into two portfolios. Within each group, the stocks are ranked on their Book/Market ratio (BM) in ascending order and grouped into three portfolios. This gives six size/BM portfolios (S/L, S/M, S/H, B/L, B/M, B/H). Equally weighted monthly returns are estimated over the next 12 months. This process is repeated each year until July 1995. For the final year, equally weighted returns are calculated for eighteen months. The portfolio returns between 1982 and 1996 are used to form the size and BM portfolios. The size factor is the difference in the mean return of three small firm portfolios and the mean return of three large portfolios (SMB). The BM factor is the difference in the mean return of the two high BM portfolios and the mean return of the two low BM portfolios (HML).

The fourth factor is constructed in a similar manner to Carhart (1997). Beginning in January 1982, all securities on LBS are ranked on the basis of their cumulative return over the past 11 months and grouped into the top and bottom third of companies. Over the next month, a portfolio is formed which is the difference between the mean return of the top 1/3 and the mean return of the bottom 1/3. This is repeated each month to generate a time-series of return observations on the momentum factor¹⁵ (Mom).

The third model is based on the APT. This follows from Connor and Korajczyk (1991) who develop an approach to form mimicking portfolios of the factors by combining the use of statistical factors based on principal components analysis (Connor and Korajczyk,

¹³ A number of studies show that similar patterns exist in UK stock returns. Fama and French (1998) show that a book to market effect exists in UK stock returns (see also Strong and Xu, 1997), Strong and Xu (1997) for the size effect and Rouwenhorst (1998) for a momentum effect in UK stock returns.

¹⁴ This obviously does not include the complete list of anomalies in UK stock returns or the mechanical trading strategies that investors follow.

¹⁵ The importance of the size effect has declined in recent years in UK stock returns (Dimson and Marsh, 1999). However the inclusion of the size factor can be motivated also by the arguments in Elton et al. (1993) that mutual funds invest in asset categories not included in the market index e.g. small stocks. In unreported results, the majority of trusts have a significant positive exposure to the SMB factor. In addition, many more trusts have significant coefficients with respect to the SMB factor than either the HML or momentum factors. This implies that the SMB factor is a relevant factor in unit trust returns.

1986) and pre-specified economic factors as in [Chen et al. \(1986\)](#). The following four factors are used:

- (i) Excess return on the FTA index.
- (ii) Term structure—difference in monthly returns between 15 year UK government bonds (obtained from Datastream) minus the UK risk-free return.
- (iii) Monthly percentage change in UK industrial production (obtained from Datastream).
- (iv) Monthly percentage change in UK inflation (obtained from LBS).

[Shanken \(1992\)](#) points out that if the factor is a portfolio return, then the excess return is the best estimate of the risk premium. The [Connor and Korajczyk \(1991\)](#) technique is used to form mimicking portfolios for factors (iii) and (iv). The statistical factors are estimated using the first five eigenvectors of the cross-products matrix of excess returns of all stocks on the LBS database between January 1982 and December 1996 that allows for stocks to have missing return observations (see [Heston et al., 1995](#) for details). The first step is to regress the de-meaned (actual factor realization minus the average value) factors (iii) and (iv) on the de-meaned eigenvectors. The coefficients from the regression are then multiplied by the original eigenvectors to get the estimated factor portfolios.¹⁶

To estimate the conditional measures in Eq. (3), the information set of investors requires to be specified.¹⁷ Following [Ferson and Warther \(1996\)](#) and [Christopherson et al. \(1998\)](#), the following two information variables¹⁸ are used:

- 1. Lagged dividend yield on the FTA index (obtained from LBS).
- 2. Lagged 1-month risk-free return.

4. Empirical results

The first test of persistence examines the persistence of the relative rankings of trusts using the annual cumulative excess returns similar to [Brown and Goetzmann \(1995\)](#). Does the relative rank of the trust over the past year provide information about the relative rank in the next year? Winners are defined as greater than the median annual cumulative excess returns across all trusts. [Table 1](#) reports the repeat winner tests over consecutive 1-year periods. The table includes the number of trusts within each of the four categories, the log-odds ratio and *z* test.

¹⁶ We also estimated performance using a three-factor model similar to [Elton et al. \(1993\)](#). This includes the excess returns on the FTA index, a small stock index (bottom decile on the LBS database) and a UK government bond index.

¹⁷ [Keim and Stambaugh \(1986\)](#) and [Fama and French \(1988\)](#) amongst others show that stock returns are partly predictable over time.

¹⁸ [Blake and Timmermann \(1998\)](#) use similar instruments in their study.

Table 1
Repeat winner tests: annual returns

	WW	WL	LW	LL	Log-odds	z
82–83	54	63	63	53	– 0.33	– 1.24
83–84	71	46	46	71	0.86	3.24 *
84–85	80	41	41	79	1.32	4.87 *
85–86	80	56	56	79	0.71	2.84 *
86–87	105	46	46	105	1.65	6.60 *
87–88	93	74	74	93	0.46	2.07 *
88–89	83	96	96	82	– 0.3	– 1.43
89–90	119	74	74	119	0.95	4.53 *
90–91	102	97	97	102	0.1	0.5
91–92	98	102	102	98	– 0.08	– 0.39
92–93	95	101	101	94	– 0.13	– 0.66
93–94	119	76	76	119	0.89	4.32 *
94–95	94	102	102	93	– 0.17	– 0.86
95–96	134	70	70	134	1.29	6.23 *
All	1327	1044	1044	1321	0.47	8.11 *

The repeat winner tests of [Brown and Goetzmann \(1995\)](#) are estimated over consecutive annual intervals between 1982 and 1996. Winners are defined as above the median annual excess returns across all trusts. Columns two to five include the number of trusts in the winner/winner (WW), winner/loser (WL), loser/winner (LW) and loser/loser (LL) categories. The log-odds ratio is defined as the $\ln[(WW*LL)/(WL*LW)]$ and tests the null hypothesis of no persistence in performance. The final column is the *z* test of the log-odds ratio, which has a standard normal distribution.

* Significant at 5%.

[Table 1](#) shows that there is a significant persistence in the relative performance rankings using annual excess returns for 8 out of the 14 periods. In addition, the aggregate results also indicate that there is significant persistence in performance. This implies that Winners (Losers) are more likely to remain Winners (Losers) than to move to the Losers (Winners) category. Much of the persistence is concentrated in the 1980s where for five successive intervals there is significant persistence (1983/1984 to 1987/1988). This is consistent with [Malkiel \(1995\)](#) and [Brown and Goetzmann \(1995\)](#) who find that persistence can be clustered in time for US mutual funds. [Rhodes \(2000\)](#) also documents that persistence is clustered in the early 1980s for UK trusts. There is no evidence of any significant reversal in performance.

The robustness of persistence is also evaluated where different performance measures are used. This includes comparing the annual excess returns of the trust to other trusts within the same investment sector¹⁹ and the Jensen (conditional Jensen) measures relative to the FTA and Carhart models over consecutive 2-year subperiods. Using different performance measures tends to have little impact on the inferences in [Table 1](#), except for the conditional Jensen measure for the FTA model. We also considered whether any of the trust characteristics at the start of the year predict the relative annual excess returns.

¹⁹ This is the most popular performance league table for UK trusts reported by financial periodicals such as Money Management.

Table 2

Repeat winner tests: market adjusted returns

	WW	WL	LW	LL	Log-odds	<i>z</i>
82–83	29	35	100	69	– 0.55	– 1.89
83–84	28	105	12	89	0.68	1.82
84–85	31	14	82	114	1.12	3.18 *
85–86	92	41	65	73	0.92	3.64 *
86–87	162	14	85	41	1.72	5.09 *
87–88	52	226	9	47	0.18	0.46
88–89	2	67	11	277	– 0.28	– 0.36
89–90	7	10	79	290	0.94	1.85
90–91	14	79	51	254	– 0.12	– 0.38
91–92	18	48	82	252	0.14	0.46
92–93	54	47	152	138	0.04	0.18
93–94	104	105	49	132	0.98	4.52 *
94–95	35	119	37	200	0.46	1.76
95–96	42	34	87	245	1.24	4.75 *
All	670	944	901	2221	1.749	27.27 *

The repeat winner tests of [Brown and Goetzmann \(1995\)](#) are estimated over consecutive annual intervals between 1982 and 1996. Winners are defined as a positive market-adjusted return. The market-adjusted return is the cumulative annual excess return of the trust minus the annual excess return on the FTA index. Columns 2–5 include the number of trusts in the winner/winner (WW), winner/loser (WL), loser/winner (LW) and loser/loser (LL) categories. The log-odds ratio is defined as the $\ln[(WW*LL)/(WL*LW)]$ and tests the null hypothesis of no persistence in performance. The final column is the *z* test of the log-odds ratio which has a standard normal distribution.

* Significant at 5%.

Winners in the first period are defined as above the median characteristic across all trusts and in the second period as above the median annual excess return. There is little evidence that large trusts, trusts with higher charges predict the relative rankings of annual excess returns.

The findings in [Table 1](#) suggest that when trusts are ranked relative to one another, that there is significant persistence in repeat winners and losers. Does this reflect superior investment ability by the repeat winners? This is examined by comparing the performance of the trusts to an absolute benchmark. Winners are defined where the annual excess returns of the trust exceed that of the FTA market index²⁰, i.e. market-adjusted returns are positive. If there is persistence in repeat winners, then this may provide some evidence of superior performance ability. [Table 2](#) reports the repeat winner tests for the market adjusted returns over consecutive 1-year periods. The table includes the number of trusts within each of the four categories, the log-odds ratio and *z* test.

[Table 2](#) shows that for the overall period, there is significant persistence in the performance of the trusts relative to the FTA index. However, this persistence is driven primarily by repeat underperformance. The number of repeat losers is over three times higher than the number of repeat winners. A similar picture emerges over the individual

²⁰ The results were also estimated by reducing the annual excess return of the FTA by a notional annual management charge of 1/2% or 1%. These are common annual charges of a number of index tracking trusts within the UK.

Table 3
Performance of quartile portfolios

	Mean	Standard deviation	FTA α	β	APT α
1 Losers	0.264	4.76	−0.269 (−2.18) *	0.903 (36.44) *	−0.166 (−1.48)
2	0.356	4.66	−0.183 (−2.16) *	0.911 (53.86) *	−0.102 (−1.30)
3	0.417	4.69	−0.128 (−1.59)	0.921 (57.48) *	−0.036 (−0.51)
4 Winners	0.466	4.74	−0.069 (−0.61)	0.906 (39.50) *	0.098 (1.05)
W–L	0.202	1.41	0.200 (1.82)	0.003 (0.15)	0.264 (2.33) *

The performance of the quartile portfolios of trusts formed on the basis of prior year excess return is estimated between January 1983 and December 1996. Columns 2 and 3 include the mean and standard deviation of monthly excess returns of the quartile portfolios. Columns 4 and 5 report the Jensen performance measure (α) and betas with respect to the FTA model. The final column contains the Jensen performance with respect to the APT model. The t statistics, in parentheses, are adjusted for heteroskedasticity using White (1980). The bottom row reports the performance of a zero-cost portfolio (W–L) which is long in the top quartile portfolio of trusts and short in the bottom quartile portfolio of trusts. All of the performance numbers are monthly percentages.

* Significant at 5%.

periods. In 11 out of 14 cases, the number of repeat losers is higher than repeat winners. There are some individual periods where there is evidence of repeat winners driving the significant persistence (e.g. 1987/1988) but this tends to be the exception. Similar results occur when the trusts are compared to a notional index fund.²¹ These results are interesting because it implies that trusts are generally unable to outperform the market in spite of all the evidence of the existence of mechanical trading strategies that appear to generate abnormal returns.²²

The robustness of the results were also examined where winners are defined as positive Jensen (conditional Jensen) performance relative to either the FTA or Carhart models. A similar picture emerges as in Table 2. The significant persistence is driven by repeat underperformance. This implies that persistence in performance is not primarily a reflection of superior investment skill. In contrast the persistence of inferior performance is more common. This is consistent with Carhart (1997) and Christopherson et al. (1998) amongst others.

The tests in Tables 1 and 2 are based on simple 2*2 sorts of trusts. The next set of tests employs the methodology of Carhart (1997) to examine the persistence in performance of portfolios of trusts, based on finer sorts, formed on the basis of prior year excess returns. Table 3 reports summary statistics of performance between January 1983 and December 1996 for the quartile portfolios of trusts formed on the basis of prior year excess returns described earlier.²³ This includes the mean and standard deviation of monthly excess returns and the Jensen performance and betas of the portfolios with

²¹ These results are consistent with the observations in the financial press that different types of funds have struggled to outperform market aggregates in recent years. This is one reason for the growth of index funds in recent years.

²² See also the discussion in Cochrane (1999).

²³ We do not report the results for the decile portfolios in any of the subsequent as they yield a similar picture to the quartile portfolios.

Table 4

Performance of quartile portfolios: Carhart model

	α	β_{FTA}	β_{SMB}	β_{HML}	β_{Mom}
1 Losers	−0.180 (−1.28)	0.959 (39.37)*	0.272 (5.21)*	0.036 (0.49)	−0.199 (−4.87)*
2	−0.189 (−1.92)	0.946 (55.34)*	0.165 (4.52)*	0.079 (1.52)	−0.116 (−4.05)*
3	−0.194 (−2.07)*	0.965 (59.37)*	0.202 (5.80)*	0.076 (1.54)	−0.059 (−2.17)*
4 Winners	−0.297 (−2.35)*	0.994 (45.42)*	0.391 (8.35)*	0.124 (1.86)	0.015 (0.42)
W–L	−0.117 (−0.87)	0.035 (1.52)	0.119 (2.42)*	0.088 (1.25)	0.214 (5.53)*

The performance of the quartile portfolios of trusts formed on the basis of prior year excess return is estimated between January 1983 and December 1996 relative to the Carhart model. The table includes the Jensen performance (α) and betas with respect to the four factors in the Carhart model. The t statistics, in parentheses, are adjusted for heteroskedasticity using White (1980). The bottom row reports the performance of a zero-cost portfolio (W–L) which is long in the top quartile portfolio of trusts and short in the bottom quartile portfolio of trusts. All of the performance numbers are monthly percentages.

* Significant at 5%.

respect to the FTA model. The final column includes the Jensen performance relative to the APT model. The t statistics, in parentheses, are corrected for heteroskedasticity using White (1980) in this and subsequent tables. Table 4 includes the Jensen performance and betas with respect to the four factors in the Carhart model. The final row of each table present the performance results of a self-financing portfolio which estimates the performance of a portfolio (W–L) that is long in the top quartile (Winners) of trusts and short in the bottom quartile (Losers) of trusts. Under the null hypothesis of no persistence in performance, the performance of the W–L portfolio should equal zero. All performance numbers are monthly percentages.

Table 3 shows that there is a monotonic relationship in the monthly mean excess returns across the quartile portfolios. The annualized average excess return of the W–L portfolio is 2.42%. The quartile portfolios of trusts have similar standard deviations. This implies that the Sharpe (1966) performance of the top quartile of trusts is nearly 40% greater than the bottom quartile of trusts. The pattern in the mean monthly excess returns suggests that there is significant persistence in the relative rankings of trusts for finer sorts of trusts than the simple winners/losers sort used earlier.

The Jensen performance of the quartile portfolios relative to the FTA model reveals a similar picture to the mean excess returns. The Jensen performance across the quartiles exhibits a monotonic pattern. The bottom two quartiles exhibit significant underperformance relative to the FTA model. The performance for the W–L portfolio is similar to the mean monthly excess returns, although it is only statistically significant at the 10% significance level. The inability of the FTA model to explain the persistence in performance can be explained by the market betas of the quartile portfolios showing little variation. Using the APT model to estimate performance tends to strengthen the findings of positive persistence. The positive performance of the W–L portfolio relative to the APT model is greater than the corresponding mean monthly excess returns and the performance of the FTA model. Furthermore the performance is significantly positive.²⁴

²⁴ Using a model similar to Elton et al. (1993) reveals a similar picture to that in Table 3.

The evidence in Table 3 suggests that there is significant persistence in the performance of portfolios formed on the basis of prior year excess returns. Furthermore, neither the FTA nor APT models cannot explain the persistence in performance. These results are generally consistent with Carhart (1997) or Blake and Timmermann (1998). It is interesting to note in Table 3, that although the performance of the W–L portfolio is economically large, that the top quartile of trusts does not exhibit significant positive performance relative to either the FTA or APT models. The performance of the W–L portfolio is due more to the underperformance of the bottom quartile portfolio. This provides suggestive evidence that the persistence in performance is not a reflection of superior performance ability.

Table 4 shows that when performance is estimated relative to the Carhart model, that the persistence in performance is eliminated. There is more evidence of reversals in performance. The performance of the W–L portfolio is no longer positive and is now actually insignificantly negative. The top quartiles of trusts exhibit significant underperformance relative to the Carhart model. The beta estimates in Table 4 relative to the different factors show a great deal of variation across the quartile portfolios. The top quartile portfolio has a significant positive exposure on the SMB and HML factors (at the 10% significance level) and a small positive exposure on the momentum factor. This implies that the top quartile portfolio of trusts is exposed to smaller companies and companies with higher book to market ratios. The bottom quartile portfolio has a significant positive exposure to the SMB factor and a significantly negative exposure on the momentum factor. This implies that the bottom quartile portfolio is more exposed to stocks with poor recent returns. The beta estimates of the W–L portfolio illustrate that these differences are statistically significant for the SMB and momentum factors.

The impact of additional factors in the Carhart model on the performance of the quartile portfolios compared to the FTA model can be examined further by using the analysis in Pastor and Stambaugh (in preparation). This is extremely useful because we can isolate the impact that each factor has on the different performance inferences between the two models. Pastor and Stambaugh (in preparation) show that for a given asset, the differences in Jensen performance between two models, where one model is a subset of the other, can be explained by two reasons. These are the betas that the asset has with respect to the additional factors and the Jensen performance of the additional factors with respect to the subset model. Within this framework, the additional factors not included in the subset model are defined as non-benchmark assets. The FTA model is a subset of the Carhart model. The difference in Jensen performance of a portfolio of trusts between the FTA and Carhart models can be written as:

$$\alpha_i - \alpha_j = B'\alpha \quad (4)$$

where α_i is the Jensen performance of the portfolio relative to the subset model (FTA), and α_j is the Jensen performance of the portfolio relative to the larger model (Carhart), B is a $(M \times 1)$ vector of portfolio betas with respect to the factors not included in the subset model, α is a $(M \times 1)$ vector of Jensen measures of the non-benchmark assets relative to the subset model and M is the number of non-benchmark assets. By comparing the Jensen performance in Tables 3 and 4, we can see that $\alpha_i - \alpha_j$ is positive for all portfolios except the bottom quartile portfolio.

The Jensen performance of the SMB, HML and Mom factors with respect to the FTA model are 0.232%, 0.990% and 0.945%, respectively. By using these Jensen measures and the betas in Table 4, we can examine which factors most explain the performance differences. The performance of the bottom quartile portfolio of trusts improves under the Carhart model relative to the FTA model because of the negative exposure to the momentum factor and the positive Jensen performance that this factor has with respect to the FTA model. In contrast, the top quartile portfolio of trusts has poorer performance under the Carhart model relative to the FTA model due to the impact of the SMB and HML factors. This is because of the positive exposures that the top quartile portfolio has with respect to the SMB and HML factors and the positive Jensen performance that these factors have with respect to the FTA model. The role of the momentum factor is limited in this case. This analysis implies that it is not the same factors that are most important explaining why the performance is different between the FTA and Carhart models across the quartile portfolios. The Jensen performance of the W–L portfolio is poorer under the Carhart model compared to the FTA model due to the impact of the HML and momentum factors. The positive betas that the W–L portfolio has with respect to these factors and the positive Jensen performance that these factors have with respect to the FTA model explains the substantial shift in performance. The momentum factor has the greatest impact on the difference in performance for the W–L portfolio. This is due to the negative exposure that the bottom quartile portfolio has on the momentum factor.

The results in Table 4 are generally consistent with Carhart (1997) but do highlight some differences. The first is that the evidence of the persistence in performance is eliminated through the use of a model similar to Carhart (1997). The second main difference is that although the momentum factor plays a major role in explaining the differences of the W–L portfolio, this is not due to the top quartile portfolio having a significant positive exposure to the momentum factor as is observed in Carhart (1997).

Table 5 examines the impact of conditional performance measurement on the persistence in performance using the measure of Ferson and Schadt (1996). The table includes

Table 5
Conditional performance of quartile portfolios

	FTA α	Wald	APT α	Wald	Carhart α	Wald
1 Losers	–0.237 (–1.89)	0.154	–0.151 (–1.29)	0.062	0.009 (0.07)	0
2	–0.181 (–2.09) *	0.987	–0.152 (–1.86)	0.086	–0.162 (–1.63)	0
3	–0.112 (–1.37)	0.580	–0.118 (–1.62)	0.009	–0.205 (–2.19) *	0
4 Winners	–0.013 (–0.11)	0.030	–0.002 (–0.02)	0.013	–0.327 (–2.48) *	0.002
W–L	0.224 (2.02) *	0.089	0.148 (1.34)	0	–0.336 (–2.59) *	0

The performance of the quartile portfolios of trusts formed on the basis of prior year excess return is estimated between January 1983 and December 1996 using the conditional performance measure of Ferson and Schadt (1996). For each model, the conditional performance (α) and a Wald test (p value) of the hypothesis that the conditional portfolio beta is constant through time are reported. The t statistics, in parentheses, are adjusted for heteroskedasticity using White (1980) as is the Wald test. The following models are used: FTA (Columns 2 and 3), APT (Columns 4 and 5) and Carhart (Columns 6 and 7). The bottom row reports the performance of a zero-cost portfolio (W–L) which is long in the top quartile portfolio of trusts and short in the bottom quartile portfolio of trusts. All of the performance numbers are monthly percentages.

* Significant at 5%.

the estimated conditional Jensen performance and corresponding t statistics, in parentheses, for the quartile portfolios relative to the FTA, APT and Carhart models. The table also includes a Wald test (p value) which examines the hypothesis that the conditional portfolio betas are constant through time. Under the null hypothesis that the conditional portfolio betas are constant, the slope coefficients on the interaction terms between the information variables and factors should be jointly equal to zero.

The use of the conditional Jensen measure with the FTA model has no impact on the inferences about performance persistence. There is a monotonic relationship in performance across the quartile portfolios. The performance of the W–L portfolio is significantly positive and is greater than the corresponding Jensen performance. The lack of impact of the conditional Jensen measure is confirmed by the Wald test, which is generally unable to reject the null hypothesis that the conditional portfolio betas are constant through time.

The use of the conditional Jensen measure has more of an impact with the APT model. The degree of performance persistence is reduced compared to that in Table 3. The performance of the W–L portfolio is lower than that of the Jensen performance and is no longer significantly positive. The Wald test highlights the increased importance of the conditional measure by rejecting the null hypothesis that the conditional portfolio betas are constant through time at the 10% significance level.

When the conditional Jensen measure is used for the Carhart model, there is stronger evidence of significant reversals in performance. There is a monotonic relationship in performance across the quartile portfolios that is in the opposite direction to that which we observe in the mean monthly excess returns. The top quartile portfolios exhibit significant underperformance relative to the Carhart model. The conditional Jensen performance of the W–L portfolio is significantly negative and economically large. The Wald test²⁵ is able to reject the null hypothesis of constant conditional betas for all quartile portfolios providing further support for the significance of conditioning information. The evidence in Table 5 suggests that conditional performance measurement have some impact on inferences.²⁶ Overall the results in Tables 3–5 suggest that the persistence in trust performance is not a manifestation of superior stock selection ability.

Since the performance persistence does not appear to reflect superior stock selection ability, are there any differences in the characteristics of the trusts in the quartile portfolios? Panel A of Table 6 reports the characteristics of the quartile portfolios. This includes the time-series average of the cross-sectional means of the annual charge, load charge and size of the trusts at the beginning of each evaluation year between 1983 and 1996. The table also includes the time-series average of the cross-sectional mean % cash flow for the quartile portfolios in the evaluation year between 1983 and 1995. This is calculated as in Del Guercio and Tkac (2000). The % cash flow = $(TNA_{it} - TNA_{it-1} * (1 + R_{it})) / TNA_{it-1}$ where TNA_{it} is the size of the trust at the end of the year, TNA_{it-1} is the size of the trust at the start of the year and R_{it} is the

²⁵ Using F test yields similar inferences to the Wald test.

²⁶ The positive persistence in performance is also eliminated when the conditional performance is estimated using the three-factor model similar to Elton et al. (1993).

Table 6
Characteristics of quartile portfolios

Panel A	Annual	Load	Size	% Cash flow
1 Losers	1.026	4.999	48.45	0.188
2	1.010	5.044	64.88	0.095
3	0.995	5.069	70.10	0.208
4 Winners	1.002	5.021	46.89	0.447
Panel B	1 Losers	2	3	4 Winners
1st year	0.196	0.287	0.351	0.396
2nd year	0.249	0.306	0.256	0.246
3rd year	0.236	0.211	0.237	0.250

At the start of each year between 1983 and 1996, trusts are ranked on the basis of their prior year excess returns and grouped into quartile portfolios. The average values of various characteristics of the trusts in the different quartile portfolios are calculated. This includes the annual charge, load charge and size at the start of the year. The percentage cash flow during the year is also estimated. The mean monthly excess return is calculated in the 1st year, 2nd year and 3rd year for the quartile portfolios after the initial ranking. Panel A reports the time-series mean for each characteristic for the respective quartiles (the % cash flow is calculated to the end of 1995). Panel B describes the time-series average of the mean monthly excess returns of the quartile portfolios in the first, second and third years after ranking. Trusts are ranked up to the start of 1994 for the results in panel B.

annual return of the trust. [Panel B of Table 6](#) examines whether the persistence in performance as reflected in the mean monthly excess returns of the quartile portfolios extends beyond the evaluation year. This includes the time-series average of the mean monthly excess returns of the quartile portfolios in the first, second and third years after ranking.

[Table 6](#) shows that there is little difference in the average charges of the trusts across the quartile portfolios. This implies that cross-sectional differences in average expenses explain none of the persistence in the mean monthly excess returns. This differs from [Elton et al. \(1996\)](#) and [Carhart \(1997\)](#) where the poorest performing funds tended to have the highest expense ratios. Furthermore there is no noticeable difference between the average size of the trusts in the top and bottom quartiles. The one area of difference is the average % cash flow. The top quartile of trusts attracts the largest proportion of cash flows. Furthermore there is little penalty for the poorest performing quartile. This is consistent with [Sirri and Tufano \(1998\)](#). By attracting the lion's share of new cash flows, extreme winning funds increase their net assets under management and maximise fee income. A number of studies examine the implications that the tournament structure can have on trading behaviour of mutual funds (see [Brown et al., 1996](#); [Karceski, 2000](#) amongst others). This raises an interesting question as to whether trust managers are aware of this and adopt a trading strategy which helps them perform well in annual league tables of returns. This argument would suggest that the composition of trusts in the top quartile portfolio may be stable through time and the performance in mean returns extends beyond the evaluation year. An alternative argument put forward by [Carhart \(1997\)](#) is that persistence may be a short-run phenomenon and decile portfolios exhibit frequent turnover. Mutual funds end up in particular deciles because they happen to hold stocks, which are doing well or badly.

The mean monthly excess returns in [panel B of Table 6](#) show that the persistence in performance is a short-run phenomenon. In the year after ranking, there is a wide spread in the mean monthly excess returns across the quartile portfolios. However in the second and third year, this spread disappears. The mean monthly excess returns are fairly similar across the quartile portfolios. This is consistent with [Carhart \(1997\)](#). There is also a high degree of turnover the quartile portfolios from year to year. This provides suggestive evidence that the short-run persistence is driven not by trusts following mechanical trading strategies but by trusts which happen to hold stocks that are currently doing well or badly.

5. Conclusions

This paper has examined the persistence in performance of UK unit trusts between January 1982 and December 1996. Consistent with the prior research of mutual funds and unit trusts, the paper finds evidence of significant persistence in the relative rankings of unit trusts in annual excess returns. This persistence is generally robust to alternative performance measures. However, when the rankings of the trusts depend upon performance compared to an absolute benchmark, then the significant persistence is driven by repeat underperformance.

The paper also documents that there is significant persistence in the performance of portfolios, which are formed on the basis of prior year excess returns, when performance is evaluated with mean monthly excess returns and various factor models such as the CAPM or APT. These results are consistent with prior research. However, when performance is evaluated relative to the Carhart model, this persistence in performance is eliminated. Using a conditional performance measure leads to the finding of significant reversals in performance. The performance of the W–L portfolio is significantly negative and economically large. These results are stronger than that observed in US mutual funds as in [Carhart \(1997\)](#). The change in the performance inferences of the W–L portfolio in moving from the FTA model to the Carhart model can be attributed to the bottom quartile portfolio having a significant negative exposure to the momentum factor and being less exposed to value stocks. This differs from [Carhart \(1997\)](#) where the top decile portfolio has a significantly positive exposure to the momentum factor.

Subsidiary findings of the paper highlight that there are no differences in the average charges of the trusts in the top and bottom quartiles and the persistence in performance of the mean monthly excess returns is a short-run phenomenon. The findings of the paper suggests that the persistence in performance of UK trusts is not a manifestation of superior stock selection strategy and can be explained by factors that are known to capture cross-sectional differences in stock returns.

Acknowledgements

Helpful comments received from Paul Draper, seminar participants at the University of Strathclyde and an anonymous reviewer.

References

- Allen, D.E., Tan, M.L., 1999. A test of the persistence in the performance of UK managed funds. *Journal of Business Finance and Accounting* 26, 559–595.
- Blake, D., Timmermann, A., 1998. Mutual fund performance: evidence from the UK. *European Finance Review* 2, 57–77.
- Brown, S.J., Goetzmann, W., 1995. Performance persistence. *Journal of Finance* 50, 679–698.
- Brown, S.J., Goetzmann, W., Ibbotson, R.G., Ross, S.A., 1992. Survivorship bias in performance studies. *Review of Financial Studies* 4, 553–580.
- Brown, K., Harlow, W.V., Starks, L.T., 1996. Of tournaments and temptations: an analysis of managerial incentives in the mutual fund industry. *Journal of Finance* 5, 85–110.
- Brown, G., Draper, P.R., McKenzie, E., 1997. Consistency of UK pension fund performance. *Journal of Business Finance and Accounting* 24, 155–178.
- Brown, S.J., Goetzmann, W., Ibbotson, R.G., 1999. Offshore hedge funds: survival and performance 1989–95. *Journal of Business* 71, 91–118.
- Carpenter, J.N., Lynch, A.W., 1999. Survivorship bias and attrition effects in measures of performance persistence. *Journal of Financial Economics* 54, 337–374.
- Carhart, M., 1997. On persistence in mutual fund performance. *Journal of Finance* 52, 57–82.
- Carhart, M., Carpenter, J.N., Lynch, A.W., Musto, D.K., 2000. Mutual fund survivorship. Unpublished manuscript.
- Chen, N., Roll, R., Ross, S.A., 1986. Economic forces and the stock market. *Journal of Business* 59, 383–403.
- Christopherson, J.A., Ferson, W.E., Glassman, D., 1998. Conditioning manager alphas on economic information: another look at the persistence of performance. *Review of Financial Studies* 11, 111–142.
- Christopherson, J.A., Ferson, W.E., Turner, A.L., 1999. Performance evaluation using conditional alphas and betas. *Journal of Portfolio Management* 25, 59–72.
- Cochrane, J.H., 1999. New facts in finance. *Journal of Economic Perspectives* 23, 36–58.
- Connor, G., Korajczyk, R.A., 1986. Performance measurement with the Arbitrage Pricing Theory: a new framework for analysis. *Journal of Financial Economics* 15, 373–394.
- Connor, G., Korajczyk, R.A., 1991. The attributes, behavior and performance of US mutual funds. *Review of Quantitative Finance and Accounting* 1, 5–26.
- Del Guercio, D., Tkac, P.A., 2000. The determinants of the flow of funds of managed portfolios: mutual funds versus pension funds. Working Paper, University of Oregon.
- Dimson, E., Marsh, P., 1999. Murphy's law and market anomalies. *Journal of Portfolio Management* 25, 53–69.
- Elton, E.J., Gruber, M., Das, S., Hlvrka, M., 1993. Efficiency with costly information: a reinterpretation of evidence from managed portfolios. *Review of Financial Studies* 6, 1–22.
- Elton, E.J., Gruber, M., Blake, C., 1996. The persistence of risk-adjusted mutual fund performance. *Journal of Business* 69, 133–159.
- Fama, E.F., French, K.R., 1988. Dividend yields and expected stock returns. *Journal of Financial Economics* 22, 3–25.
- Fama, E.F., French, K.R., 1993. Common risk factors in the returns of stocks and bonds. *Journal of Financial Economics* 33, 3–56.
- Fama, E.F., French, K.R., 1998. Value versus growth: the international evidence. *Journal of Finance* 53, 1975–1999.
- Ferson, W.E., Schadt, R.W., 1996. Measuring fund strategy and performance in changing economic conditions. *Journal of Finance* 51, 425–462.
- Ferson, W.F., Warther, V.A., 1996. Evaluating fund performance in a dynamic market. *Financial Analysts Journal* 52, 20–28.
- Goetzmann, W.N., Ibbotson, R.G., 1994. Do winners repeat? Patterns in mutual fund performance. *Journal of Portfolio Management* 20, 9–17.
- Grinblatt, M., Titman, S., 1992. The persistence of mutual fund performance. *Journal of Finance* 47, 1977–1984.
- Gruber, M.J., 1996. Another puzzle: the growth in actively managed mutual funds. *Journal of Finance* 51, 783–810.

- Hendricks, D., Patel, J., Zeckhauser, R., 1993. Hot hands in mutual funds: short-run persistence of relative performance, 1974–88. *Journal of Finance* 48, 93–130.
- Heston, S.L., Rouwenhorst, K.G., Wessels, R.E., 1995. The structure of international stock markets and the integration of capital markets. *Journal of Empirical Finance* 2, 173–199.
- Jegadeesh, N., Titman, S., 1993. Returns to buying winners and selling losers: implications for stock market efficiency. *Journal of Finance* 48, 65–91.
- Jensen, M.C., 1968. The performance of mutual funds in the period 1945–1964. *Journal of Finance* 23, 389–416.
- Karceski, J., 2000. Returns—chasing behavior, mutual funds and beta's death. Working Paper, University of Florida.
- Keim, D.B., Stambaugh, R.F., 1986. Predicting return in the bond and stock market. *Journal of Financial Economics* 17, 357–390.
- Lunde, A., Timmermann, A., Blake, D., 1999. The hazards of mutual fund underperformance: a cox regression analysis. *Journal of Empirical Finance* 6, 121–152.
- Malkiel, B.G., 1995. Returns from investing in equity mutual funds 1971 to 1991. *Journal of Finance* 50, 549–572.
- Pastor, L., Stambaugh, R.F., 2001. Mutual fund performance and seemingly unrelated assets. *Journal of Financial Economics*, in preparation.
- Rhodes, M., 2000. Past imperfect? The performance of UK equity managed funds. Occasional Paper Series, Financial Services Authority.
- Rouwenhorst, G., 1998. International momentum strategies. *Journal of Finance* 53, 267–284.
- Shanken, J., 1992. On the estimation of beta pricing models. *Review of Financial Studies* 5, 1–55.
- Sharpe, W.F., 1966. Mutual Fund Performance. *Journal of Business* 39, 119–138.
- Sirri, E.R., Tufano, P., 1998. Costly search and mutual fund flows. *Journal of Finance* 53, 1589–1622.
- Strong, N., Xu, X.G., 1997. Explaining the cross-section of expected UK stock returns. *British Accounting Review* 29, 1–24.
- White, H., 1980. A heteroskedasticity-consistent covariance matrix and a direct test for heteroskedasticity. *Econometrica* 48, 817–883.
- Zheng, L., 1999. Is money smart? A study of mutual fund investors' fund selection ability. *Journal of Finance* 54, 901–933.