

# Deconstructing Performance Persistence: Benchmark- and Peer-Adjusted Mutual Fund Evaluation

**Preliminary work. Do not quote without permission**

## **Abstract**

This paper investigates features of performance persistence in actively managed mutual funds, when fund performance is measured relative to both funds' self-reported benchmarks and peer groups. Using a novel benchmark- and peer-group-adjusted performance model (MMT), we evaluate 735 US Large Cap Growth funds from 1990 to 2018. Our findings show that MMT-identified winner funds deliver persistent outperformance across 6-, 12-, and 24-month holding periods, with the strongest results for shorter rebalancing intervals. Crucially, this persistence holds across market regimes—including crisis periods—and survives realistic transaction cost estimates, supporting the model's practical feasibility. Outperformance is not limited to the select few top funds; even diversified portfolios of MMT winners generate superior returns. The MMT model's strength lies in its ability to isolate persistent fund manager skill by adjusting for benchmark misalignment and peer-group dynamics, offering a robust alternative to traditional alpha-based models.

**Keywords:** US equity mutual funds, Benchmark- and peer-group-adjusted performance persistence, holding periods, market conditions, trading costs

**JEL:** G11, G12, G2

## 1. Introduction

Research on active equity mutual funds has long grappled with performance persistence and the question of whether ‘winner’ funds continue to outperform. The answer depends on how performance is measured, i.e. what constitutes a winner, the horizon over which persistence is evaluated, and whether any alpha survives trading frictions. Classic persistence studies generally define winners using past returns or Fama-French (1993)/Carhart (1997) factor-model alphas; however, recent work argues that persistence is more likely to be detected when winner definitions reflect the manager’s actual competitive environment. Mateus, Mateus, and Todorovic (2023, 2025) re-align the definition of ‘true winner’ funds to both fund managers’ and investment client’s perspective, as those funds that simultaneously outperform their self-declared benchmarks and peer group. The benchmark and peer group-adjusted alpha model developed in Mateus et al. (2025), the ‘MMT model’ hereafter, documents that ‘true’ winners, outperforming both the benchmark and the peer-group, persist in performance in time horizons 2- and 3-years post portfolio formation. While this is good news for investment community, several important questions remain unanswered. Specifically, it is not clear whether the winner outperformance is sustained over 24- and 36-month periods tested or it is perhaps front-loaded, concentrated only in early months post-portfolio formation? Is it concentrated in certain subperiods, including known financial markets crisis periods? Does the number of funds in the winner portfolio matter? Do transaction costs erode the MMT benchmark and peer-group adjusted alphas? In this paper, we address these questions and contribute to the literature of mutual fund performance persistence, by providing detailed insight into persistence in the context of benchmark- and peer-group-adjusted alphas. Let us now further discuss the literature underpinning these important questions.

A central concern in recent literature on mutual fund performance and persistence is that mutual fund alpha reported from standard traditional pricing models such as the Fama-French-Carhart (FFC hereafter) framework is model-dependent. FFC models have been shown to generate non-zero alpha estimates when applied to passive indices—benchmarks that are often used to evaluate mutual fund performance (Matallín-Sáez, 2007; Chan et al., 2009; Cremers et al., 2012; Chinthalapati et al., 2017). Since fund managers often structure portfolios in close alignment with their benchmarks, when the benchmark itself underperforms (or outperforms), the FFC models may inadvertently deflate (or inflate) the perceived alpha of a fund. This raises concerns about the accuracy of the models widely accepted as ‘standard’ in the mutual fund

literature. In response, Angelidis et al. (2013) and Chinthalapati et al. (2017) introduced modifications that adjust for benchmark effects—though not for peer group performance. Both studies reported notably improved benchmark-adjusted alpha estimates, when benchmark indices yielded systematically negative alphas under the FFC specification. These findings were corroborated in the UK in Mateus et al. (2016) and imply that standard FFC models may misclassify skilled managers as underperformers during benchmark downturns, and vice versa. However, the benchmark-adjusted models do not account for commonalities arising from the fact that funds within the same category share common investment styles and constraints. The investors are interested in distinctiveness of individual fund alphas and would like to identify funds within the peer group that would perform better than others given the commonalities. Hunter et al. (2014) address this issue by introducing active peer benchmarks in US mutual fund peer-groups: returns constructed from the average performance of funds within the same category. They show that including active peer benchmarks in the FFC model improves performance evaluation by capturing shared active components and reducing residual correlation. This is corroborated by Mateus et al. (2019) in the UK market.

Most recently, Mateus et al. (2023) recognise the performance persistence benefits of selecting mutual fund winners on the basis of both outperforming the fund's self-declared benchmark and the peer-group jointly. This approach is formalised in Mateus et al. (2025), who develop the MMT model which formally augments the FFC specification to account for the fund's benchmark and the fund's peer-group. They show that forming unbiased 'true winner' portfolios based on 36 months rolling benchmark- and peer-group-adjusted MMT alphas, we can achieve persistent performance over the next 36- and 24-month holding period. While the MMT winner fund alphas are significant in the evaluation holding period, it raises the question whether they are period-dependent, concentrated in the shorter horizon that drives this result or whether they will disappear after transaction costs are taken into account.

The length of the holding period critically influences the detection and interpretation of performance persistence in mutual funds. Carhart (1997), using a one-year evaluation and holding period, found that short-term persistence was largely attributable to momentum, which fades beyond a year. Bollen and Busse (2005) show that performance persistence is limited to short, three-month holding windows, but reverses at longer, annual horizons. In contrast, using a bootstrap alpha methodology, Kosowski et al. (2006) identified significant persistence among the top funds using annual holding period, but they have not tested shorter or longer holding

horizons. Most recently, Cuthbertson et al. (2022) document persistence at short horizons, up to 6 months, using standard factor models and practitioners' models such as the seven-factor 'index model' of Cremers et al (2013), which is particularly pronounced amongst portfolios of winners comprising a small, limited number (e.g. 5) of past top performers. The performance diminishes as the holding horizon is extended. Thus, studies based on standard FFC models, not accounting for benchmark or peer-group adjustments consistently show that short holding periods (6–12 months) tend to reveal more persistence, whereas longer horizons (1 year and beyond) often wash out performance signals completely. In contrast, the aforementioned benchmark- and peer-group adjusted model studies, including the MMT model of interest in this paper, all document persistence in performance at 1-year, 2-year or 3 year horizons (e.g. Hunter et al., 2014, Mateus et al. 2019, 2023, 2025) but it remains an open question whether the persistence in this longer horizons is actually driven by early spike in outperformance in shorter horizons following winner portfolio formation.

Further, even though benchmark- and peer-group adjusted models document persistence over longer horizons, it is not clear whether they can identify the persistent winners in the bear and crisis periods. There is some evidence in Mateus et al. (2016), who evaluate benchmark-adjusted persistence in performance, that top quartile funds do not necessarily outperform the bottom quartile funds in each of their 21 rolling sub-periods, across some of the investment styles. However, this paper does not define the winner funds as the MMT model does: as those outbidding the peer-group as well as the benchmark. In this paper, we will formally test the persistence of winners from selected by the MMT over different market conditions.

Finally, A persistent concern is whether any alpha survives implementation. Carhart (1997) provides early evidence that trading and expenses materially erode returns: he reports a sizeable trading-cost drag (approximately 0.95%) and documents underperformance of load funds relative to no-load funds. More recently, Cuthbertson et al. (2022) bring implementability closer to the persistence design by reporting switching intensity. They show that strategies that generate stronger alphas - shorter holding periods, smaller winner sets - imply more switching and therefore potentially higher transaction costs. Benchmark-adjustment performance studies implicitly also underline the importance of costs: Chinthapati et al. (2017) show that tracker funds underperform by approximately the average size of the passive fund expense ratios, indicating that alpha specification improvement does not necessarily remove the economic sting of costs. This notion is yet to be examined in the context of winner fund selection-based

benchmark and peer-group adjusted alphas, as in the MMT model. In this paper, we address the question of whether costs erode the persistence of MMT winner funds, and whether the size of the winner portfolio (as indicated by Cuthbertson et al., 2022) matters.

We contribute to the literature on mutual fund performance persistence by providing further empirical insight into the new Mateus et al. (2025) MMT model that can identify the funds that outperform both their benchmark *and* their peer group simultaneously and persist in performance over 36m and 24m investment horizon. We evaluate whether the MMT model can select unbiased winners whose persistence in performance:

- i) is not limited only to any particular short or long investment horizons;
- ii) is significant regardless of the size of winner portfolios
- iii) holds over various market conditions, including known crisis periods such as the dot.com bubble burst and global financial crisis (GFC hereafter);
- iv) holds after accounting for trading costs

To evaluate the model, we conduct our empirical analysis on a sample of 735 US Large Cap Growth mutual funds, as classified by Morningstar. The period of analysis is from January 1990 to June 2018. Using the MMT model, which adjusts for both benchmark- and peer-group effects, we identify “true winner” funds and track their performance across 6-, 12-, and 24-month holding periods. The results reveal persistent outperformance of winners relative to both benchmarks and peers, with the highest returns observed under shorter rebalancing frequencies. Crucially, this persistence remains evident across different market regimes—including the dot-com bubble and the Global Financial Crisis—and survives realistic estimates of transaction costs, confirming its practical implementability. Notably, the superior returns are not driven solely by a few standout funds. Even diversified equally weighted portfolios comprising all identified winners—often more than 60 funds—exhibit substantial excess returns, reducing the need for high turnover or narrow fund selection. The MMT model’s strength lies in its ability to detect enduring fund manager skill by capturing peer-relative sensitivity and benchmark-relative style exposures, such as small-cap and growth tilts, that are missed by traditional factor-based models. We show that loser funds diverge more from the characteristics of the peer-group compared to winners, by overweighting small cap and value stocks, which if not timed properly may have led to their inferior performance.

The rest of the paper is organised as follows. Section 2 presents the data and our new methodology. Section 3 lays out the results and Section 4 provides conclusions.

## 2. Data and Methodology

### 2.1.Data

Our data set comprises 785 active US long-only equity mutual funds in the Large Cap Growth category, selected as the largest peer-group of US equity mutual funds reported in Morningstar. The number is obtained after screening out tracker funds and imposing a minimum requirement of 36 months of returns for each fund to conform to the research design. The dataset used in this study mitigates survivorship bias by including both active and defunct mutual funds. Consistent with standard practice in mutual fund research (e.g., Cuthbertson et al., 2022, 2023; Clare et al., 2021), the analysis is based on each fund's oldest available share class. Data spans from January 1990 to June 2018. The total monthly fund returns net of fees, funds' Morningstar style category, funds' self-reported primary prospectus benchmark, the total returns of primary prospectus benchmarks, the total returns of the representative benchmark for the Large Cap Growth category (Russell 1000 Growth) are all from Morningstar. Fama-French-Carhart factors used in the MMT model are from Kenneth French's data library<sup>1</sup>.

The fund characteristics in Table 1 indicate that there are 34 different primary prospectus benchmarks used by 735 funds in the Large Cap Growth category, that an average fund in the category spends nearly 15 years in our sample, it has a net expense ratio of just over 1% and average manager tenure of around 10 years.

**Table 1: Description of fund characteristics – US Large Cap Growth**

The table reports the number of unique funds, average number of months in the sample, average fund size (in million of USD), average net expense ratio and average manager tenure (in years) for the Large Cap Value and Large Cap growth peer group.

|                                                                | <b>Large Cap Growth</b> |
|----------------------------------------------------------------|-------------------------|
| Number of unique funds                                         | 785                     |
| Average number of months in sample                             | 178                     |
| Number of self-reported primary benchmarks within the category | 34                      |
| Fund Size (in million USD, average)                            | 7,234                   |
| Net Expense Ratio (Average)                                    | 1.02%                   |
| Manager Tenure (average, in years)                             | 10.03                   |

<sup>1</sup> [http://mba.tuck.dartmouth.edu/pages/faculty/ken.french/data\\_library.html](http://mba.tuck.dartmouth.edu/pages/faculty/ken.french/data_library.html)

## 2.2. The Mateus, Mateus, Todorovic (2025) - MMT Model

To construct unbiased winner portfolios, this study adopts the methodology proposed by Mateus, Mateus, and Todorovic (2025), the MMT model, which introduces a performance evaluation framework that reflects the dual criteria commonly employed in professional fund assessment: outperformance relative to a fund's self-declared benchmark index and relative to its peer group. While the Carhart (1997) four-factor model remains a widely accepted standard tool for estimating mutual fund alpha, its application is in absolute terms and it fails to consider these two essential dimensions of comparative, relative performance. As a result, standard models may misclassify funds—labelling as underperformers those that outperform both their benchmark and peers during periods of adverse market or benchmark conditions.

The MMT approach involves two key steps. First, an Active Peer Benchmark (APB) is constructed by calculating the equally weighted return of all funds within a given peer group. It is then assessed whether the APB generates excess returns relative to the general characteristics of investment style it is meant to represent. Specifically, the performance commonalities among peer funds that are not captured by the benchmark index representing the peer group<sup>2</sup> are isolated. To do this, alpha of the APB is estimated using an augmented Carhart model, modifying the standard specification by replacing the risk-free rate with the benchmark return on the left-hand side of the regression<sup>3</sup>, as in benchmark-adjusted model of Angelides et al. (2013).

The final step is the estimation of a benchmark- and peer-adjusted performance model. This augmented specification modifies both sides of the Carhart equation to account for relative performance. Specifically, the dependent variable becomes the fund's return in excess of its self-reported benchmark, following Angelidis et al. (2013), while the impact of benchmark-adjusted APB return is included as the additional 'peer-group' factor ( $\alpha_{APB,b} + e_{APB,b,t}$ ), to capture peer-group dynamics beyond those explained by the peer-group's benchmark. The resulting MMT model is given by:

---

<sup>2</sup> Russell 1000 Growth in this case

<sup>3</sup> As follows:

$$R_{APB,t} - R_{b,t} = \alpha_{APB,b} + \beta_{APB,b,M}(R_{M,t} - R_{f,t}) + \beta_{APB,b,SMB}SMB_t + \beta_{APB,b,HML}HML_t + \beta_{APB,b,WML}WML_t + e_{APB,b,t}$$

Where  $R_{APB,t}$  is the return of APB portfolio of the peer group,  $R_{b,t}$  is the return of the representative benchmark for our peer-group, Russell Russell 1000 Growth;  $\alpha_{APB,b}$  is the benchmark-adjusted alpha of the peer-group and  $e_{APB,b,t}$  is the difference of the idiosyncratic risk of the peer-group and of the benchmark index with respect to the common four factors in the model. If factor loadings in this model differ from zero - the APB portfolio of large cap value peer group has the different risk characteristics from those of the corresponding Russell 1000 Value benchmark.

$$R_{i,t} - R_{bi,t} = \alpha_{i,MMT} + \beta_{ib,M}(R_{M,t} - R_{f,t}) + \beta_{ib,SMB}SMB_t + \beta_{ib,HML}HML_t + \beta_{ib,WML}WML_t + \beta_{ib,APBp}(\alpha_{APB,b} + e_{APBp,b,t}) + \omega_{i,MMT,t} \quad (1)$$

where  $R_{i,t}$  denotes the return on fund  $i$ , and  $R_{bi,t}$  represents the return on the fund's self-designated benchmark.  $(R_{M,t} - R_{f,t})$  is the excess return on the aggregate US equity market over the one-month Treasury bill rate. The US equity market is a value-weighted portfolio of all firms listed on the NYSE, AMEX, and NASDAQ. This market portfolio proxy encompasses a broader investment universe than any benchmark index employed in this study; consequently, each fund's benchmark index constitutes a subset of the market portfolio. Consistent with Fama and French (1993), the SMB factor captures size-related risk and is constructed as the return differential between small-capitalization and large-capitalization stock portfolios, while the HML measures value risk and is defined as the return spread between portfolios with high and low book-to-market ratios. Following Carhart (1997), WML represents the momentum factor, calculated as the difference in returns between portfolios of prior-period winner and loser stocks. All factor definitions adhere to standard conventions in the literature.

The term  $(\alpha_{APB,b} + e_{APB,b,t})$  is introduced as additional, 'peer-group' factor to the standard Carhart four-factors to capture performance and risk commonalities within a fund's peer group, that are not reflected in the peer-group's benchmark, the Russell 1000 Growth Index. The coefficient  $\alpha_{i,MMT}$  is the benchmark- and peer-group augmented MMT alpha of fund  $i$ ; while  $\omega_{i,MMT,t}$  reflects deviations between the fund's idiosyncratic risk exposures and those implied by its prospectus benchmark given the five factors in this model.  $\beta_{ib,APB}$  measures the fund's sensitivity to the 'peer-group' factor, adjusted for benchmark effects.  $\beta_{ib,M}, \beta_{ib,SMB}, \beta_{ib,HML}, \beta_{ib,WML}$  denote the fund's excess loadings on the market, size, value, and momentum factors, respectively, relative to its self-reported benchmark.

This augmented model produces a more accurate estimate of fund manager skill by correcting for structural biases identified in traditional factor models. As demonstrated by Mateus et al. (2025), it enables more reliable identification of outperforming funds—those that generate excess returns relative to both industry benchmarks and peers—and provides a more effective basis for studying performance persistence.



### 3. Persistence in Performance: Empirical Design and Results

#### 3.1 Persistence on a monthly basis, month 1 to month 36

We deploy the MMT model to evaluate performance of our funds by computing the MMT alphas using rolling windows approach, consistent with Mateus et al. (2005). Specifically, the model is estimated monthly using a 36-month rolling window, starting with the initial period from January 1990 to December 1992. For each window, mutual funds are ranked based on the t-statistics of their estimated MMT alphas. So, the first MMT alpha is estimated using data from January 1990 to December 1992, the second one using data from February 1990 to January 1993 and so on.

Funds with MMT alpha t-statistics exceeding +1.65 or falling below -1.65 are designated as ‘winners’ or ‘losers’, respectively, based on statistical significance at the 10% level. For instance, funds identified as winners in the window ending December 1992 enter the post-ranking period starting January 1993, which is considered ‘month 1’ of their holding period. This labelling convention is consistently applied across the sample, such that if a regression spans July 1994 to June 1997, the corresponding ‘month 1’ for the post-ranking performance is July 1997. This allows for the construction of a panel of month-specific returns: ‘month 1’ (first month post-ranking), ‘month 2’ (second month post-ranking), all through to ‘month 36’ (three years post-ranking) for each cohort of historical ‘winners’ and ‘losers’ established at the end of each rolling window. The purpose of this approach is to evaluate whether average abnormal returns vary systematically across individual post-ranking months. In total, there are 300 rolling windows and therefore 300 ‘month 1’, ‘month 2’, ‘month 3’ returns and so on.

We also compute the excess returns for each winner/loser fund in each month post portfolio formation as:

$$R_{Wi/Li,t} - (\max(R_{bi,t}, R_{APB,t})) \quad (2)$$

Where  $R_{Wi/Li,t}$  is the return of winner fund (W) or loser fund (L)  $i$  in month  $t$  ( $t = 1$  to 36) post ranking,  $R_{bi,t}$  is the return on the fund’s self-designated benchmark and  $R_{APB,t}$  is the return of the active peer group portfolio in the same month  $t$ . Excess return is calculated for each month separately, from ‘month 1’ to ‘month 36’ relative to the higher of the two benchmarks: fund’s self-reported benchmark or peer-group. The results are presented in Table 2 and Figure 1. We average individual monthly excess returns for each ‘winner’ and ‘loser’ fund over semi-annual

periods: 1-6 months, 7-12 months etc. The results show that on average, ‘winners’ provide a consistent return above the peer group and benchmark of around 45 basis points per month in the three years following winner/loser fund classification. In contrast, losers on average underperform by approximately 30 basis points in each month. The winner/loser spread is persistently in excess of 65bps per month. From the investor’s perspective, these results are excellent news, indicating that if the MMT model classifies a fund as a ‘winner’ - the timing of investment within 36 months following fund classification does not dictate the outcome – consistent abnormal returns vis-a-vie the benchmark(s) are achieved in each month throughout the period for the next 3 years.

**Table 2: Excess returns of winner and loser funds**

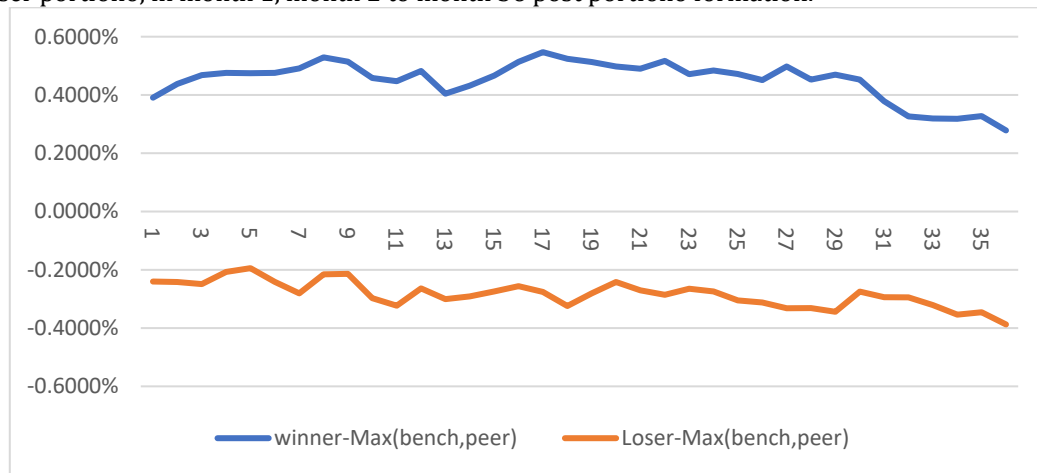
Returns are in basis points (bps) per month, where 100 bps = 1%. ‘Winners’ are defined as funds with statistically significant positive alpha (t-stat > 1.65) over a t-36-month estimation window using MMT model; ‘losers’ are those with significantly negative alpha (t-stat < -1.65). Excess return is calculated for each month, from Month 1 to Month 36 relative to each fund's self-reported benchmark and/or peer-group. They are then averaged for each semi-annual period: 1-6 months, 7-12 months etc. The “Spread” column shows the return differential between winner and loser portfolios.

| <b>Post-Selection Period (Months)</b> | <b>Winners (bps/mth)</b> | <b>Losers (bps/mth)</b> | <b>Spread (bps)</b> |
|---------------------------------------|--------------------------|-------------------------|---------------------|
| 1–6                                   | 45                       | -23                     | <b>68</b>           |
| 7–12                                  | 48                       | -27                     | <b>75</b>           |
| 13–18                                 | 48                       | -28                     | <b>76</b>           |
| 19–24                                 | 50                       | -27                     | <b>77</b>           |
| 25–30                                 | 47                       | -32                     | <b>79</b>           |
| 31–36                                 | 33                       | -33                     | <b>66</b>           |

The results in Figure 1 illustrate that excess returns relative to both the benchmark and peer group are always positive statistically significant in the first month following selection, with patterns persisting across subsequent months.

### Figure 1: Excess returns of winner and loser funds relative to benchmark/peer-group

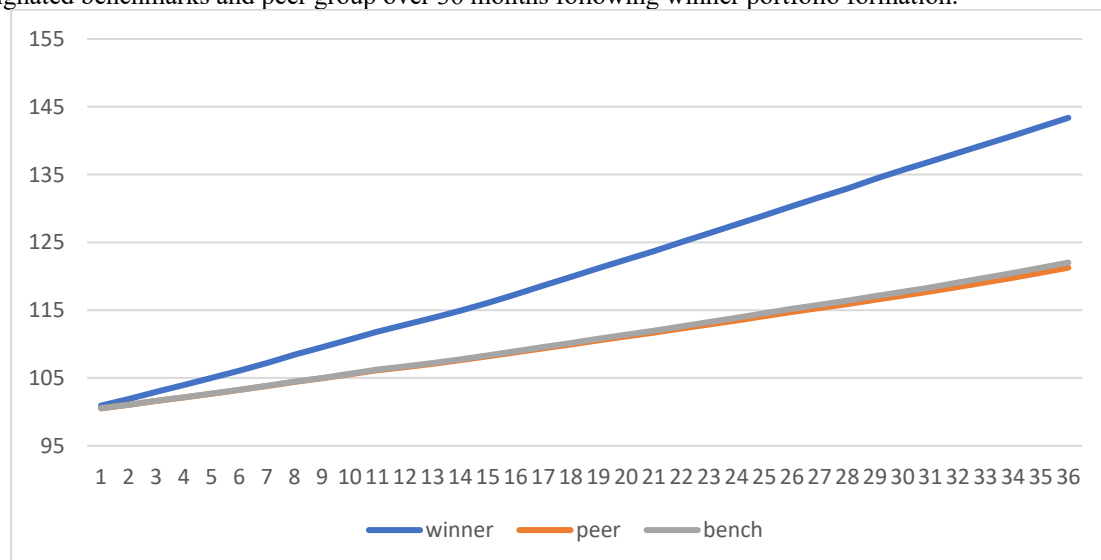
The figure shows average monthly returns in excess of  $\max(R_{bi,t}, R_{APB,t})$  across all the funds in the winner and loser portfolio, in month 1, month 2 to month 36 post portfolio formation.



Further, we examine the cumulative return in excess of the risk-free rate of winner funds compared to the peer group and self-reported benchmark(s). We invest 100 USD in ‘winner’ funds in each rolling window and on average accumulate to 143 USD by the end of 36 months, well above the peer-group and benchmark (benchmarks accumulate slightly more than the peer-group, USD122). This is presented in Figure 2. The compound annual growth rate (CAGR) corresponding the values in Figure 2 is 12.76% p.a. on average for winner funds, 6.86% p.a. for the funds’ self-designated benchmarks and 6.64% p.a. for the peer-group. The benchmark and the peer-group CAGR are in line with long term estimates of equity risk premium in developed markets, including US.

### Figure 2: Cumulative excess returns of winners, benchmarks and peer-group

The figure shows the average cumulative returns in excess of risk-free rate of winner portfolios, funds’ self-designated benchmarks and peer group over 36 months following winner portfolio formation.

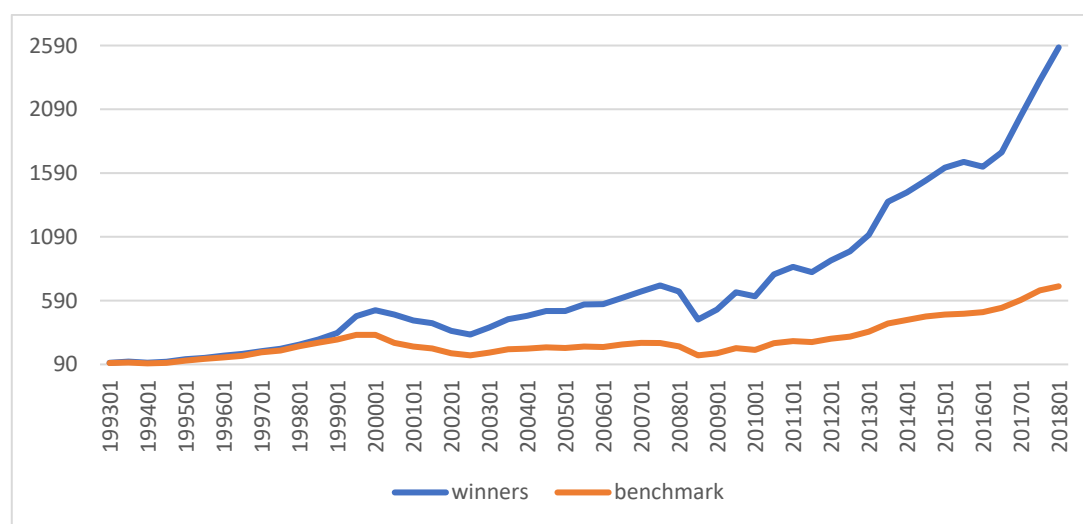


## 3.2. Does the length of the holding period play a role in capturing persistence of winners?

### 3.2.1. 6 months holding period

To assess performance of winners over six month holding period horizons, we start by creating create equally weighted portfolios of *all* winner funds. Winners are defined as described Section 3.1. We rebalance this ‘All winners’ portfolio every six months. We then compute cumulative compounded investment value over the sample period as follows: we invest 100 USD in January 1993 in the first equally weighted portfolio of ‘All winner’ funds (based on MMT alpha t-statistics from January 1990 – December 1992) and compound the investment until June 1993. We then rebalance the winner portfolio in July 1993 based on the new set of MMT alpha t-states from the 36 months preceding that month and continue compounding for the next six months, until December 1993. We repeat the process every six months until we reach the end of the sample period and generate series of compounded returns of equally weighted portfolio of all winner funds. These are presented in Figure 3 and compared to cumulative returns of equally weighed portfolio of funds’ self-reported benchmarks<sup>4</sup>. The figure shows that the that portfolio of winners generates close to 2590 USD while the self-designated benchmark portfolio accumulates just over 700 USD at the end of our sample.

**Figure 3: Cumulative compounded returns of ALL winners’ portfolio vs benchmark, 6 months rebalancing**



<sup>4</sup> Note that cumulative performance of the benchmarks, Russell 1000 Growth and the peer-group portfolio are very much aligned, so we report cumulative performance relative to the benchmarks – if the winners have outperformed the benchmarks, they have outperformed the peer-group and the Russell 1000 Growth index.

Next, we zoom into the top five winner funds only and repeat construction of equally weighted winner portfolio but this time of top 5 funds only, and the computation of cumulative compounded returns when the portfolio is rebalanced every 6 months. The results are presented in Figure 4.

**Figure 4: Cumulative compounded returns of Top 5 winners portfolio vs benchmarks, 6 months rebalancing**

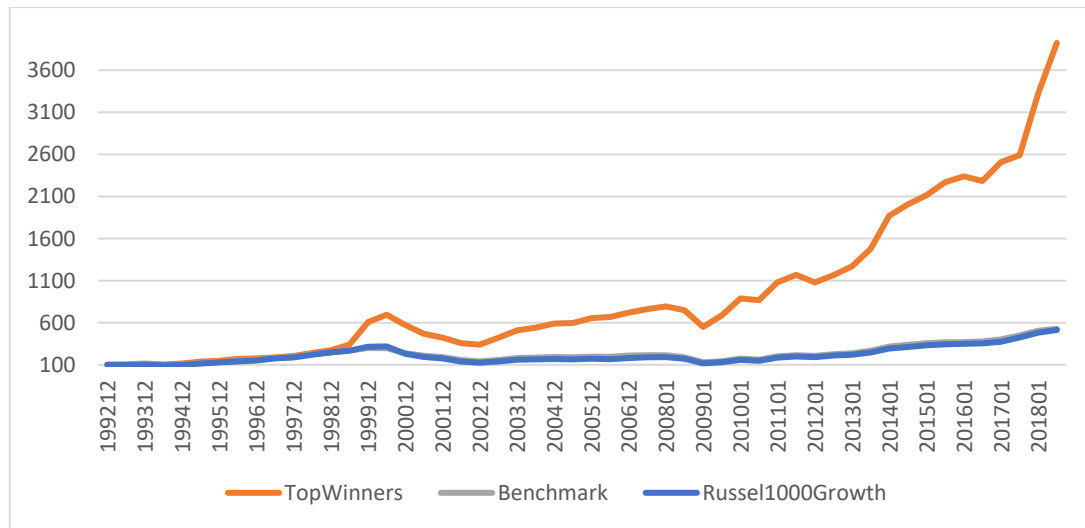


Figure 4 shows that investing 100 USD in the winner portfolio comprised of top 5 winners only, rebalanced every 6 months, will generate nearly 4000 USD at the end of horizon, while their benchmarks accumulate only just over 500 USD. We find the average monthly excess return relative to the benchmarks/peer-group for top 5 winners using equation (2) to be 0.43% (t-stat = 2.87%, significant at 1% level). On the annualised basis, top 5 winners have on average 5.24% higher return than the self-declared benchmark/peer group.

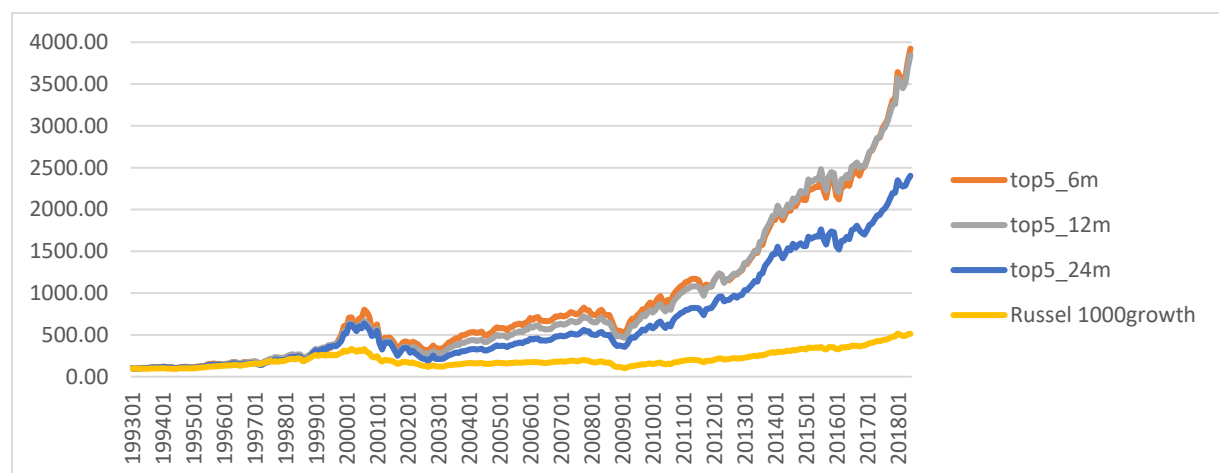
Figure 3 and Figure 4 also documents that the economic value of cumulative performance of ‘All winner’ portfolio for investor is strong, so one does not have to focus on a small number of top funds only to outperform the benchmark(s). Focusing on small number of funds in a portfolio can lead to more rebalancing and greater transaction costs (Cuthebertson et al., 2022).

### 3.2.2. 12 months and 24 months holding period and winner portfolios of 5, 10 and all winners

Previous studies on performance persistence reviewed in the introduction to this paper, where winners were picked based on standard FFC or industry models such as Cremers et al. (2013), suggest that there is more evidence of persistence of winners over shorter horizons of up to one year. To test this within the MMT framework, we extend our analysis from Section 3.1.1. and introduce winner portfolio rebalancing at 12 and 24 months. Using top 5 winner funds for the equally weighted winner portfolio, 12 and 24 month holding period, we compute compounded cumulative returns as in Section 3.1.1., relative to Russell 1000 Growth Index - recall footnote 4, Russell 1000 Growth generates very similar compounded cumulative returns as the benchmark and the peer group, so the results with the other two benchmarks would be qualitatively and quantitatively comparable.

Figure 5 shows the results, including the 6-month holding period rebalancing for comparison. It is evident that compounded cumulative returns from 6 months and 12-month horizon rebalancing are very much aligned while 24-month rebalancing still indicates good, albeit weaker persistence in performance than over shorter holding periods. Specifically, investing 100 USD in January 1993 in equally weighted portfolio of top 5 winners will earn just under 4000 USD at the end of our sample if rebalancing every 6 or every 12 months, while 24-monthly rebalancing will almost half the end of period value to just over 2400 USD. Nevertheless, this is still nearly 5 times more than one would earn if invested in Russell 1000 Growth index over the same period (just over 500 USD).

**Figure 5: Cumulative compounded returns of Top 5 winners portfolio vs benchmarks, 6-, 12- and 24-month rebalancing**



We extend the analysis to the portfolio of top 10 winner funds and all winner funds, and plot Figure 6, where we observe similar pattern as in Figure 5. For all portfolio sizes, top 5, top 10 or all winners (which often have in excess of 60 funds), the results show the highest cumulative performance over 6-month horizon. Unsurprisingly, the portfolio of top 5 winners generates the highest cumulative value at the end of the sample. This is very much in line with prior studies on performance persistence. However, many studies based on standard FFC models find that persistence wanes at 12 months holding periods or longer. This is not the case when winners are defined by the MMT model. The persistence weakens over longer horizons but the minimum cumulative performance recorded for all winners and 24 months holding period shows almost three times higher end of period value of 100 USD investment (in excess of 1400 USD) compared to that of Russell 1000 Growth index (just over 500 USD). This implies that MMT model is capable of identifying winners whose performance can persist over shorter and longer horizons, even though it is stronger over shorter, 6-month horizons.

**Figure 6: Cumulative compounded returns of Top 5, Top 10 and All winners portfolio vs benchmarks, 6-, 12- and 24-month rebalancing**

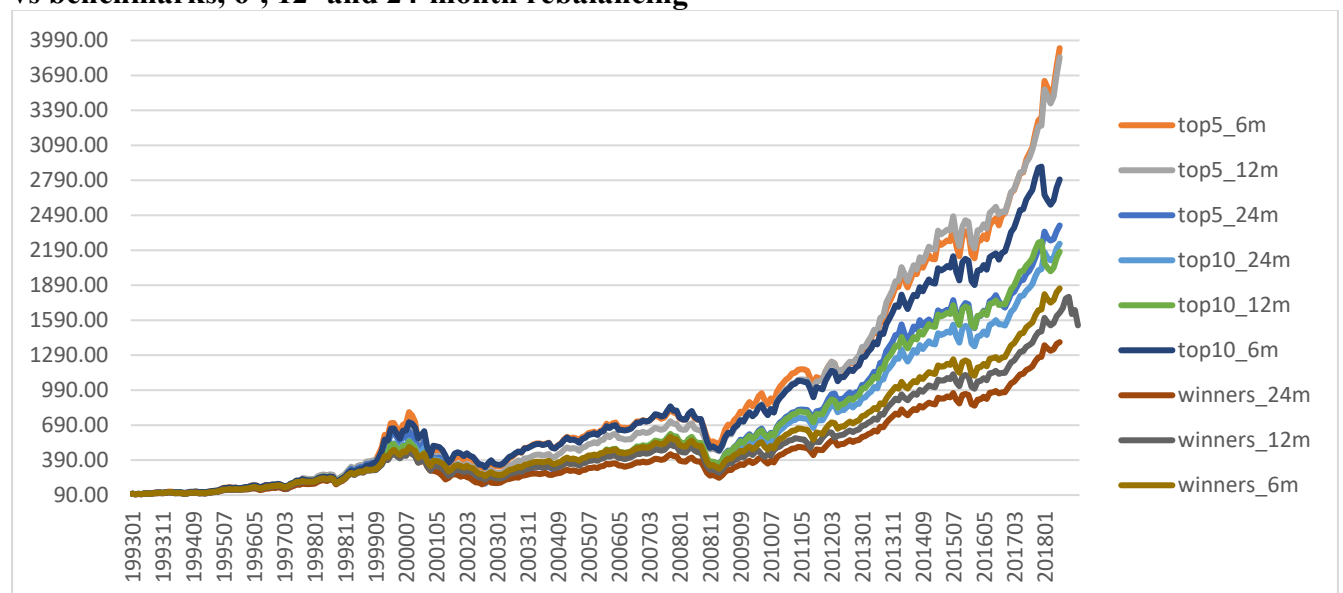


Table 3 below more formally sets out the performance of top 5, top 10 and all winners portfolios over the evaluation periods of 6, 12 and 24 months relative to self-declared benchmarks and the Russell 1000 Index. The table presents averages of monthly and annualised mean returns, standard deviations and Sharpe ratios. Corroborating results from Figure 5 and Figure 6, all winner portfolios across all holding periods generate substantially higher returns than the benchmarks, but also higher standard deviations. Nevertheless, all Sharpe ratios associated with winner portfolios across various holding periods are at between 41% and 64% higher than

average Sharpe ratio of self-declared benchmarks. Higher Sharpe ratios are the highest over 6- and 12-month horizons for all portfolios, similar to what cumulative performance was suggesting in Figure 6. Sharpe ratios over 24-month horizons are only marginally weaker, confirming the strong ability of the MMT model to select persistent winners.

**Table 3: Performance indicators for Top 5, Top 10 and All Winners funds, over 6-, 12- and 24- month holding periods**

The table shows monthly return and standard deviation, annual return and standard deviation and the Sharpe ratio for Top 5, Top 10 and All winners portfolios rebalanced over 6-, 12- and 24-month horizon, their self-declared benchmarks and Russell 1000 Growth index. \*\*\*, \*\*, \* denotes significance at 1%, 5% and 10% (t-test).

|                          | Monthly<br>return | St. dev.<br>monthly | Return p.a. | St.Dev. p.a. | Sharpe ratio |
|--------------------------|-------------------|---------------------|-------------|--------------|--------------|
| top5_6m                  | 1.39%***          | 6.01%               | 18.00%***   | 20.82%       | 0.86         |
| top5_12m                 | 1.37%***          | 5.83%               | 17.76%***   | 20.19%       | 0.88         |
| top5_24m                 | 1.24%***          | 6.22%               | 15.95%***   | 21.54%       | 0.74         |
| top10_6m                 | 1.24%***          | 5.47%               | 16.00%***   | 18.94%       | 0.84         |
| top10_12m                | 1.16%***          | 5.35%               | 14.79%***   | 18.53%       | 0.80         |
| top10_24m                | 1.18%***          | 5.51%               | 15.05%***   | 19.08%       | 0.79         |
| All winners_6m           | 1.08%***          | 4.79%               | 13.70%***   | 16.60%       | 0.83         |
| All winners_12m          | 1.04%***          | 4.74%               | 13.26%***   | 16.43%       | 0.81         |
| All winners_24m          | 1.02%***          | 4.89%               | 12.99%***   | 16.94%       | 0.77         |
| Self-declared benchmarks | 0.64%**           | 4.35%               | 7.97%**     | 15.07%       | 0.53         |
| Russel 1000 Growth       | 0.65%**           | 4.67%               | 8.05%**     | 16.17%       | 0.50         |

#### 4. Persistence of winners from the MMT model across various market conditions

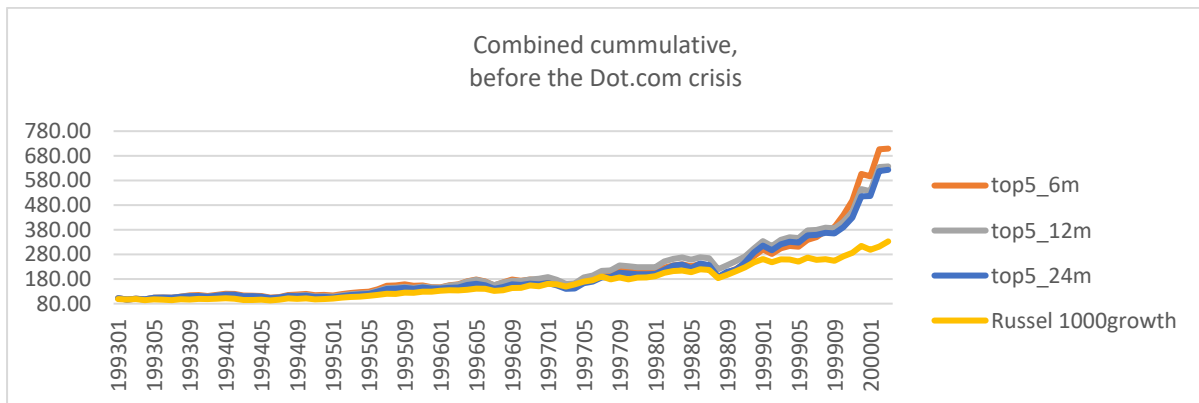
To further investigate the time variation in fund performance persistence, we partition the sample into distinct sub-periods based on major market events. This segmentation allows for a clearer interpretation of return dynamics under different economic conditions. First, we define the pre-dot-com period as running from the beginning of the sample (January 1990) up to March 2000. The dot-com crisis period itself spans from April 2000 to September 2002, reflecting the sharp market decline and subsequent volatility. The post-dot-com, pre-Global Financial Crisis (GFC) period covers the relative recovery and expansion phase from October 2002 to October 2007. The GFC is captured between November 2007 and March 2009, a timeframe marked by severe market downturn. Finally, the post-GFC period extends from April 2009 through the end of our sample in June 2018 encompassing the long post-crisis recovery and expansion. This sub period-based breakdown enables us to examine whether performance persistence patterns differ meaningfully across market regimes.



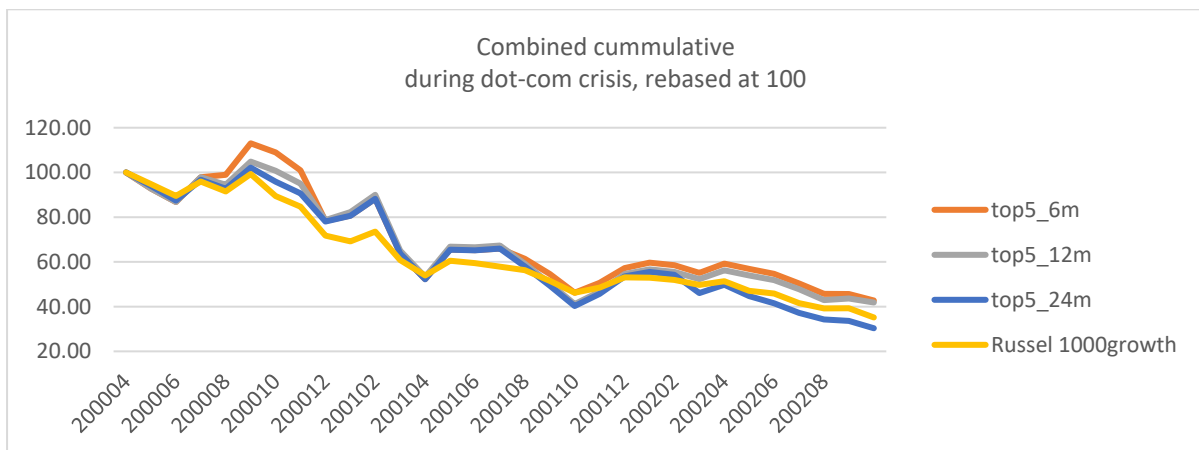
Figure 7 follows the same methodology for computing cumulative compounded returns of a portfolio of top 5 winner funds as used in Figure 5. All parts of Figure 7 are rebased to 100 and show superior cumulative compounded returns corresponding to Top 5 winner portfolios relative the Russell 1000 Growth Index. An exception to the overall outperformance pattern appears in panel (b) of Figure 7, which covers the dot-com bubble burst period. During this crisis, the top 5 winner portfolio, rebalanced every 24 months, underperforms the Russell 1000 Growth Index. This underperformance is not unexpected, given that growth stocks—which were disproportionately affected during the dot-com crash—form the core of the large-cap growth fund category from which these top-performing funds are selected. The extended rebalancing horizon implies that the portfolio continued to hold these previously high-performing funds as their underlying growth-oriented holdings experienced sharp declines. As a result, the adverse impact of the dot-com crisis on growth stocks contributed to the observed drop in cumulative compounded returns for the top 5 winners portfolio rebalanced every 24 months during this period. Also unsurprisingly, there is a decline in cumulative compounded returns during of winner funds the two crisis periods in part (b) and (d) of the figure, but they decline less than those of the benchmark. In the more stable market environments depicted in panels (a), (c), and (e) of Figure 7, there is no consistent pattern indicating a single optimal holding period. For instance, the 6-month rebalancing strategy performs best in the early 1990s and during the dot-com expansion phase (panel a), while the 24-month horizon dominates in the period between the dot-com crisis and the Global Financial Crisis (panel c), and the 12-month rebalancing strategy outperforms in the post-GFC era (panel e). Despite this variation, the results collectively demonstrate robust evidence of performance persistence across different market regimes. Notably, all three holding period strategies yield cumulative returns that exceed Russell 1000 Growth benchmark, reinforcing the notion that past winners from the MMT model continue to generate abnormal returns.

**Figure 7: Cumulative compounded returns of Top 5 winners portfolio over different market conditions, 6-, 12- and 24-month rebalancing**

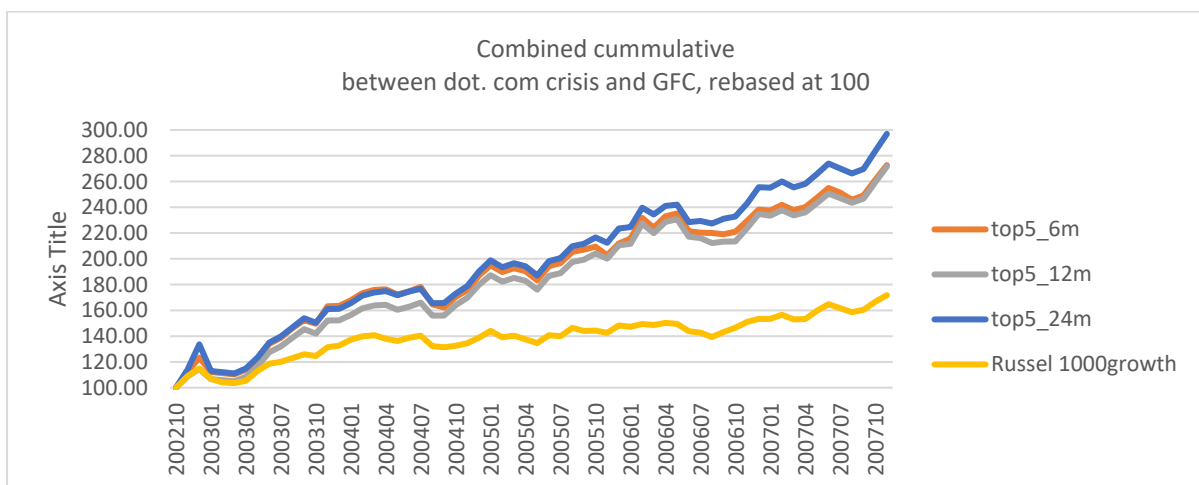
**Part (a) Period prior to the dot.com crisis**



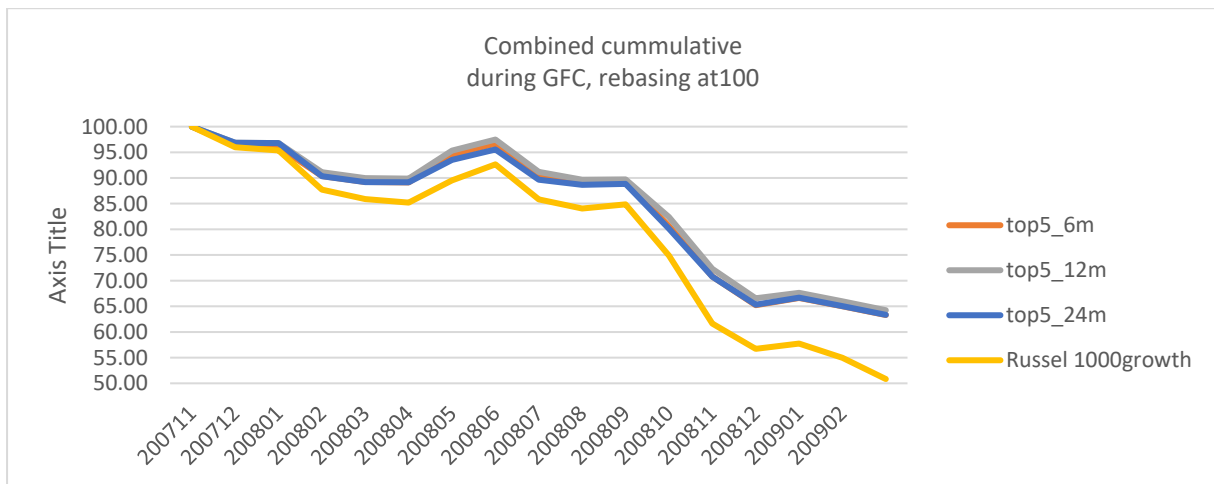
**Part (b): Dot.com crisis**



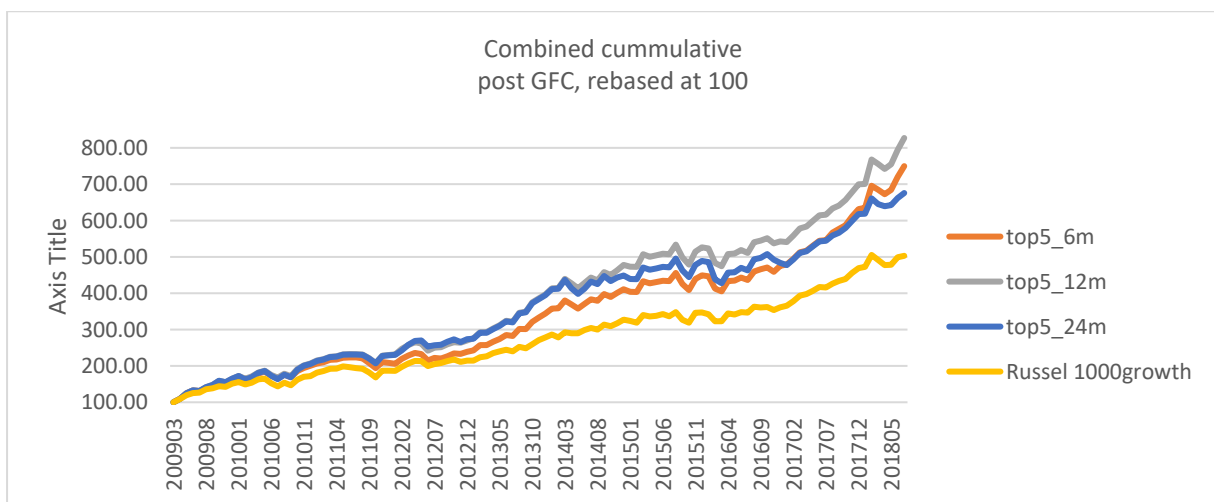
**Part (c): Period between dot.com crisis and global financial crisis (GFC)**



#### Part (d): Global Financial Crisis (GFC)



#### Part (e) Post GFC period



#### 5. Is there persistence after trading costs? Whole sample, bull and crises periods.

To assess the practical relevance of performance persistence, it is essential to evaluate whether abnormal returns survive once trading costs are considered. This section examines whether the return differentials between past winners and losers remain economically meaningful after accounting for realistic transaction costs, offering a more accurate reflection of net performance available to investors.

Our fund returns are already given net of fees, but our winner portfolios are rebalanced at 6-, 12- and 24-month periods. To assess if such strategies of buying and rebalancing winner portfolios are feasible and do not erode performance, we compute break-even Sharpe ratios

showing how much the return of the portfolio should be reduced to equalise Sharpe ratio of the winner portfolio to that of the Russell 1000 Value Index and funds' self-designated benchmarks.

Table 4 reports performance indicators such as annual returns, standard deviations, Sharpe ratios and break-even Sharpe ratios for top 5, top 10 and all winners' portfolios over 6-, 12- and 24-month holding periods, across the whole sample period and subperiods analysed in Section 4. Note that when annual returns and Sharpe ratios are negative during the crisis periods, the break-even Sharpe ratios are not computed. Busse et al. (2017) document that average annual transaction costs per year of actively managed US equity mutual funds is around 0.75%. In a more recent study Busse et al. (2021) state that trading cost impact large cap funds much less and cause a performance drag of around 0.40% per annum. Our fund returns are net of fees. All break-even transaction costs reported in Table 4 are well above values reported in literature, signalling that trading the winner portfolios constructed using the MMT model which will not wipe out the outperformance.

**Table 4: Break even transaction costs for Top 5, Top 10 and All Winners portfolios, over 6-, 12-, and 24-month holding periods, across sub-periods**

The table shows annual return and standard deviation, the Sharpe ratio and break-even Sharpe ratio for Top 5, Top 10 and All winners portfolios rebalanced over 6-, 12- and 12-month horizon, their self-declared benchmarks and Russell 1000 Growth index, across the market regimes: pre-dot.com crisis, dot.com crisis, dot.com crisis to GFC, GFC period and post GFC period. \*\*\*, \*\*, \* denotes significance at 1%, 5% and 10% (t-test).

| Period                                             | Return p.a. | St. Dev. p.a. | Sharpe ratio | Break-even Transaction Costs |
|----------------------------------------------------|-------------|---------------|--------------|------------------------------|
| <b>Whole Sample</b>                                |             |               |              |                              |
| Top 5 Winner-6m                                    | 18.00%***   | 20.82%        | 0.8648       | 7.64%                        |
| Top 5 Winner-12m                                   | 17.76%***   | 20.19%        | 0.8795       | 7.71%                        |
| Top 5 Winner-24m                                   | 15.95%***   | 21.61%        | 0.7382       | 5.19%                        |
| Top 10 Winner-6m                                   | 16.00%***   | 18.94%        | 0.8445       | 6.57%                        |
| Top 10 Winner-12m                                  | 14.79%***   | 18.51%        | 0.7341       | 5.57%                        |
| Top 10 Winner-24m                                  | 15.05%***   | 19.09%        | 0.7884       | 5.55%                        |
| All Winner-6m                                      | 13.26%***   | 16.43%        | 0.8075       | 5.08%                        |
| All Winner-12m                                     | 13.70%***   | 16.60%        | 0.8254       | 5.44%                        |
| All Winner-24m                                     | 12.99%***   | 16.94%        | 0.7667       | 4.56%                        |
| <i>Self-declared Benchmark</i>                     | 7.97%**     | 15.07%        | 0.5288       | ---                          |
| <i>Russel 1000 Growth</i>                          | 8.05%**     | 16.17%        | 0.4978       | ---                          |
| <b>Dot.com crisis, April 2000 – September 2002</b> |             |               |              |                              |
| Top 5 Winner-6m                                    | -22.98%     | 39.24%        | ---          | ---                          |
| Top 5 Winner-12m                                   | -23.35%     | 40.31%        | ---          | ---                          |
| Top 5 Winner-24m                                   | -32.45%     | 40.74%        | ---          | ---                          |
| Top 10 Winner-6m                                   | -20.47%     | 32.09%        | ---          | ---                          |
| Top 10 Winner-12m                                  | -24.10%     | 33.92%        | ---          | ---                          |

|                                                                   |           |        |        |       |
|-------------------------------------------------------------------|-----------|--------|--------|-------|
| Top 10 Winner-24m                                                 | -26.73%   | 32.73% | ---    | ---   |
| Winner-6m                                                         | -21.80%   | 24.94% | ---    | ---   |
| Winner-12m                                                        | -20.48%   | 24.11% | ---    | ---   |
| Winner-24m                                                        | -24.79%*  | 25.15% | ---    | ---   |
| <i>Self-declared Benchmark</i>                                    | -27.61%** | 21.23% | ---    | ---   |
| <i>Russel 1000 Growth</i>                                         | -31.98%** | 25.30% | ---    | ---   |
| <b>After Dotcom Bubble until GFC October 2002 to October 2007</b> |           |        |        |       |
| Top 5 Winner-6m                                                   | 22.96%*** | 13.96% | 1.6449 | 6.98% |
| Top 5 Winner-12m                                                  | 22.63%*** | 12.45% | 1.8171 | 8.38% |
| Top 5 Winner-24m                                                  | 25.41%*** | 16.04% | 1.5846 | 7.06% |
| Top 10 Winner-6m                                                  | 21.59%*** | 12.77% | 1.6901 | 6.97% |
| Top 10 Winner-12m                                                 | 22.28%*** | 12.05% | 1.8491 | 8.49% |
| Top 10 Winner-24m                                                 | 22.69%*** | 14.11% | 1.6085 | 6.54% |
| Winner-6m                                                         | 19.14%*** | 11.05% | 1.7331 | 6.50% |
| Winner-12m                                                        | 19.10%*** | 11.30% | 1.6910 | 6.17% |
| Winner-24m                                                        | 19.50%*** | 12.10% | 1.6122 | 5.66% |
| <i>Self-declared Benchmark</i>                                    | 11.73%**  | 10.05% | 1.1676 | ---   |
| <i>Russel 1000 Growth</i>                                         | 11.80%**  | 10.31% | 1.1446 | ---   |
| <b>GFC November 2007 to February 2009</b>                         |           |        |        |       |
| Top 5 Winner-6m                                                   | -28.05%** | 16.51% | ---    | ---   |
| Top 5 Winner-12m                                                  | -27.32%** | 15.93% | ---    | ---   |
| Top 5 Winner-24m                                                  | -28.09%** | 15.90% | ---    | ---   |
| Top 10 Winner-6m                                                  | -34.31%** | 20.18% | ---    | ---   |
| Top 10 Winner-12m                                                 | -32.89%** | 18.97% | ---    | ---   |
| Top 10 Winner-24m                                                 | -33.59%** | 19.55% | ---    | ---   |
| Winner-6m                                                         | -36.44%** | 20.31% | ---    | ---   |
| Winner-12m                                                        | -36.31%** | 20.59% | ---    | ---   |
| Winner-24m                                                        | -35.85%** | 19.89% | ---    | ---   |
| <i>Self-declared Benchmark</i>                                    | -39.64%** | 20.11% | ---    | ---   |
| <i>Russel 1000 Growth</i>                                         | -38.45%** | 20.57% | ---    | ---   |
| <b>After Financial Crisis (GFC) March 2009 to June 2018</b>       |           |        |        |       |
| Top 5 Winner-6m                                                   | 25.19%*** | 13.59% | 1.8539 | 4.68% |
| Top 5 Winner-12m                                                  | 26.53%*** | 13.68% | 1.9390 | 5.88% |
| Top 5 Winner-24m                                                  | 23.94%*** | 14.36% | 1.6670 | 2.26% |
| Top 10 Winner-6m                                                  | 22.07%*** | 13.70% | 1.6112 | 1.39% |
| Top 10 Winner-12m                                                 | 22.67%*** | 13.79% | 1.6442 | 1.86% |
| Top 10 Winner-24m                                                 | 24.16%*** | 14.16% | 1.7060 | 2.78% |
| Winner-6m                                                         | 22.31%*** | 13.48% | 1.6555 | 1.97% |
| Winner-12m                                                        | 22.06%*** | 13.42% | 1.6433 | 1.80% |
| Winner-24m                                                        | 22.05%*** | 13.69% | 1.6101 | 1.38% |
| <i>Self-declared Benchmark</i>                                    | 19.05%*** | 12.63% | 1.5094 | ---   |
| <i>Russel 1000 Growth</i>                                         | 19.84%*** | 11.50% | 1.7246 | ---   |

## 6. How different are winner fund and loser fund characteristics?

In previous sections, we have established that the strategy of purchasing winner funds that have outperformed their self-declared benchmarks and the peer group based on MMT model leads to superior cumulative compounded returns and Sharpe ratios across holding periods of 6-, 12- and 24-month horizons. The results hold for various sizes of winner portfolios – top 5, top 10 and all winners, under varying market - stable periods as well as market crises, when our

winners performance deteriorates with the market, but less than the benchmarks. We have also established that outperformance persists at any reasonable level of transaction costs. Let us look at the differences in how winner and loser funds invest.

In Table 5, we compute the average values for the mutual fund beta coefficients from the MMT model (eq.(1)) used to rank funds as winners, based on 36 months of data prior to the ranking month. The table shows that both all winners, and top 5 winners funds depart from their benchmarks' characteristics in a different way than the loser funds. Specifically, top 5 winners overweight small cap stocks the most relative to their benchmarks - their SMB excess beta ( $\beta_{ib,SMB}$  from eq. (1)) shows 33% excess exposure to small caps (significant at 1%), all winners have 22% more exposure to small caps (although economically significant, it is statistically not significant), while loser funds overweight small caps only by 13% (significant at 1%). The differences between the winner and loser group are even more pronounced when we look at the exposure to value and growth stocks relative to the benchmarks – both groups of winners overweight growth stocks (the HML excess beta,  $\beta_{ib,HML}$  from eq. (1), is significant at 1% level for both winner groups), while loser funds have 3.3% more value stocks in their portfolios (significant at 5% level) relative to their benchmarks. Finally, top 5 winners have the greatest sensitivity to the 'peer-group' factor adjusted for benchmark returns ( $\alpha_{APB,b} + e_{APB,b,t}$ ), followed by all winner funds, while loser funds exhibit the least sensitivity to the peer group (all significant at 1% level). The differences between the SMB excess beta, HML excess beta and peer-group coefficients obtained from eq. (1) are all statistically significant at least at 5% level for top 5 winners vs. all winners, top 5 winners vs losers and all winners vs. losers using Wilcoxon rank/sum z-test, as reported in Table 5. The three groups do not exhibit any difference in terms of general market exposure or momentum.

This shows that funds with weaker performance, classified as losers in this study exhibited greater divergence from both their benchmarks in the large cap growth category, by investing in small cap and value stocks, thus having less sensitivity to the peer group. Small cap and value dimensions have traditionally benefited investors, as the large body of literature from the 1980s and 1990s supports. However, the cyclical behaviour of those styles is a stylised fact (e.g. the dot.com bubble was driven by outperformance of growth stocks), and it is open to further analysis, beyond the scope of this paper, whether loser funds have poorly timed their exposure to value and small cap, compared to winner funds.

**Table 5: Characteristics of winners and losers based on MMT model**

The table shows R-squared and coefficients from the MMT model defined in equation (1) for top 5 winners, all winners and all loser funds. The beta coefficients in the model show exposure to factors in excess of funds' self-declared benchmark. Significance of coefficients is based on t-test (\*\* denoting 1%, \*\*5% and \*1% level of significance). The differences between top 5 and all winners, top 5 and all losers and all winners vs. all losers is gouged with Wilcoxon rank/sum z-test (\*\* denoting 1%, \*\*5% and \*1% level of significance)

| Model<br>MMT model<br>(eq. (1))                   | Top 5<br>winners       | Funds                  |                        | Wilcoxon rank/sum test (z-statistic) |                     |                              |
|---------------------------------------------------|------------------------|------------------------|------------------------|--------------------------------------|---------------------|------------------------------|
|                                                   |                        | All winners            | All losers             | Top 5<br>vs. All<br>winners          | Top 5<br>vs. Losers | All<br>winners<br>vs. Losers |
| R2                                                | 0.5827***<br>(38.82)   | 0.6173***<br>(69.16)   | 0.5121***<br>(37.81)   | 2.359**                              | -3.383***           | -5.247***                    |
| MMT Alpha                                         | 0.0100***<br>(16.43)   | 0.0056***<br>(19.50)   | -0.0043***<br>(-25.68) | -6.067***                            | -8.704***           | -8.704***                    |
| Mkt - rf                                          | 0.0579**<br>(2.46)     | 0.0295***<br>(2.81)    | 0.0150***<br>(3.35)    | 0.084                                | -0.137              | -0.020                       |
| SMB                                               | 0.3286***<br>(9.96)    | 0.2197<br>(11.05)      | 0.1274***<br>(11.95)   | -2.526**                             | -5.598***           | -3.708***                    |
| HML                                               | -0.2694***<br>(-10.37) | -0.1553***<br>(-10.44) | 0.0334**<br>(2.29)     | 3.310***                             | 8.061***            | 7.248***                     |
| MOM                                               | 0.0766***<br>(3.82)    | 0.0580***<br>(4.64)    | 0.0602***<br>(5.67)    | 0.258                                | 0.532               | 0.388                        |
| Peer group:<br>( $\alpha_{APB,b} + e_{APB,b,t}$ ) | 2.1605***<br>(21.79)   | 1.5749***<br>(40.32)   | 0.5548***<br>(23.10)   | -5.190***                            | -8.657***           | -8.506***                    |

## 7. Conclusions

This study extends the literature on mutual fund performance persistence by examining the characteristics of persistence of ‘true winner’ and ‘true loser’ mutual funds, as identified by the benchmark- and peer-group adjusted alphas from the Mateus, Mateus and Todorovic (2025) novel factor model (MMT). Unlike standard Fama-French-Carhart (FFC) approaches, may misclassify performance, the MMT model offers a more industry-aligned framework. It simultaneously adjusts for both the fund’s self-declared benchmark and fund’s peer-group, thereby yielding a more accurate identification of outperforming funds that investors are interested in. Our study contributes to the literature by testing whether MMT winners persist at different investment horizons, over varying market conditions and after trading costs are accounted for. We also consider whether outperformance is driven by a handful of top performing funds.

We perform our empirical tests on a sample of 735 US Large Cap Growth mutual funds, classified by Morningstar, as the largest US peer-group of funds. The sample spans from January 1990 – June2018. Our empirical findings provide robust evidence that the MMT-identified winner portfolios consistently outperform their benchmarks and peer groups,

regardless of holding period length (6, 12, or 24 months), with the strongest performance observed for shorter rebalancing frequencies. Importantly, this performance persists across a range of market conditions, including crisis periods such as the dot-com bubble and the Global Financial Crisis. While returns are naturally lower during these downturns, MMT winners still outperform their benchmarks on a risk-adjusted basis. Additionally, we show that this outperformance survives plausible levels of transaction costs, making the strategy implementable in real-world settings.

Furthermore, we find that performance persistence is not driven solely by a handful of top-performing funds. Even broadly diversified portfolios comprising all MMT winners (often 60 funds or more) deliver substantial outperformance, reducing the need for excessive fund turnover or overly concentrated bets. The model's ability to identify persistently skilled funds stems in part from its capacity to capture stylistic tilts—such as exposure to small-cap and growth stocks—and sensitivity to peer dynamics not captured in traditional models.

Overall, this paper contributes to the evolving discussion on mutual fund evaluation by demonstrating the efficacy of a dual-adjusted alpha model, the MMT, in predicting future fund performance. Future research may explore how the MMT model interacts with other dimensions of manager behaviour, such as market timing or fund flow dynamics, to further refine performance assessment and fund selection strategies.



## Bibliography

- Angelidis, T., Giamouridis, D. and Tessaromatis, N., 2013. Revisiting mutual fund performance evaluation. *Journal of Banking & Finance*, 37(5), pp.1759-1776.
- Bollen, N. P. B., & Busse, J. A. (2005). *Short-term persistence in mutual fund performance*. The Review of Financial Studies, 18(2), 569–597.
- Busse, J. Chordia, T., Jiang, L.; and Tang, Y. (2017). Mutual Fund Trading Costs and Diseconomies of Scale. 1-90. Available at: Available at: [https://ink.library.smu.edu.sg/lkcsb\\_research/4536/](https://ink.library.smu.edu.sg/lkcsb_research/4536/)  
Accessed: 18<sup>th</sup> Jauary 2026.
- Carhart, M. M., 1997, On persistence in mutual fund performance, *Journal of Finance* 52, pp. 57-82.
- Chan, L. K. C., S. G. Dimmock, and J. Lakonishok. 2009. “Benchmarking Money Manager Performance: Issues and Evidence.” *Review of Financial Studies* 22(11): 4553–4599.
- Chen, Z. and Knez P.J., (1996), Portfolio Performance Measurement: Theory and Applications, *Review of Financial Studies*, 2:511-555.
- Chinthalapati, V.L., Mateus, C. and Todorovic, N., 2017, Alphas in disguise: A new approach to uncovering them, *International Journal of Finance and Economics*, 22 (3), pp. 234-243.
- Cremers, M., Petajisto, A. and Zitzewitz, E., 2012. Should Benchmark Indices Have Alpha? Revisiting Performance Evaluation, *Critical Finance Review*, 2, pp.1-48
- Cremers, K. J. M., Fulkerson, J. A. and Riley, T. B., 2019, Challenging the Conventional Wisdom on Active Management: A Review of the Past 20 Years of Academic Literature on Actively Managed Mutual Funds, *Financial Analysts Journal*, 75(4), pp. 8-32.
- Cuthbertson, K., Nitzsche D. and O’Sullivan N., 2010, Mutual fund performance: measurement and evidence, *Financial Markets, Institutions and Instruments*, 19 (2), pp. 95-187.
- Cuthbertson, K., Nitzsche, D. & O’Sullivan, N. 2022. Mutual fund performance persistence: Factor models and portfolio size. *International Review of Financial Analysis*, 81, pp.102-133.
- Fama, E. F., and French, K. R., 1993, Common risk factors in the returns on stocks and bonds, *Journal of Financial Economics* 33, pp. 3-56.
- Fletcher J. and D. Forbes, 2002, An exploration of the persistence of UK unit trust performance, *Journal of Empirical Finance*, 9 (5), pp. 475-493.
- Hunter, D., Kandel E., Kandel S. and Wermers R., 2014, Mutual fund performance evaluation with active peer benchmarks, *Journal of Financial Economics*, 112(1), pp. 1-29.
- Kosowski, R., Timmermann, A., Wermers, R., & White, H. (2006). *Can mutual fund “stars” really pick stocks? New evidence from a bootstrap analysis*. The Journal of Finance, 61(6), 2551–2595
- Matallin-Saez, J. C. 2007. “Portfolio Performance: Factors or Benchmarks.” *Applied Financial Economics* 17 (14): 1167–1178.
- Mateus, I.B., Mateus, C. and Todorovic, N., 2016, UK equity mutual fund alphas make a comeback. *International Review of Financial Analysis*, 44, pp. 98-110.
- Mateus, I.B., Mateus, C. and Todorovic, N., 2019a, Review of new trends in the literature on factor models and mutual fund performance. *International Review of Financial Analysis*, 63, pp. 344–354.
- Mateus, I.B., Mateus, C. and Todorovic, N., 2019b. Use of active peer benchmarks in assessing UK mutual fund performance and performance persistence. *The European Journal of Finance*, 25(12), pp.1077-1098.
- Mateus, I.B., Mateus, C. and Todorovic, N., 2023. Searching for mutual fund winners? The strategy is to outbid both, the benchmark and the peer group”, *Applied Economics*, 56(11), pp. 1268–1282.
- Mateus, I.B., Mateus, C. and Todorovic, N., 2025. Mutual fund performance: the model for selecting persistent winners. *The European Journal of Finance*, 31(5), pp.647-669.