



Code & Kapital Research Department

52-Week High Recency 12 Signal

Code&Kapital Research

Tim Bastian

September 11, 2026



Abstract

This paper sets out a practical framework for evaluating candidate cross-sectional equity signals in a consistent and repeatable way. The goal is not only to measure whether a predictor appears statistically interesting, but also to determine whether its behaviour survives the kinds of implementation frictions that often separate a promising backtest from a usable trading signal.

Within this framework, each proposed predictor is examined in relation to portfolio construction, trading costs, concentration, turnover, and capacity-related constraints, while also being compared against neighbouring signals that may capture similar effects. This matters because a signal that looks strong in a basic return sort can still lose much of its appeal once realistic implementation constraints are introduced.

By combining a disciplined testing protocol with reproducible documentation, the aim is to make signal research easier to review, easier to compare across ideas, and less vulnerable to conclusions that look attractive in-sample but are difficult to carry into practice. The broader objective is to identify signals that are not only statistically interesting, but also robust enough to remain useful under realistic portfolio conditions.

About Code&Kapital

Code&Kapital is a quantitative research and software business built for serious practitioners in systematic investing. The company sits at the intersection of investment research, financial data engineering, and research infrastructure, with a focus on turning good ideas into dependable workflows rather than isolated notebook results.

The broader mission is to close the gap between attractive market research and real implementation. That means treating data integrity, validation design, execution assumptions, portfolio construction, and operational repeatability as core parts of the research process rather than as details to be addressed later. In practice, Code&Kapital publishes applied research, develops tooling for signal validation and portfolio analytics, and builds data systems that support repeatable, production-minded quantitative work. The intended audience includes quants, engineers, portfolio managers, allocators, and technically serious investors who care not only about whether an idea works, but also about how the result was produced and whether the workflow can be trusted over time. Across research, software, data, and educational content, the emphasis remains the same: practical implementation over elegant but fragile abstractions.

The signal testing framework behind this paper follows that same philosophy. Each report is designed to evaluate a candidate signal through a consistent research standard: return behaviour, robustness, implementation frictions, cross-signal context, and the practical constraints that determine whether a signal is useful outside of a clean backtest. The aim is not simply to identify statistically interesting predictors, but to understand where they are durable, where they are fragile, and how they fit within a broader research stack. That includes asking whether a signal complements existing ideas, whether its performance is likely being driven by hidden overlap with known effects, and whether implementation constraints such as turnover, concentration, and trading frictions materially change the conclusion.

This paper is part of an ongoing research publication process and is intended to be updated on a monthly basis as new observations, diagnostics, comparisons, and implementation notes are added.

Contents

Abstract	i
About Code&Kapital	ii
Contents	iii
1 Introduction	1
1.1 Strategy Overview	1
1.2 What Matters Most	2
2 Signal Diagnostic	3
3 Return Predictability	6
3.1 Baseline Portfolio Sort	6
3.2 Exposure	7
3.3 Sector Evidence	7
3.4 Robustness Across Construction Choices	9
3.5 Size-Conditioned Evidence	10
3.6 Investability	11
4 Comparison With Documented Anomalies	19
4.1 Relationship To Known Anomalies	19
4.2 Performance Benchmark	20
Appendices	24
A Signal Catalog	25
B Turnover-Constrained Optimizer	27
C Exposure-Matching Optimizer	28
Bibliography	29

CHAPTER 1

Introduction

This report presents the **52-Week High Recency 12** strategy inside the Code&Kapital signal testing framework for cross-sectional equity research. The purpose of the framework is to evaluate signals in a way that is not only statistically informative, but also useful for practical portfolio work. At a high level, the framework is designed to move from signal construction, to diagnostic validation, to portfolio results, and then further into the conditions under which a signal is more or less effective. Rather than asking only whether a backtest looks strong in isolation, the broader goal is to understand how a signal behaves across the investable universe, whether it provides enough dispersion to sort stocks in a meaningful way, and whether the resulting portfolios show patterns that are robust enough to matter for research and implementation decisions. The research is carried out on the **Russell 1000**, which serves both as the underlying reference index and as the investable universe for the signal tests. Using a clearly defined universe helps keep the evaluation tied to a realistic opportunity set, while also ensuring that comparisons across signals are made on a common and interpretable base. As a first step, the next chapter turns to signal diagnostics, focusing on how the signal's value metrics evolve through time and how much coverage the signal achieves across the Russell 1000 universe. From there, the report moves into three linked portfolio tests. Section 3.1 asks whether the signal produces a clean and economically meaningful spread in the baseline quintile sort. Section 3.2 then asks how much of that result survives once the portfolios are compared against market and style proxies in the factor-model regressions. Section 3.3 turns to sector evidence, asking whether the signal works broadly across the market or whether the headline result is being driven mainly by a narrower subset of industries. Taken together, those sections move from raw payoff, to explained payoff, to breadth of payoff. The remaining sections extend that core sequence into the broader framework questions the report is meant to answer. After the baseline sort, the analysis checks start-date sensitivity, asks what part of the payoff survives standard factor controls, and tests whether the evidence is broad across sectors or narrowly concentrated. It then studies robustness across alternative portfolio construction choices in Section 3.4, and adds a dedicated size-conditioned section in Section 3.5 to show whether the signal behaves differently across the market-cap spectrum once the Russell 1000 is broken into large- and small-name buckets. From there, the report turns more directly to implementation in Section 3.6. That section moves from transaction-cost evidence to turnover-constrained portfolio construction and path dependency, then to benchmark-aware index implementations based on sector matching and broader exposure matching. Finally, Section 4 places the signal in a broader empirical context by comparing it with the anomaly signals covered by the workflow, the collection often referred to as the “zoo” in the literature. Taken together, the framework is designed to answer not only whether the signal works in a clean cross-sectional ranking exercise, but also where it works, how stable it is, how sensitive it is to implementation frictions and starting conditions, whether it is concentrated in particular parts of the universe, and whether it still adds value once the portfolio is pushed toward a realistic benchmark-aware implementation and then benchmarked against documented anomaly strategies.

1.1 Strategy Overview

The 52-week-high-recency strategy implemented in this report follows Bhootra and Hur [BH13], who argue that the timing of a stock's trailing 52-week high contains incremental information about future returns. Their central idea is that investors exhibit recency bias: when a stock has reached its 52-week high only recently, investors remain reluctant to fully incorporate the information that pushed the

stock to that level. As a result, stocks with more recent 52-week highs tend to outperform stocks whose 52-week highs were reached further in the past.

Operationally, the signal measures how recently the trailing 52-week high was attained within the prior 12 months:

$$RR_{i,t} = 1 - \frac{\text{DaysSinceHigh}_{i,t}}{365}.$$

Higher values therefore indicate that the stock reached its trailing 52-week high more recently, while lower values indicate that the high occurred further back in time.

The portfolio construction preserves the paper's direction. At each monthly rebalance, the cross-section is sorted on the recency ratio, the long side is tilted toward stocks with the most recent trailing 52-week highs, and the short side is tilted toward stocks with the most distant highs. The paper also shows that conditioning the standard 52-week-high ratio strategy on this recency measure further strengthens its profitability.

Inside the Russell 1000 universe, the ranked cross-section is sorted into quintiles at each rebalance. The top quintile contains stocks with the most recent trailing 52-week highs, the bottom quintile contains those with the most distant highs, and the long-short portfolio buys the top bucket while shorting the bottom bucket. The empirical question for this report is whether the recency channel highlighted by the paper remains broad, interpretable, and persistent.

1.2 What Matters Most

The performance summary provides the quickest view of the main takeaways. In the current configuration, the strategy shows a total return of **1.4%**, an annualized return of **1.9%**, a Sharpe ratio of **0.2**, and a maximum drawdown of **-21.5%**. The strongest long-short alpha currently appears in the **Market** specification at **0.4%** per month. At the sector level, the strongest result appears in **Technology** at **4.8%** average monthly excess return, while the weakest appears in **Communication Services** at **-2.3%**. In the six-factor model, the largest absolute long-short loading is on **Size** at **3.07**.

That headline view becomes more informative when the signal is compared across the portfolios it is meant to separate. Table 1.1 contrasts the strongest-ranked bucket, weakest-ranked bucket, long-short spread, and the Russell 1000 benchmark across a small set of metrics that help frame return, risk, and implementation quality at a glance.

Table 1.1: Headline comparison across sorted signal portfolios and the Russell 1000 benchmark

Metric	Top Quintile	Bottom Quintile	Long-Short	Russell 1000
Start	Dec 2025	Dec 2025	Dec 2025	Dec 2025
Total Return	3.5%	0.7%	1.4%	12.1%
CAGR	4.6%	0.8%	1.9%	15.8%
YTD	4.5%	-1.5%	4.7%	11.6%
Max Drawdown	-13.3%	-13.6%	-21.5%	-9.1%
VaR 95	-1.7%	-1.8%	-2.8%	-1.3%
Sharpe	0.14	-0.09	0.2	0.93

Description: Top Quintile contains the highest-ranked names on the signal, Bottom Quintile contains the lowest-ranked names, Long-Short is the dollar-neutral spread that buys the top bucket and sells the bottom bucket, and Russell 1000 is the benchmark universe.

The *Top Quintile* column refers to the portfolio formed from the highest-ranked names on the signal, while *Bottom Quintile* captures the lowest-ranked names. *Long-Short* represents a dollar-neutral spread that buys the top names and sells the bottom names, making it the clearest summary of whether the ranking adds value. *Russell 1000* provides the benchmark view of the underlying investable universe. Across those columns, *Start* marks the first date at which the return series is available, *Total Return* reports the cumulative return over the full sample, *CAGR* reports the annualized compound growth rate over that same period, *YTD* shows the return earned in the current calendar year, *Max Drawdown* captures the worst peak-to-trough decline, *VaR 95* shows the loss threshold exceeded only in the worst five percent of periods, and *Sharpe* measures return relative to volatility. Together, these comparisons provide a compact first view of payoff profile, downside risk, and how clearly the signal separates stronger from weaker names.

CHAPTER 2

Signal Diagnostic

This chapter evaluates the basic statistical behaviour and practical availability of the **52-Week High Recency 12** signal before turning to deeper portfolio interpretation. A useful signal should show enough cross-sectional variation to separate stronger from weaker names, while also covering a broad enough share of the investable universe to support realistic implementation.

Figure 2.1 reports signal coverage across the Russell 1000 universe. The signal is available for **99.0%** of names, representing **98.5%** of total market capitalisation. Coverage should be read in two ways: by count, it shows how broadly the signal reaches across the universe; by market value, it shows how much of the investable capital base is actually represented in the signal.

At each rebalance, extreme cross-sectional outliers are excluded. Where the signal definition requests cross-sectional standardization, the remaining values are expressed as z-scores; otherwise the diagnostic retains the signal's native scale. This keeps unusually large observations from dominating the score without implying that every signal uses the same transformation.

Figure 2.2 then summarises the time-series behaviour of the cleaned cross-sectional distribution actually used by the signal. The average cross-sectional mean (median) of the signal is **0.6 (0.7)**, while the interquartile range runs from **0.26 to 0.92**. Its average cross-sectional standard deviation is **0.35**. The mean, standard deviation, median, and interquartile range describe the realized signal scale and show whether the cleaned score remains balanced or retains meaningful asymmetry through time.

Figure 2.3 closes the diagnostic with the latest cross-sectional distribution of the signal. A histogram view is helpful because it makes it easier to see whether the signal is skewed to one side, whether a small number of observations dominate the tails, and whether the cross-section is tightly clustered or more broadly spread out. The latest cross-section is skewed to the left, with a longer negative tail than positive tail. The tails appear relatively light, with fewer extreme observations than a bell-shaped distribution would imply. The current spread sits close to its typical historical range.

Together, the coverage, factor-statistics, and distribution-shape views provide an early diagnostic of whether the signal is rich enough to sort the universe, stable enough to interpret, and broad enough to matter in practice across different market environments.

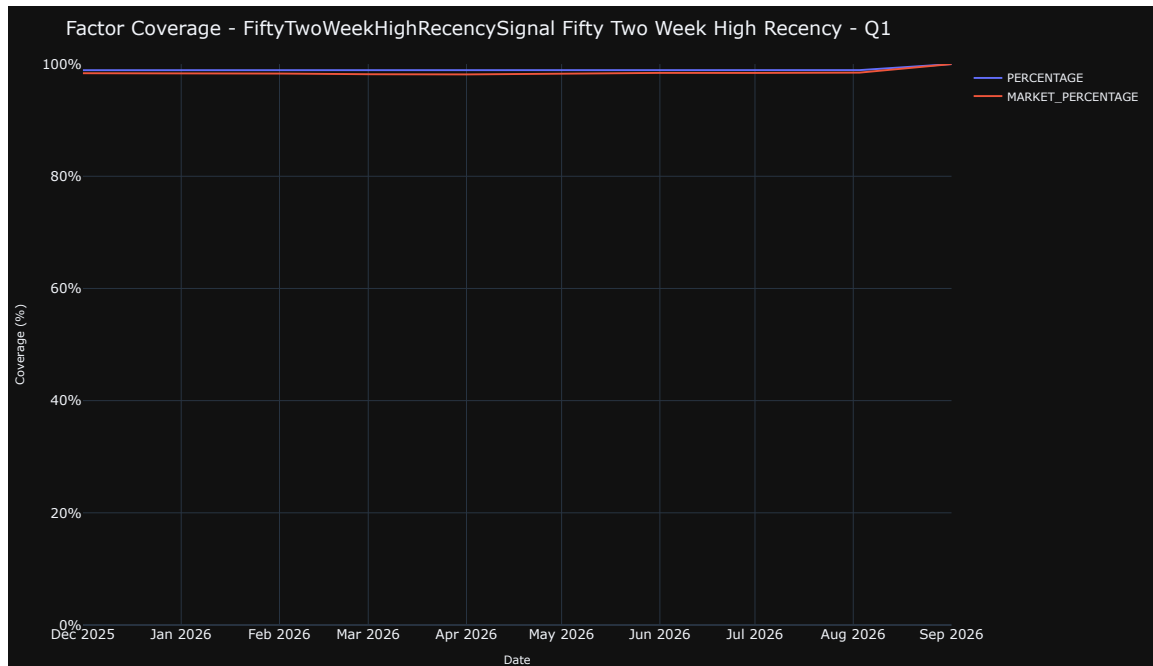


Figure 2.1: Coverage of the signal across the Russell 1000 universe

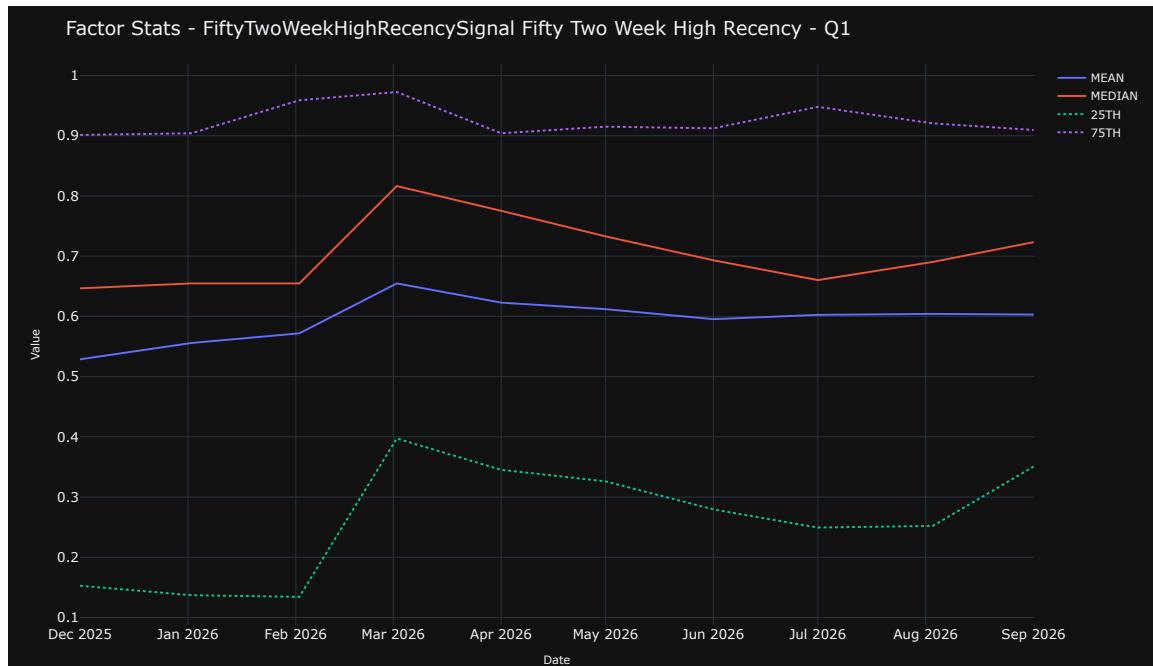


Figure 2.2: Time-series of cross-sectional statistics for the signal

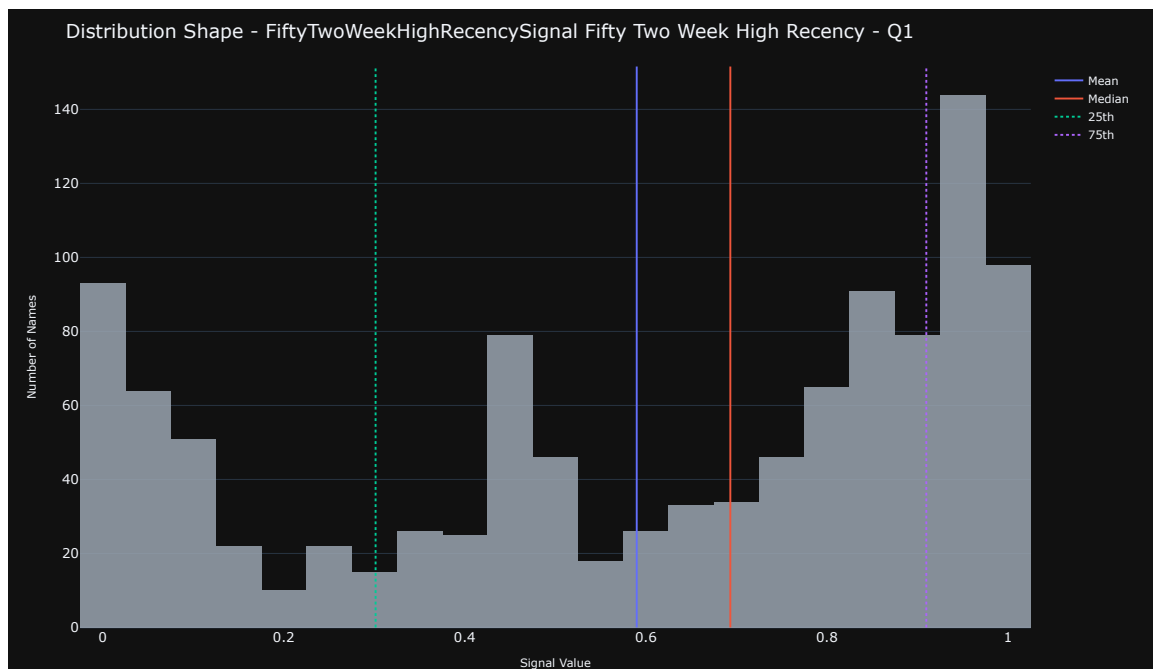


Figure 2.3: Latest cross-sectional distribution of the signal

CHAPTER 3

Return Predictability

This chapter turns from signal diagnostics to the portfolio evidence itself. The goal is not only to ask whether the signal produces attractive spread returns, but also whether those returns remain meaningful once we look at factor exposure, alternative portfolio constructions, trading frictions, and the role of firm size.

3.1 Baseline Portfolio Sort

The baseline implementation sorts the Russell 1000 universe into quintiles using the signal and forms market-weighted portfolios within each bucket. It rebalances those portfolios monthly and then compares the strongest-ranked and weakest-ranked portfolios, together with the long-short portfolio that buys the top bucket and sells the bottom bucket. This gives a direct first read on whether the signal is separating stronger and weaker names in an economically meaningful way.

Under that baseline construction, the high-ranked long-only portfolio delivers an average monthly return of **0.6%** with an annualized Sharpe ratio of **0.14**. The corresponding long-short portfolio delivers an average monthly return of **0.7%** and an annualized Sharpe ratio of **0.2**. On average, the top bucket holds **200** names, while the bottom bucket holds **199.3**. The average market cap of the top bucket is **\$102.0B**, versus **\$34.7B** for the bottom bucket. Because size can matter for both signal strength and implementation, Section 3.5 returns to this question directly in the dedicated size-conditioned analysis. Figure 3.1 shows how the top, bottom, and long-short portfolios evolve cumulatively through time. Table 3.1 shows how that performance builds across the full quintile sort, from the high-ranked bucket through the low-ranked bucket and the resulting long-short portfolio. Taken together, these baseline results show a positive high-minus-low spread, but the quintile pattern is not especially monotonic from the high-ranked bucket to the low-ranked bucket. That points to some separation power, though the cross-sectional ordering looks less consistent before turning to the factor-model breakdown that follows in the next section.

Table 3.1: Baseline long-short evidence under the main portfolio construction

Measure	T	2	3	4	B	T-B
Average Monthly Return	0.6%	0.8%	2.9%	1.3%	-0.1%	0.7%
Average Monthly Excess Return (T, B) / Return (T-B)	0.3%	0.5%	2.6%	1.1%	-0.4%	0.7%
Annualized Sharpe Ratio	0.14	0.42	1.97	0.78	-0.09	0.2
Average Securities Held	200	199.9	199.7	199.4	199.3	399.3
Average Market Cap	\$102.0B	\$92.8B	\$92.6B	\$53.3B	\$34.7B	\$68.4B

Notes: Portfolios are sorted monthly on the signal and are market-weighted within each quintile. Long-only bucket excess returns are measured relative to the 3-month U.S. Treasury bill rate, proxied by the FRED ‘DTB3’ series and aligned to the backtest dates; the dollar-neutral T-B spread is shown as its raw self-financing return. Bold values in the monthly return rows indicate statistical significance at the 5% level based on a two-sided t-test, using the exact Student-*t* critical value for the available number of monthly observations in each portfolio series.

To show how sensitive the return estimate is to the chosen sample window, Table 3.2 restarts the measurement every six months and reports the long-only excess returns and self-financing T-B return from that point through the end of the available sample.

Table 3.2: Monthly return by sample start date

Start Date	T	B	T-B
Dec 2025	0.2%	-0.1%	0.3%
Jun 2026	-2.5%	0.9%	-3.3%

Notes: Each row uses the same monthly market-weighted portfolio construction as the baseline sort, but starts the sample at the stated month and carries it through the final available month. The T and B columns are excess returns relative to the 3-month U.S. Treasury bill rate, while T-B is the raw self-financing spread return.

3.2 Exposure

Strong spread performance is more informative when we can also see what kinds of systematic exposure are sitting behind it. To make that comparison timely and operational, the workflow estimates monthly factor models using tradable ETF and index proxies. In this setup, the market proxy is the Russell 1000, while the style legs use investable proxies for size, value, momentum, profitability, and investment. The one-factor specification controls only for the market proxy; the three-factor version adds size and value; the four-factor version adds momentum; the five-factor version uses profitability and investment alongside market, size, and value; and the six-factor version includes all six exposures together.

Table 3.3 reports the monthly alpha estimates for each portfolio across the available proxy-model specifications. Table 3.4 reports the corresponding factor loadings. Together, those tables show whether the return spread survives broader controls and which risk dimensions are doing the most explanatory work. For the long-short portfolio, the highest monthly alpha appears in the **Market** specification at **0.4%** with a t-statistic of **0.16**. Across the available specifications, the long-short alpha is statistically significant in **0** of **5** cases. In the six-factor model, the largest absolute loading is on **Size** at **3.07**, with a t-statistic of **2.05**.

Table 3.3: Monthly alpha estimates across ETF-proxy factor models

Model	T	2	3	4	B	T-B
Market	-0.4%	-0.3%	1.3%	0.2%	-0.8%	0.4%
3-Factor	-2.8%	0.6%	2.5%	0.1%	0.3%	-3.1%
4-Factor	-2.5%	0.8%	2.7%	-1.0%	-0.4%	-2.1%
5-Factor	-2.9%	0.6%	2.7%	0.0%	0.2%	-3.1%
6-Factor	-2.4%	0.7%	2.4%	-0.3%	0.0%	-2.4%

Notes: Long-only sleeve alphas are estimated from monthly excess returns. The dollar-neutral long-short spread is self-financing, so its raw spread return is regressed on factor excess returns. ETF or index proxies supply the factor legs. Bold values indicate statistical significance at the 5% level based on the regression t-statistic.

Table 3.4: Factor loadings across ETF-proxy factor models

Factor	T	2	3	4	B	T-B
Market	-1.03	1.04	1.24	1.28	1.54	-2.58
Size	2.01	0.36	-1.92	-2.31	-1.06	3.07
Value	0.62	-0.45	-0.28	0.33	-0.26	0.87
Profitability	-1.92	0.45	3.39	0.11	-0.11	-1.81
Investment	0.32	-0.43	-1.28	2.49	1.44	-1.12
Momentum	0.41	0.05	-0.28	-0.25	-0.17	0.58

Notes: Factor loadings are reported from the six-factor specification estimated on the same regressions as the alpha table. Bold values indicate statistical significance at the 5% level based on the regression t-statistic.

3.3 Sector Evidence

To understand whether the signal works broadly across the market or leans on a smaller subset of industries, the framework also runs the long-short portfolio inside each Russell 1000 sector separately.



Figure 3.1: Cumulative return paths for the top-ranked portfolio, bottom-ranked portfolio, and long-short portfolio

This helps distinguish a signal that is widely useful from one whose headline result is being driven mainly by just a few favorable sector regimes. The portfolio construction remains the same as in the baseline test: stocks are sorted into quintile portfolios on the signal, those portfolios are rebalanced monthly, and the sector long-short portfolio buys the top quintile and sells the bottom quintile within each sector only.

Table 3.5 summarizes the sector-by-sector long-short results. Figure 3.2 ranks sectors by average monthly return, while Figure 3.3 places those same returns against absolute max drawdown to show the basic sector-level risk/return trade-off.

Table 3.5: Sector long-short evidence

Sector	Avg Monthly Return	Sharpe	Max DD	Total Return
Utilities	-1.1%	-0.65	-18.9%	-13.4%
Industrials	-0.3%	-0.26	-22.3%	-7.1%
Technology	4.8%	0.66	-51.7%	17.7%
Real Estate	-1.1%	-0.99	-15.3%	-14.0%
Healthcare	-0.9%	-0.91	-29.9%	-14.5%
Consumer Cyclical	-1.3%	-0.54	-21.5%	-9.9%
Financial Services	0.9%	0.51	-19.2%	5.7%
Consumer Defensive	0.8%	0.67	-12.6%	8.2%
Energy	0.2%	-0.09	-21.2%	-2.7%
Basic Materials	-0.9%	-0.17	-27.2%	-8.1%
Communication Services	-2.3%	-1.2	-36.9%	-26.4%

Notes: Each row reports the sector-specific long-short portfolio that buys the top signal bucket and sells the bottom signal bucket within that sector only. Portfolios are rebalanced monthly and weighted using the same market-weighted schedule as the baseline tests. Average monthly return is the raw self-financing spread return; it is not reduced by the 3-month U.S. Treasury bill rate.

Sector performance is strongest in **Technology**, where the long-short portfolio earns an average monthly return of **4.8%**. The weakest result appears in **Communication Services** at **-2.3%**. The highest Sharpe ratio comes from **Consumer Defensive** at **0.67**. In total, **4** of **11** sectors post positive

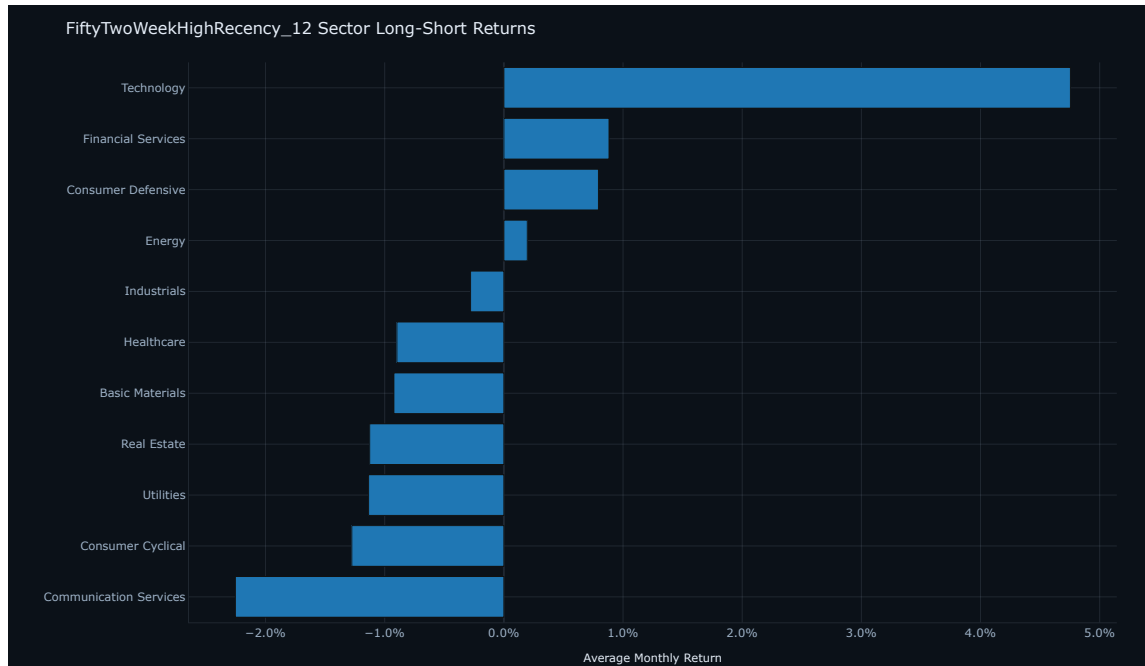


Figure 3.2: Average monthly return of the sector long-short portfolios

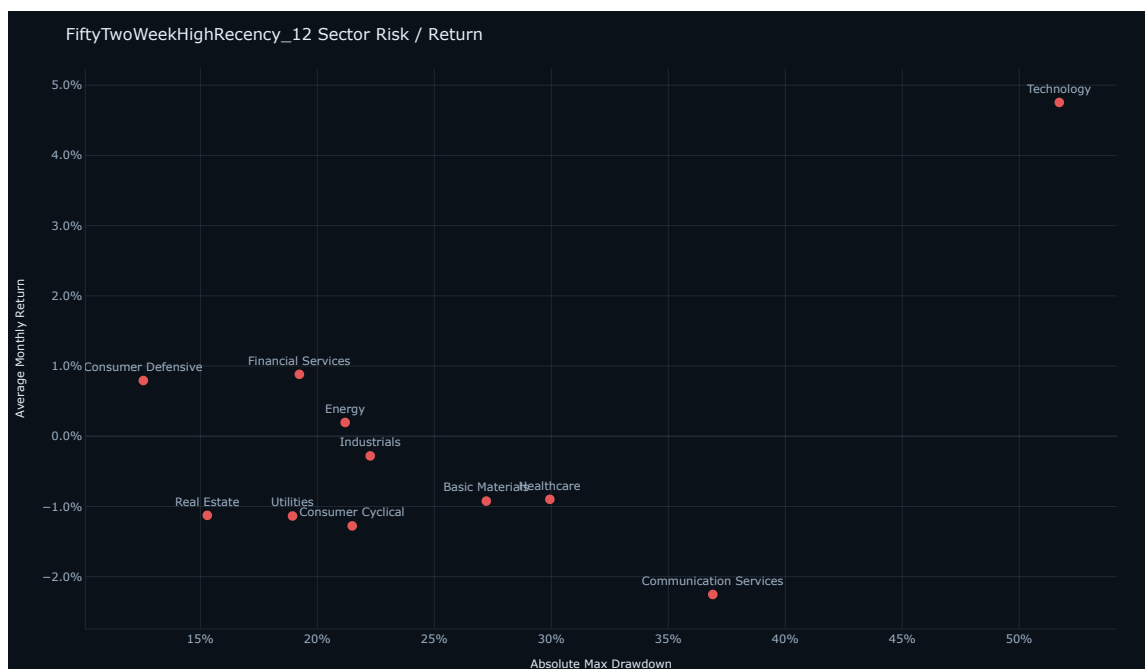


Figure 3.3: Sector long-short average monthly return versus absolute max drawdown

average monthly returns, which points to a fairly concentrated result, with most of the strength coming from a smaller subset of sectors.

3.4 Robustness Across Construction Choices

Signals can look stronger on paper simply because of how the portfolios are formed. For that reason, the framework also compares the baseline result with a range of alternative constructions, including decile instead of quintile sorting, equal instead of market weighting, an S&P 500 universe instead of

the Russell 1000, and daily or quarterly rebalancing instead of the baseline monthly schedule. This is an important step, because a signal that survives those variations is much more likely to reflect something robust rather than a narrow implementation choice.

Across the robustness checks, long-short average monthly returns range from **-0.0%** to **9.7%** per month, with the weakest specification still showing a t-statistic of **-0.01**. Across the tested construction-model combinations, **0** out of **30** alpha estimates remain statistically meaningful.

The strongest average monthly return comes from the **Quintile, Market, Daily, R1000** specification at **9.7%**, while the weakest result comes from **Quintile, Market, Quarterly, R1000** at **-0.0%**. Most reported variants remain positive, which still points to a fairly stable signal across the main implementation choices. The highest model-adjusted alpha appears in the **Quintile, Market, Daily, R1000** specification under the **6-Factor** model at **11.7%**.

Table 3.6: Robustness across construction variants

Sort	Weights	Frequency	Universe	Average Monthly Return	Market	3-Factor	4-Factor	5-Factor	6-Factor
Quintile	Market	Monthly	R1000	0.7%	0.4%	-3.1%	-2.1%	-3.1%	-2.4%
Quintile	Equal	Monthly	R1000	0.7%	0.8%	-1.6%	-0.6%	-1.5%	-1.2%
Decile	Market	Monthly	R1000	0.4%	0.3%	-2.7%	-1.2%	-2.9%	-1.5%
Quintile	Market	Monthly	S&P 500	0.7%	0.3%	-3.2%	-2.2%	-3.3%	-2.5%
Quintile	Market	Daily	R1000	9.7%	8.3%	8.2%	7.9%	6.7%	11.7%
Quintile	Market	Quarterly	R1000	-0.0%	-0.4%	-3.3%	-2.2%	-3.2%	-2.6%

Notes: Average Monthly Return reports the raw self-financing return of the long-short portfolio. The remaining columns report long-short alphas across the full factor-model set, including the six-factor specification. Bold entries indicate statistical significance at the 5% level.

3.5 Size-Conditioned Evidence

Because many cross-sectional effects become much weaker once the smallest and least liquid names are removed, the framework also looks at signal performance within size buckets. This helps separate a broadly useful signal from one that is largely concentrated in parts of the market that are harder to scale.

Table 3.7 reports the within-size excess returns, Table 3.8 reports the corresponding factor-model alphas, and Table 3.9 shows the realized market-cap profile in each cell. Here, Size quintile 1 refers to the largest names, Size quintile 5 to the smallest, Signal Q1 to the strongest signal values, and Signal Q5 to the weakest. The strongest size-conditioned excess-return cell appears in **Size quintile 5, Signal Q3**, where the average monthly excess return reaches **2.6%**. Across the within-size long-short portfolios, the highest factor-adjusted alpha comes from **Size quintile 1** under the **Market** specification at **2.0%**. Taken together, these tables show whether the signal remains present after conditioning on size rather than relying only on the smallest corner of the universe.

Table 3.7: Size-conditioned excess returns

Size Bucket	Signal Q1	Signal Q2	Signal Q3	Signal Q4	Signal Q5
Size quintile 1	2.3%	0.7%	1.6%	1.9%	-0.2%
Size quintile 2	1.3%	0.1%	0.2%	0.9%	0.2%
Size quintile 3	-0.3%	0.1%	2.1%	-0.6%	1.3%
Size quintile 4	0.6%	-0.2%	1.9%	0.7%	1.1%
Size quintile 5	1.1%	1.4%	2.6%	-0.6%	0.2%

Notes: Each cell reports the average monthly excess return for the corresponding size-signal portfolio. Size quintile 1 contains the largest market-cap names and size quintile 5 the smallest; Signal Q1 contains the strongest signal values and Signal Q5 the weakest. Bold entries indicate statistical significance at the 5% level.

Table 3.8: Within-size long-short alphas

Size Bucket	Market	3-Factor	4-Factor	5-Factor	6-Factor
Size quintile 1	2.0%	0.3%	1.2%	0.3%	0.5%
Size quintile 2	1.1%	-1.5%	-0.5%	-1.3%	-0.9%
Size quintile 3	-1.2%	-3.0%	-2.5%	-3.2%	-3.2%
Size quintile 4	-0.3%	-2.0%	-0.7%	-2.0%	-1.2%
Size quintile 5	0.8%	-0.7%	-0.2%	-0.8%	-1.1%

Notes: Each cell reports the long-short alpha within a given size bucket under the indicated factor model, where the long leg buys Signal Q1 and the short leg sells Signal Q5. Size quintile 1 contains the largest market-cap names and size quintile 5 the smallest; Signal Q1 contains the strongest signal values and Signal Q5 the weakest. Bold entries indicate statistical significance at the 5% level.

Table 3.9: Average market capitalization within size-signal portfolios

Size Bucket	Signal Q1	Signal Q2	Signal Q3	Signal Q4	Signal Q5
Size quintile 1	\$82.5B	\$82.3B	\$80.5B	\$87.0B	\$81.4B
Size quintile 2	\$29.5B	\$29.5B	\$29.6B	\$29.3B	\$29.6B
Size quintile 3	\$15.7B	\$15.7B	\$15.5B	\$15.5B	\$15.3B
Size quintile 4	\$9.4B	\$9.3B	\$9.3B	\$9.4B	\$9.3B
Size quintile 5	\$5.1B	\$5.1B	\$4.7B	\$4.6B	\$4.2B

Notes: Each cell reports the average market capitalization of the names held in the corresponding size-signal portfolio. Size quintile 1 contains the largest market-cap names and size quintile 5 the smallest; Signal Q1 contains the strongest signal values and Signal Q5 the weakest.

3.6 Investability

Strong headline results are more useful when they can survive the practical constraints that come with implementation. This section therefore moves from return predictability in the abstract to return predictability under trading, turnover, and benchmark-aware portfolio design. The goal is to understand not just whether the signal works, but whether it remains attractive once the portfolio has to behave more like something that could actually be run.

Transaction Costs

Investability matters not just through holdings breadth and size, but also through what remains after implementation frictions are applied. In the current setup, average net returns after trading costs still range from **-0.1%** to **9.3%** per month. That helps distinguish a signal that is only attractive before frictions from one that remains usable once a more realistic trading lens is applied.

After applying a flat one-way transaction cost of 10 bps, the strongest average monthly return comes from the **Quintile, Market, Daily, R1000** specification at **9.3%**, while the weakest result comes from **Quintile, Market, Quarterly, R1000** at **-0.1%**. Most reported variants remain positive after costs, which suggests that the signal largely survives implementation frictions even though the payoff is compressed. Relative to the gross robustness results in Section 3.4, trading costs reduce the average monthly return by about **0.2%** on average across the tested constructions. The highest post-cost alpha appears in the **Quintile, Market, Daily, R1000** specification under the **6-Factor** model at **11.2%**.

Table 3.10: Investability after transaction costs

Sort	Weights	Frequency	Universe	Average Monthly Return	Market	3-Factor	4-Factor	5-Factor	6-Factor
Quintile	Market	Monthly	R1000	0.5%	0.2%	-3.3%	-2.3%	-3.3%	-2.6%
Quintile	Equal	Monthly	R1000	0.5%	0.6%	-1.8%	-0.8%	-1.7%	-1.3%
Decile	Market	Monthly	R1000	0.1%	0.0%	-2.9%	-1.5%	-3.2%	-1.8%
Quintile	Market	Monthly	S&P 500	0.5%	0.1%	-3.4%	-2.4%	-3.5%	-2.7%
Quintile	Market	Daily	R1000	9.3%	7.8%	7.7%	7.4%	6.3%	11.2%
Quintile	Market	Quarterly	R1000	-0.1%	-0.5%	-3.4%	-2.3%	-3.3%	-2.6%

Notes: Average Monthly Return reports the raw self-financing return of the long-short portfolio after applying a flat one-way transaction cost of 10 bps to traded notional on each rebalance. The remaining columns report long-short alphas across the full factor-model set, including the six-factor specification, using the post-cost strategy returns. Bold entries indicate statistical significance at the 5% level.

Turnover Constraints

Turnover is best introduced in stages. The first question is descriptive: how much trading does the unconstrained strategy actually require on its own implemented backtest path? The second question is more operational: what happens to performance once that same strategy is forced to respect an explicit turnover budget at each rebalance? The third question is more path-dependent: how sensitive are turnover-constrained outcomes to the starting portfolio from which implementation begins? Those three steps move from measurement, to direct implementation control, to a broader assessment of how much the signal's realized results depend on the path taken into the target holdings.

Actual Realized Turnover

Table 3.11 reports realized turnover on the actual implemented path of the backtest. Turnover is summarized only across dates with positive trading activity, so the table reflects the size of actual implementation events instead of averaging in inactive periods. The **long-short spread** averages **86.5%** turnover per trading event, reaches a maximum of **128.4%**, and compounds to about **1037.5%** of annualized turnover. For reference, the top and bottom long-only buckets average **58.6%** and **38.0%** turnover per trading event, respectively, while the long-short spread records **10** trading events over the sample.

Table 3.11: Realized turnover on the actual implemented path

Metric	T	B	T-B
Trading Events	9	9	10
Average Turnover	58.6%	38.0%	86.5%
Median Turnover	62.4%	33.8%	94.4%
95th Percentile Turnover	78.3%	67.2%	126.5%
Maximum Turnover	81.2%	67.5%	128.4%
Annualized Turnover	703.7%	456.5%	1037.5%

Notes: Turnover is computed from realized backtest trades as the smaller of aggregate buys and aggregate sells divided by portfolio NAV. Summary statistics are based on dates with positive turnover only, so the table describes actual trading events rather than averaging in inactive periods. Annualized turnover is the sum of realized event turnover scaled by the elapsed sample length.

Turnover-Constrained Implementation

The next step is to move from describing turnover to constraining it directly inside the strategy implementation. The purpose of this test is to ask how much of the baseline signal survives once the portfolio is no longer allowed to rebalance fully into the desired target weights at every decision point. Instead, the strategy is forced to trade only up to a specified turnover budget, which means the implemented portfolio may move only part of the way toward the unconstrained target on any given rebalance date.

This matters because many signals look attractive when evaluated as if the portfolio can always jump frictionlessly into the desired holdings. In practice, however, real implementations often face turnover budgets driven by capacity, market impact, transaction costs, and product constraints. A turnover-constrained version of the strategy therefore provides a cleaner read on how quickly the

signal decays once implementation frictions are embedded directly into the portfolio path rather than applied only as an ex-post cost adjustment.

In the implementation tested here, the analysis stays with the long-short quintile spread used earlier in the paper. Each sleeve is initialized from the benchmark's market-cap weights on its first rebalance. After that initial step, the portfolio is re-optimized monthly inside each selected quintile to maximize directional signal exposure subject to a one-way turnover cap versus the previously implemented sleeve, while a soft quadratic concentration penalty discourages extreme weight concentration within the sleeve. Appendix B states the exact optimization problem.

Table 3.12 reports the long-short quintile strategy when each monthly rebalance is capped by a one-way turnover budget and each sleeve is initialized from market-cap benchmark weights on its first rebalance. After initialization, the long and short books maximize signal exposure within their selected quintiles subject to the turnover cap, while applying a soft quadratic concentration penalty that discourages extreme sleeve weights. Relative to the unconstrained baseline return of **0.7%**, the turnover-constrained implementations range from **0.4%** at the **5%** budget to **-1.0%** at the **50%** budget. The strongest average monthly return in the constrained set comes from the **5%** budget at **0.4%**. The strongest reported alpha appears under the **Market** model for the **5%** budget at **-0.1%**.

Table 3.12: Turnover-constrained implementation

Turnover Budget	Average Monthly Return	Market	3-Factor	4-Factor	5-Factor	6-Factor
5%	0.4%	-0.1%	-1.0%	-0.5%	-1.1%	-1.0%
10%	-0.3%	-0.6%	-2.0%	-1.3%	-2.2%	-1.9%
20%	-0.9%	-0.7%	-3.1%	-2.0%	-3.1%	-2.5%
50%	-1.0%	-0.5%	-4.5%	-3.5%	-4.8%	-4.1%

Notes: Each row reports the long-short quintile implementation when each sleeve is initialized with benchmark market-cap weights on its first rebalance and then re-optimized monthly to maximize directional signal exposure subject to a one-way turnover cap versus the currently implemented sleeve. Average Monthly Return reports the raw self-financing return of the long-short portfolio, and the remaining columns report long-short alphas across the full factor-model set, including the six-factor specification. Bold entries indicate statistical significance at the 5% level.

Path Dependency

Even a turnover-constrained implementation still depends on where the portfolio starts. A strategy that begins from yesterday's implemented holdings can face a very different transition problem from one that starts from a benchmark-like portfolio, an equal-weight portfolio, or some other admissible book. That is the motivation for the path-dependency framework: instead of evaluating turnover-constrained performance from a single starting path, the analysis evaluates many plausible starting portfolios and follows how each one evolves under the same turnover budget and target signal process.

In this report, the path-dependency exercise follows the same long-short quintile methodology as the turnover-constrained section above, but changes the starting point. Instead of initializing each sleeve by maximizing the signal, the analysis begins from **1000** different market-weighted random portfolios. Concretely, on the first rebalance each long and short sleeve is formed by ranking the investable universe on a random factor, selecting the top or bottom quintile exactly as in the baseline signal sort, and then assigning benchmark market-cap weights within that randomly selected sleeve. After that initial rebalance, the random factor disappears completely: every path follows the same monthly signal-maximizing implementation used in the turnover-constrained section above, with the usual target signal, the same long-short quintile structure, the same **10%** one-way turnover budget, and the same soft concentration penalty. The purpose is to test whether the measured outperformance is simply the consequence of beginning from a favorable path, or whether it survives across a broad range of admissible starting books. Looking across many such paths therefore gives a direct read on the dispersion of outcomes and the extent to which implementation performance is path-dependent rather than a one-off result of a lucky initial portfolio. Figure 3.4 shows how those random-start implementations evolve through time relative to the realized constrained path, while Figure 3.5 summarizes the cross-path distribution of average monthly returns.

Table 3.13 summarizes the same turnover-constrained long-short implementation under **1000** different random market-weighted starting portfolios, each run with the same monthly **10%** turnover budget used in the constrained implementation test above. Across those random starts, the long-short average monthly return has a mean of **-0.9%** and a median of **-0.9%**. The realized path from the main implementation section delivers **-0.3%**, which places it around the **80th** percentile of the random-start

return distribution and in the **9th** decile of that same distribution, as illustrated by the histogram in Figure 3.5. The largest median alpha across the path runs appears under the **Market** model at **-0.7%**.

Table 3.13: Path dependency under random market-weighted starts

Statistic	Average Monthly Return	Market	3-Factor	4-Factor	5-Factor	6-Factor
Realized 10% Path	-0.3%	-0.6%	-2.0%	-1.3%	-2.2%	-1.9%
Mean Random Path	-0.9%	-0.7%	-2.1%	-1.3%	-2.1%	-1.6%
Median Random Path	-0.9%	-0.7%	-2.1%	-1.3%	-2.2%	-1.7%
10th Percentile	-1.6%	-1.4%	-2.8%	-2.0%	-2.9%	-2.2%
90th Percentile	-0.1%	0.1%	-1.3%	-0.5%	-1.4%	-1.0%

Notes: The random-start rows summarize **1000** path-dependent long-short implementations. Each path begins from a random market-weighted initial quintile portfolio, then follows the same monthly signal-maximizing turnover-constrained implementation with a **10%** one-way turnover budget. The realized path row reports the single realized implementation path from the turnover-constrained section above.

Index Approach

An index-replication implementation asks a stricter question than the baseline long-short tests: can the signal still add value once the portfolio is forced to look like a realistic benchmark-aware product rather than an unconstrained signal sleeve? That distinction matters because a raw signal portfolio can pick up large incidental tilts, especially at the sector level, that would be hard to justify in an index context. A signal that appears strong only because it effectively concentrates the portfolio into one corner of the market is less informative than a signal that survives while the portfolio remains close to the benchmark it is meant to track.

For that reason, the index approach evaluates three progressively tighter constructions. The first is *Sector Name Matching*, which replicates the proportion of names selected across sectors in line with the underlying benchmark and then tilts the portfolio toward the signal while staying constrained within those sector buckets. Operationally, the portfolio takes the top decile within each sector, so the sector mix is matched through within-sector name selection rather than by matching final portfolio weights. The second is *Sector Exposure Matching*, which keeps the full benchmark universe investable rather than pre-filtering to sector winners, but constrains realized sector ETF exposures directly so that the tilt remains close to the benchmark in sector risk space. The third is *Exposure Matching*, which adds broader benchmark-alignment constraints on top of sector control by estimating stock-level proxy exposures and keeping the realized portfolio close to the benchmark in that richer exposure space. Concretely, each stock's returns are regressed over a trailing three-year window in two separate blocks: one regression on the broad market and style-factor proxies, and one regression on the sector ETF proxies. Those rolling regressions provide the stock-level loadings used by the optimizer. The portfolio is then tilted toward the signal while constraining benchmark-relative proxy exposures directly and limiting each constituent to a 1% active-weight move. In the current setup, each individual proxy exposure is required to remain within a ± 0.05 active loading band versus the benchmark. These are portfolio-level active proxy loadings, not security-level weights: the optimizer aggregates stock-level regression loadings into benchmark-relative portfolio exposures and constrains those aggregate deviations one proxy at a time. Appendix C states the corresponding optimization problem in compact form. In both cases, the portfolio is not allowed to chase the signal freely; instead, it is tilted toward the signal while remaining constrained to stay close to the benchmark. That makes the exercise a cleaner test of whether the signal itself adds value once the most obvious sources of benchmark-relative bias have been removed.

Sector Matching

Table 3.14 compares the two sector-only tilts against the benchmark, Table 3.15 reports the corresponding factor-model alphas, and Table 3.16 summarizes the latest sector weights. The first step, **Sector Name Matching**, keeps the portfolio benchmark-aware by matching the proportion of names selected in each sector, implemented as taking the top decile within each sector rather than reweighting the full universe. The second step, **Sector Exposure Matching**, keeps the whole benchmark universe investable and moves from sector-level name matching to direct sector exposure

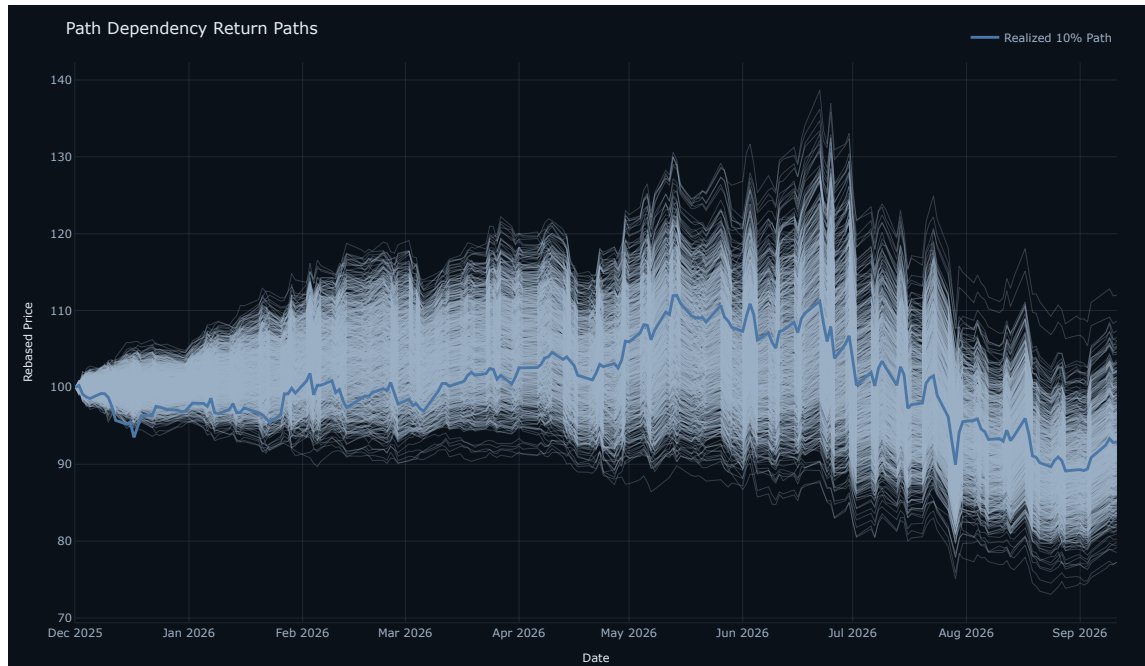


Figure 3.4: Random-start path dependency under a 10% turnover budget

control by constraining realized sector ETF loadings. Relative to **Russell 1000**, **Sector Name Matching** changes total return by **-9.1%** and average monthly excess return by **-0.7%**; **Sector Exposure Matching** changes total return by **-0.4%** and average monthly excess return by **0.1%**. Across factor models, the strongest alpha appears under the **Market** specification at **0.2%**. Taken together, these results show how the signal behaves as the sector control tightens from matching names by sector to matching sector exposures more directly before the analysis turns to the full factor-plus-sector exposure-matching setup.

Table 3.14: Sector-only index tilts versus benchmark

Metric	Sector Name Matching	Sector Exposure Matching	Russell 1000
Average Monthly Return	0.6%	1.4%	1.3%
Average Monthly Excess Return	0.3%	1.1%	1.0%
CAGR	3.9%	15.4%	15.8%
Sharpe Ratio	0.1	0.84	0.93
Max Drawdown	-16.0%	-7.6%	-9.1%
Total Return	3.0%	11.8%	12.1%

Notes: Each column reports the realized performance of the sector-only benchmark-aware tilts relative to the benchmark over the same sample. The first variant matches the number of names selected in each sector to the benchmark, while the second constrains realized sector ETF exposures directly.

Table 3.15: Sector-only tilt factor-model alphas

Strategy	Market	3-Factor	4-Factor	5-Factor	6-Factor
Sector Name Matching	-0.7%	-2.8%	-2.4%	-3.0%	-2.2%
Sector Exposure Matching	0.2%	-0.4%	-0.3%	-0.6%	-0.4%

Notes: Each cell reports the average monthly alpha for the indicated sector-only tilt under the corresponding factor model. The first variant matches the number of names selected in each sector to the benchmark, while the second constrains realized sector ETF exposures directly. Bold entries indicate statistical significance at the 5% level.

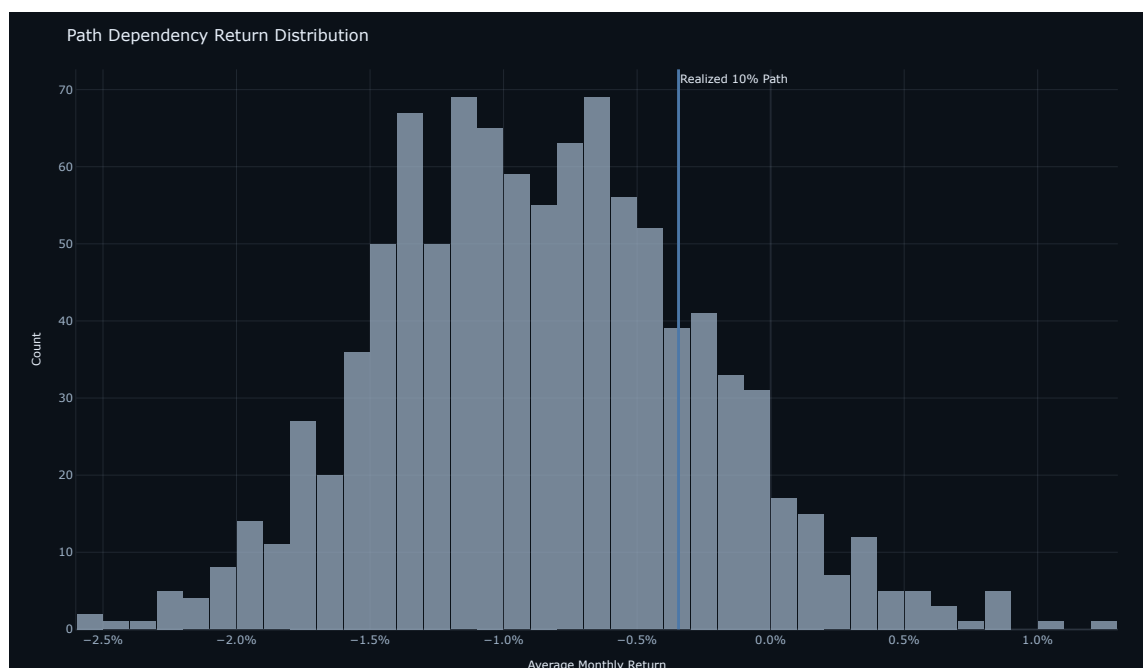


Figure 3.5: Distribution of average monthly returns across 1000 random starts under a 10% turnover budget

Table 3.16: Latest sector weights across sector-only tilts

Sector	Sector Name Matching	Sector Exposure Matching	Benchmark
Utilities	1.0%	0.8%	2.0%
Industrials	12.8%	9.2%	8.7%
Technology	36.3%	33.7%	37.1%
Real Estate	1.1%	1.1%	2.1%
Healthcare	6.3%	17.4%	9.4%
Consumer Cyclical	17.3%	7.1%	9.2%
Financial Services	11.1%	10.7%	12.2%
Consumer Defensive	8.0%	3.2%	4.3%
Energy	3.2%	7.0%	3.7%
Basic Materials	2.3%	2.8%	1.9%
Communication Services	0.6%	7.2%	9.4%

Notes: Weights are aggregated from security-level holdings to sectors on the latest date shared by the strategies and benchmark. The first variant matches the benchmark by the proportion of names selected per sector, while the second targets sector ETF exposure alignment directly.

Exposure Matching

Exposure matching asks a stricter implementation question than sector matching alone. Instead of only controlling the sector mix, it tries to keep the portfolio close to the benchmark across the broader proxy set that a realistic index product would care about. The purpose is to isolate the signal more cleanly: if the signal still adds value after the optimizer is forced to stay near the benchmark in realized exposure space, that is stronger evidence that the effect is not simply a repackaged market, style, or sector bet. The three tables below are intended to be read together. The first compares overall performance to the benchmark, the second shows whether return-based alpha survives under standard factor models, and the third reports the latest realized proxy exposures versus the benchmark. Appendix C gives the optimization problem used for this implementation test.

Table 3.17 compares the exposure-matched tilt against its benchmark, Table 3.18 reports the corresponding factor-model alphas, and Table 3.19 summarizes the latest realized proxy exposures. The methodology keeps the full benchmark universe active rather than pre-selecting a narrow signal

sleeve. At each rebalance, stock-level proxy loadings are estimated from trailing three-year time-series regressions of individual stock returns in two separate blocks: one against the broad market and style-factor proxies, and one against the sector ETF proxies. Those estimated loadings are then fed into a benchmark-relative optimizer that tilts toward the signal, enforces explicit proxy-level active exposure bounds of ± 0.05 around the benchmark, enforces long-only total weights, and limits each security to a 1% active-weight move. Relative to **Russell 1000**, **Exposure Matching** changes total return by **-0.6%** and average monthly excess return by **0.0%**. The largest realized proxy deviation appears in **RMW** with a maximum absolute active exposure of **0.05**. Across factor models, the strongest alpha appears under the **Market** specification at **0.2%**. Taken together, these results show whether the signal remains useful after explicitly keeping realized proxy exposures close to the benchmark rather than letting the portfolio absorb broad factor or sector tilts.

Table 3.17: Exposure-matched index tilt versus benchmark

Metric	Exposure Matching	Russell 1000
Average Monthly Return	1.4%	1.3%
Average Monthly Excess Return	1.1%	1.0%
CAGR	15.0%	15.8%
Sharpe Ratio	0.81	0.93
Max Drawdown	-8.5%	-9.1%
Total Return	11.5%	12.1%

Notes: Each column reports the realized performance of the exposure-matched tilt relative to the benchmark over the same sample. The portfolio uses a three-year rolling proxy regression, includes both factor and sector proxies, limits each security to a 1% active-weight deviation versus the benchmark, and constrains each portfolio-level active proxy loading to lie within ± 0.05 of the benchmark.

Table 3.18: Exposure-matched tilt factor-model alphas

Strategy	Market	3-Factor	4-Factor	5-Factor	6-Factor
Exposure Matching	0.2%	-0.3%	-0.3%	-0.5%	-0.4%

Notes: Each cell reports the average monthly alpha for the exposure-matched tilt under the indicated factor model. The portfolio uses a three-year rolling proxy regression, includes both factor and sector proxies, limits each security to a 1% active-weight deviation versus the benchmark, and constrains each portfolio-level active proxy loading to lie within ± 0.05 of the benchmark. Bold entries indicate statistical significance at the 5% level.

Table 3.19: Latest realized proxy exposures versus benchmark

Proxy	Exposure Matching	Benchmark	Active
MKT	1.108	1.158	-0.05
SMB	0.056	0.006	0.05
HML	-0.001	-0.047	0.046
MOM	-0.056	-0.006	-0.05
RMW	0.052	0.1	-0.048
CMA	-0.113	-0.163	0.05
COMM_SERVICES	0.127	0.162	-0.036
ENERGY	0.08	0.047	0.034
FINANCIALS	0.161	0.114	0.047
INDUSTRIALS	0.002	0.052	-0.05
TECHNOLOGY	0.382	0.395	-0.013
CONSUMER_STAPLES	-0.072	-0.022	-0.05
REAL_ESTATE	0.035	0.009	0.026
UTILITIES	-0.014	0.03	-0.044
HEALTH_CARE	0.156	0.106	0.05
CONSUMER_DISCRETIONARY	0.102	0.142	-0.039
MATERIALS	0.056	0.01	0.046

Notes: Reported values are taken from the latest rebalance-date snapshot of realized portfolio proxy exposures. Active exposures are strategy minus benchmark. Proxy loadings are estimated from separate three-year rolling factor and sector regressions.

CHAPTER 4

Comparison With Documented Anomalies

The earlier chapters evaluate the signal on its own terms: first as a cross-sectional ranking signal, then as an implemented portfolio subject to realistic construction and trading frictions. A different question is how those results compare with the broader set of documented anomaly strategies. That comparison matters because a signal can look attractive in isolation while still being fairly ordinary relative to the wider empirical asset-pricing literature. A cross-strategy benchmark therefore helps place the signal’s Sharpe ratio, compounded growth, and factor-adjusted alpha in context. Appendix A lists the anomaly signals in the current research set, a collection often referred to as the “zoo” in the empirical asset-pricing literature.

In the spirit of the comparison exercises used in the anomaly literature, this chapter is intended to benchmark the signal against a common reference set of published long-short anomaly strategies. Within the current workflow, the comparison is carried out at the level of each signal’s default T - B long-short strategy, so the chapter is intended to compare those baseline long-short implementations against one another. In the current setup, that means comparing the signal against **66** other anomaly strategies in that research set before placing it in the broader documented-anomaly context. The goal is not to reproduce every detail of the external source data inside this paper, but to show where the signal falls in the corresponding cross-strategy distributions once the benchmark dataset has been aligned to the same gross return convention used here.

4.1 Relationship To Known Anomalies

With so many documented anomalies, it is possible that any new cross-sectional predictor is capturing some combination of already known signals rather than introducing genuinely distinct predictive information. It is therefore natural to ask to what extent the proposed signal adds incremental content beyond the most closely related anomaly strategies. Signals that are more closely related should, in general, exhibit larger absolute correlations, so the first step is to examine where the signal sits in the cross-strategy correlation structure of the current anomaly set. Figure 4.1 displays the signal’s pairwise correlations with the other anomaly strategies in the research set, grouped into correlation bins for readability. Figure 4.2 is intended to complement that view with an agglomerative hierarchical cluster diagram, using Ward’s minimum method and a capped number of clusters, so that the signal can be located visually within the broader anomaly taxonomy. Correlation alone, however, does not establish whether related anomaly strategies fully span the signal’s implemented return. Figure 4.3 therefore is intended to report pairwise spanning tests in which the signal’s default monthly T - B return is regressed on each peer anomaly strategy one at a time:

$$r_{\text{signal},t}^{\text{T-B}} = \alpha + \beta r_{\text{peer},t}^{\text{T-B}} + \varepsilon_t.$$

In this setup, a positive alpha means that the signal continues to deliver incremental long-short return even after the return of the peer anomaly strategy has been taken into account. The resulting scatter is intended to plot residual alpha against the corresponding alpha t-statistic, so that both economic magnitude and statistical strength can be compared across the current anomaly set.

Across the current anomaly set, the signal lands in the upper-right quadrant in **52** out of **66** pairwise tests, with a median residual alpha of **0.5%** and a median alpha t-statistic of **0.25**. The residual alpha ranges from **-0.8%** to **3.9%** across peer strategies.

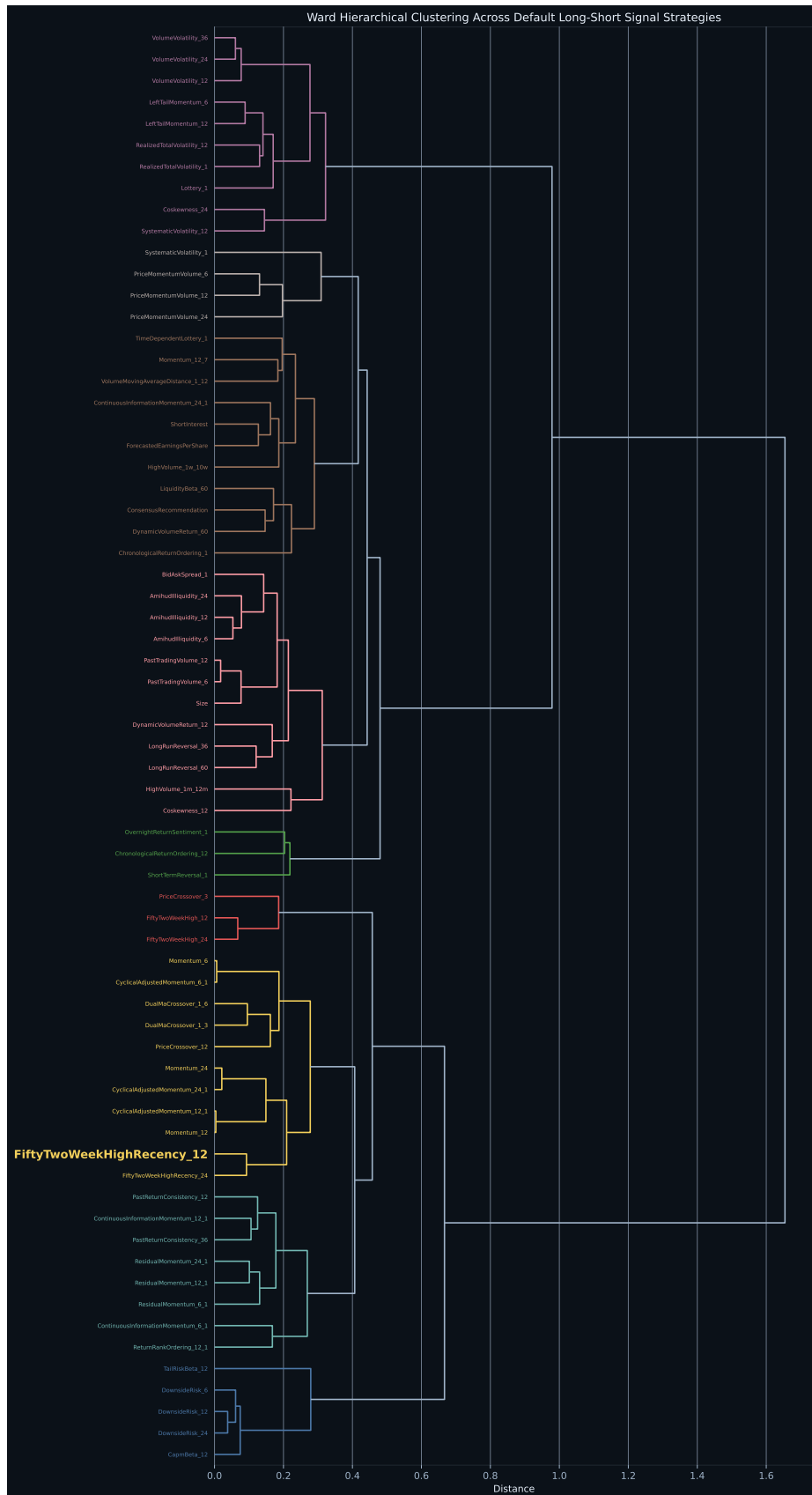


Figure 4.2: Hierarchical clustering of the current anomaly set

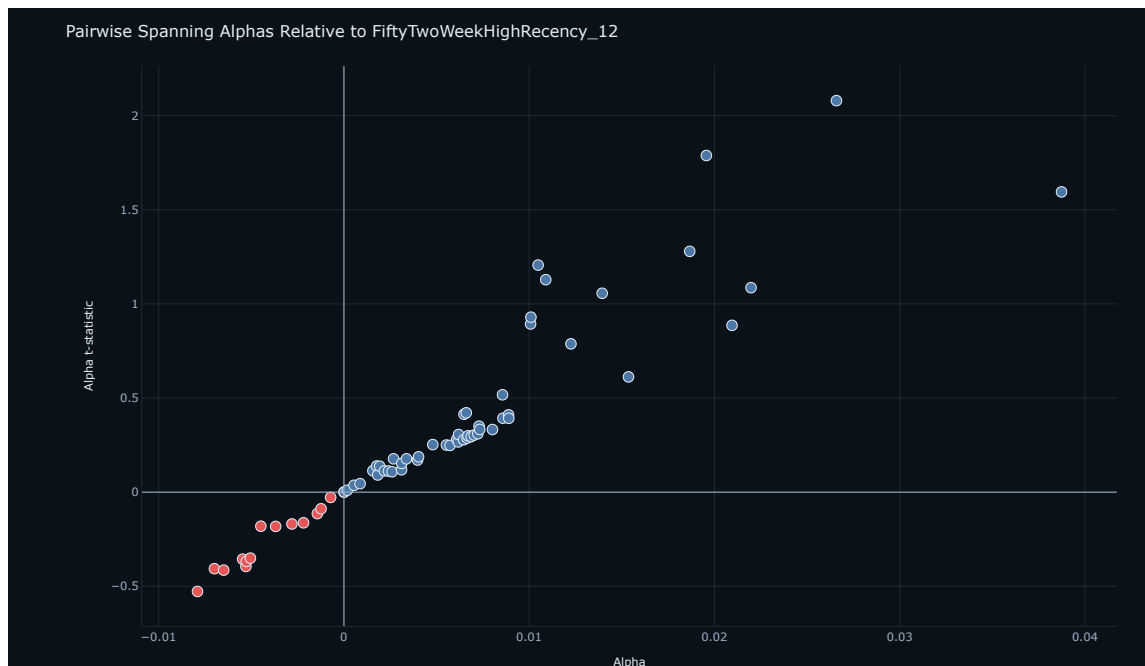


Figure 4.3: Pairwise spanning-test alphas and alpha t-statistics relative to peer anomaly strategies

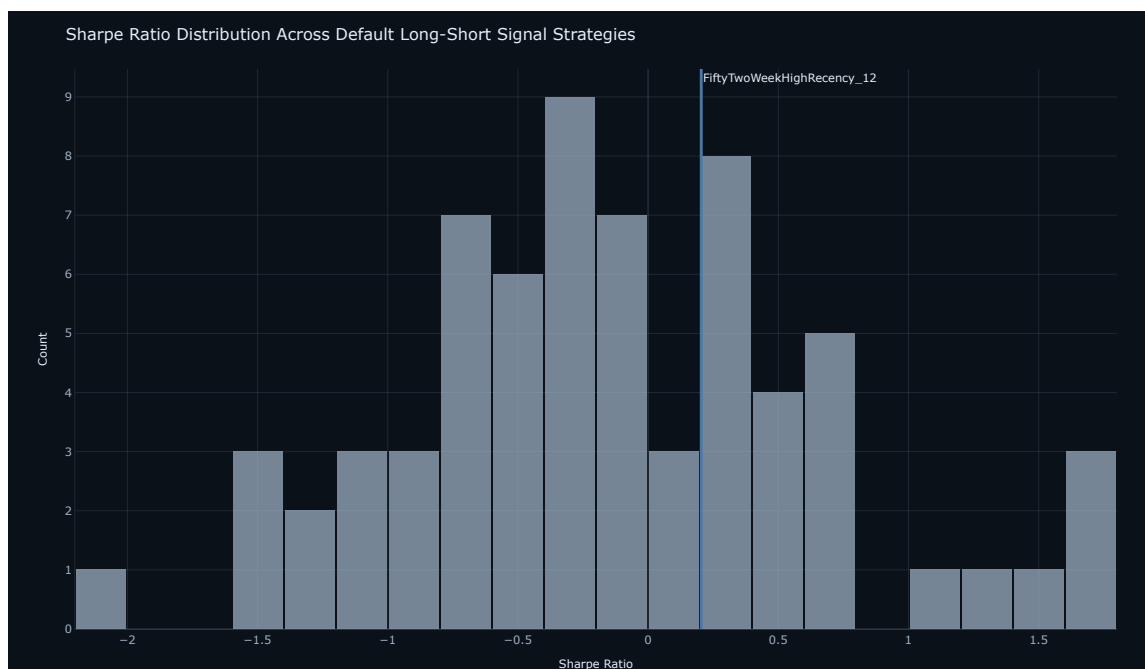


Figure 4.4: Signal Sharpe ratio relative to the reference anomaly distribution

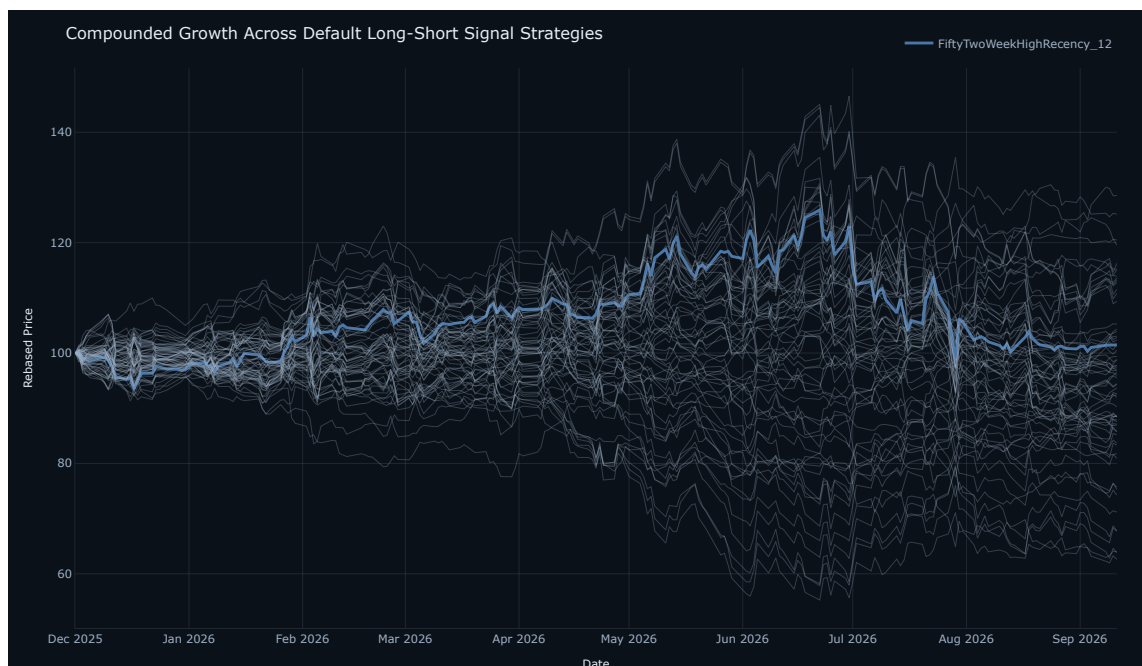


Figure 4.5: Compounded growth of the signal relative to the reference anomaly distribution

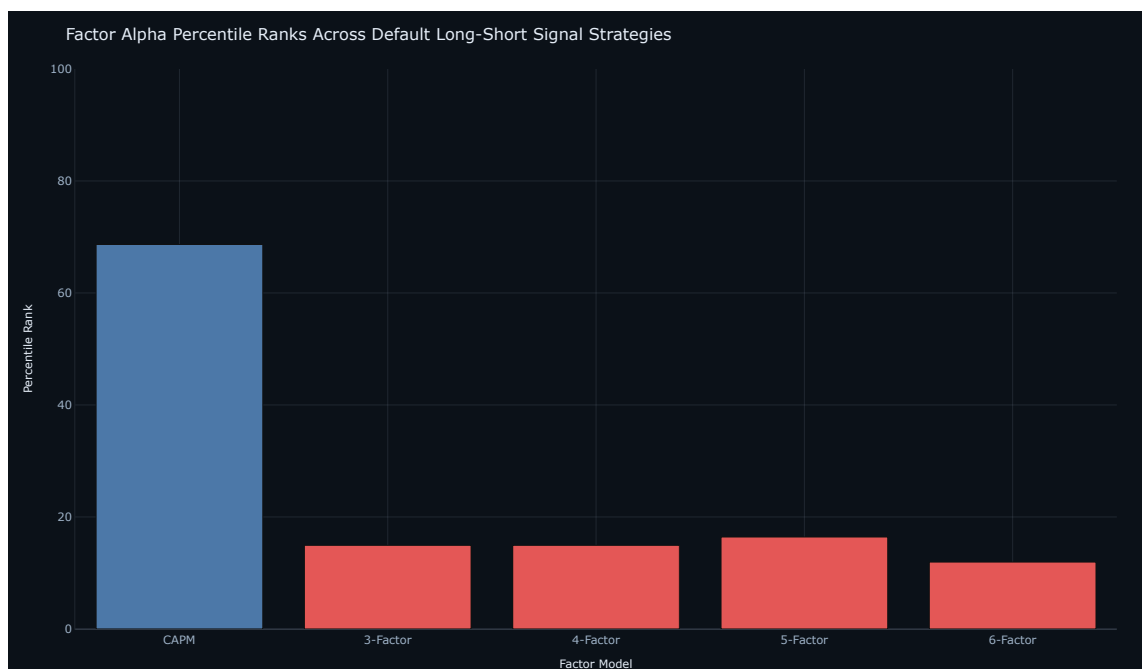


Figure 4.6: Signal alpha percentile ranks relative to the reference anomaly distribution

Appendices

APPENDIX A

Signal Catalog

Table A.1: Default signal strategies in the current research set

Signal	Family	Parameter	Ranking Direction
Momentum_6	Momentum	6M	Higher is better
Momentum_12	Momentum	12M	Higher is better
Momentum_12_7	Momentum	6M	Higher is better
Momentum_24	Momentum	24M	Higher is better
CyclicalAdjustedMomentum_6_1	Cyclical Adjusted Momentum	6M	Higher is better
CyclicalAdjustedMomentum_12_1	Cyclical Adjusted Momentum	12M	Higher is better
CyclicalAdjustedMomentum_24_1	Cyclical Adjusted Momentum	24M	Higher is better
PastReturnConsistency_12	PastReturnConsistency	12M	Higher is better
PastReturnConsistency_36	PastReturnConsistency	36M	Higher is better
ContinuousInformationMomentum_6_1	Continuous Information Momentum	6M	Higher is better
ContinuousInformationMomentum_12_1	Continuous Information Momentum	12M	Higher is better
ContinuousInformationMomentum_24_1	Continuous Information Momentum	24M	Higher is better
ResidualMomentum_6_1	Residual Momentum	6M	Higher is better
ResidualMomentum_12_1	Residual Momentum	12M	Higher is better
ResidualMomentum_24_1	Residual Momentum	24M	Higher is better
LeftTailMomentum_6	LeftTailMomentum	6M	Higher is better
LeftTailMomentum_12	LeftTailMomentum	12M	Higher is better
ChronologicalReturnOrdering_1	Return Ordering	1M	Higher is better
ChronologicalReturnOrdering_12	Return Ordering	12M	Higher is better
ReturnRankOrdering_12_1	Return Rank Ordering	12M	Higher is better
Lottery_1	Lottery	1M	Lower is better
TimeDependentLottery_1	Time-Dependent Lottery	1M	Lower is better
OvernightReturnSentiment_1	Overnight Return Sentiment	1M	Lower is better
Size	Size		Lower is better
ShortInterest	Short Interest		Lower is better
ConsensusRecommendation	Consensus Recommendation		Lower is better
ForecastedEarningsPerShare	Forecasted Earnings Per Share		Higher is better
FiftyTwoWeekHigh_12	52-Week High	12M	Higher is better
FiftyTwoWeekHigh_24	52-Week High	24M	Higher is better
FiftyTwoWeekHighRecency_12	52-Week High Recency	12M	Higher is better
FiftyTwoWeekHighRecency_24	52-Week High Recency	24M	Higher is better
PriceMomentumVolume_6	Price Momentum + Volume	6M	Higher is better
PriceMomentumVolume_12	Price Momentum + Volume	12M	Higher is better
PriceMomentumVolume_24	Price Momentum + Volume	24M	Higher is better
DynamicVolumeReturn_12	DynamicVolumeReturn	12M	Higher is better
DynamicVolumeReturn_60	DynamicVolumeReturn	60M	Higher is better
HighVolume_1w_10w	High Volume	1W/10W	Higher is better
HighVolume_1m_12m	High Volume	1M/1Y	Higher is better
PastTradingVolume_6	Past Trading Volume	6M	Lower is better
PastTradingVolume_12	Past Trading Volume	12M	Lower is better

Table A.1 continued from previous page

Signal	Family	Parameter	Ranking Direction
VolumeMovingAverageDistance_1_12	Volume MA Distance	1M/12M	Lower is better
LongRunReversal_36	Long-Run Reversal	36M	Lower is better
LongRunReversal_60	Long-Run Reversal	60M	Lower is better
PriceCrossover_3	Price Crossover	3M	Higher is better
PriceCrossover_12	Price Crossover	12M	Higher is better
DualMaCrossover_1_3	Dual MA Crossover	1M/3M	Higher is better
DualMaCrossover_1_6	Dual MA Crossover	1M/6M	Higher is better
DownsideRisk_6	Downside Risk	6M	Higher is better
DownsideRisk_12	Downside Risk	12M	Higher is better
DownsideRisk_24	Downside Risk	24M	Higher is better
CapmBeta_12	CAPM Beta	12M	Higher is better
Coskewness_12	Coskewness	12M	Lower is better
Coskewness_24	Coskewness	24M	Lower is better
SystematicVolatility_1	SystematicVolatility	1M	Lower is better
SystematicVolatility_12	SystematicVolatility	12M	Lower is better
RealizedTotalVolatility_1	Realized Total Volatility	1M	Lower is better
RealizedTotalVolatility_12	Realized Total Volatility	12M	Lower is better
TailRiskBeta_12	Tail Risk Beta	12M	Higher is better
AmihudIlliquidity_6	Amihud Illiquidity	6M	Higher is better
AmihudIlliquidity_12	Amihud Illiquidity	12M	Higher is better
AmihudIlliquidity_24	Amihud Illiquidity	24M	Higher is better
BidAskSpread_1	Bid-Ask Spread	1M	Higher is better
LiquidityBeta_60	Liquidity Beta	60M	Higher is better
VolumeVolatility_12	Volume Volatility	12M	Lower is better
VolumeVolatility_24	Volume Volatility	24M	Lower is better
VolumeVolatility_36	Volume Volatility	36M	Lower is better
ShortTermReversal_1	Momentum	1M	Lower is better

Notes: This appendix lists the default signal strategies in the current research set. Parameter reports the main lookback in months or the short/long EWMA half-life pair where applicable. Ranking Direction indicates whether larger or smaller signal values are treated as more favorable in the portfolio sort.

APPENDIX B

Turnover-Constrained Optimizer

This appendix states the sleeve-level optimization problem used in the turnover-constrained implementation exercise. The optimizer is applied separately to the long and short sleeves, with the sign of the signal flipped across sleeves so that each sleeve remains a long-only portfolio over its own selected universe.

For one sleeve at one rebalance date, let w^{prev} denote the currently implemented sleeve weights, let s denote the aligned cross-sectional signal vector, and let w denote the new implementable sleeve weights. The optimization problem is

$$\max_w \quad s^\top w - \kappa \sum_i w_i^2 \quad (\text{B.1})$$

subject to

$$\sum_i w_i = 1, \quad (\text{B.2})$$

$$w_i \geq 0 \quad \forall i, \quad (\text{B.3})$$

$$w_i = 0 \quad \text{for securities outside the eligible sleeve universe}, \quad (\text{B.4})$$

$$\frac{1}{2} \sum_i |w_i - w_i^{\text{prev}}| \leq \tau, \quad (\text{B.5})$$

where τ is the one-way turnover budget for the rebalance and $\kappa \geq 0$ is the soft concentration-penalty coefficient. The quadratic term discourages the optimizer from reallocating too much sleeve weight into a small number of names when several eligible portfolios deliver similar signal exposure.

The implementation also allows forced exits when a currently held security no longer has a usable signal value at the rebalance date. Let

$$\phi = \sum_{i: w_i^{\text{prev}} > 0, s_i \text{ unusable}} w_i^{\text{prev}} \quad (\text{B.6})$$

denote the sleeve weight that must be liquidated for data-availability reasons. The effective turnover budget for that rebalance is then

$$\tau^* = \tau + \phi. \quad (\text{B.7})$$

The optimizer is therefore solved with the same objective and portfolio constraints as above, but with τ^* replacing τ whenever forced exits are required.

APPENDIX C

Exposure-Matching Optimizer

This appendix states the benchmark-relative optimization problem used in the exposure-matching implementation exercise. The goal is to tilt the portfolio toward the cross-sectional signal while keeping the resulting portfolio close to the benchmark in both weights and realized proxy exposures.

For one rebalance date, let w^{bench} denote the benchmark weights over the eligible universe, let x denote the vector of active weights relative to that benchmark, and let s denote the aligned signal vector. The implemented portfolio is therefore

$$w = w^{\text{bench}} + x. \quad (\text{C.1})$$

The optimizer solves

$$\max_x s^\top x \quad (\text{C.2})$$

subject to

$$\sum_i x_i = 0, \quad (\text{C.3})$$

$$-\delta \leq x_i \leq \delta \quad \forall i, \quad (\text{C.4})$$

$$w_i^{\text{bench}} + x_i \geq 0 \quad \forall i. \quad (\text{C.5})$$

The first constraint enforces benchmark-relative neutrality so that the active weights sum to zero. The box bound δ limits how far any individual name can move away from its benchmark weight, while the non-negativity condition keeps the implemented portfolio long-only.

To control benchmark-relative risk, the optimizer also constrains active proxy exposures. Let $B^{(f)}$ denote the stock-level matrix of broad market and style proxy loadings, and let $B^{(s)}$ denote the stock-level matrix of sector proxy loadings. The active exposure vectors are

$$e^{(f)} = \left(B^{(f)}\right)^\top x, \quad (\text{C.6})$$

$$e^{(s)} = \left(B^{(s)}\right)^\top x. \quad (\text{C.7})$$

The implementation applies coordinate-wise bounds to those active loadings:

$$-\gamma \leq e_j^{(f)} \leq \gamma \quad \forall j, \quad (\text{C.8})$$

$$-\gamma \leq e_k^{(s)} \leq \gamma \quad \forall k, \quad (\text{C.9})$$

where γ is the allowable benchmark-relative exposure deviation per proxy. In the current paper setup, $\delta = 0.01$ and $\gamma = 0.05$.

The practical interpretation is straightforward. The signal can only improve the portfolio by reallocating weight within a tight benchmark-relative budget. It cannot create large stock-level overweights, and it cannot express large market, style, or sector bets through the proxy system. As a result, any remaining performance is better interpreted as signal-specific stock selection rather than incidental benchmark-relative exposure drift.

Bibliography

- [BH13] Bhootra, A. and Hur, J. ‘The Timing of 52-Week High Price and Momentum’. In: *Journal of Banking and Finance* vol. 37, no. 10 (2013), pp. 3773–3782.