
WORKING PAPER — VERSION 1.1

Do Durable Market Power Characteristics Predict Future Fundamentals and Stock Returns?

A preregistered public-data study of U.S. equities, 2017–2025

Hesham Mashhour

Independent Researcher

Publisher: Monopoly MOAT

Canonical research page: <https://mplymoat.com/research/durable-market-power-stock-returns/>

28 August 2026

Preregistered analysis. Version 1 and the approved version-2 addendum are bound to osf.io/fvgsr. Version 2 changed only the treatment of the outcome-blind coverage threshold.

Abstract

This study asks whether public, point-in-time characteristics associated with pricing power, operating durability, and public-company sales dominance predict future operating performance and stock returns. I construct a preregistered Durable Market Power Score (DMPS) from five SEC filing-based inputs and public-company sales share, then test it in a historically reconstructed U.S. equity universe. The formation universe contains 7,152 issuer-years from June 2017 through June 2025; after sector exclusions, 2,573 of 6,137 issuer-years have every score input. Complete observations represent 64.26% of eligible capitalization, a material selection limitation disclosed before outcomes were analyzed.

DMPS does not validate as a predictor of the two registered one-year-ahead operating outcomes. Its coefficient is -0.0150 for next-year gross margin (Holm-adjusted $p = 0.910$) and 0.0105 for next-year operating profitability (Holm-adjusted $p = 0.881$); wild year-cluster bootstrap tests reach the same conclusion. The value-weighted high-minus-low portfolio has a positive monthly six-factor alpha of 0.855%, or 10.26% annualized, but its 95% confidence interval is -2.77% to 23.29% and $p = 0.121$. The registered secondary Fama–MacBeth estimate is also positive—0.665% per month, or 7.98% annualized—but imprecise ($p = 0.180$). Joint Holm adjustment for the two return tests gives $p = 0.243$.

The results therefore do not support describing DMPS as a validated durable operating-performance construct, and they do not establish abnormal returns at the preregistered 5% level. They do leave an economically large but weakly identified positive return association for future study. The design is predictive and noncausal; public sales share is not a product-market definition or a direct markup measure.

Keywords: market power; operating profitability; gross margin; asset pricing; factor models; preregistration; public data

JEL classifications: G11; G12; G14; L11

Bottom line: the operating-validation tests are null, while the return point estimates are positive but too imprecise to distinguish a modest negative association from a large positive one.

1. Introduction

Some firms appear able to sustain unusually high margins, stable profitability, and dominant positions within their public-company peer groups. Those attributes are economically interesting for two different reasons. First, they may describe durable businesses whose operating outcomes persist. Second, if investors underappreciate that durability—or require compensation for risks correlated with it—the same attributes may predict subsequent stock returns. The two propositions are not equivalent. A durable business can be fully valued, and a profitable return strategy can reflect omitted risk rather than mispricing.

This paper tests both propositions with one fixed, interpretable composite: the Durable Market Power Score (DMPS). The score combines gross-margin level and stability, operating-profitability level and stability, and the issuer's share of revenue among eligible public companies in its broad FF12 industry. Every accounting input is selected from the latest original annual filing public by a May 31 information cutoff. The score is

formed in June and used to predict the next consecutive fiscal-year outcome and returns from July onward. No future filing, current ticker snapshot treated as historical fact, provider-adjusted return field, or outcome-trained weight enters the score.

The project is unusual in placing historical-universe and terminal-event reconstruction at the center of a public-data asset-pricing design. Market data services that are adequate for surviving tickers do not automatically supply a point-in-time security master, stable class identifiers, or delisting returns. The study therefore begins with contemporaneous SEC Official List of Section 13(f) Securities frames, maps those frames to SEC CIKs and dated ticker/exchange routes, resolves primary classes, and applies preregistered price and liquidity screens. It separately adjudicates security disappearances and successor routes from primary documents. This produces 7,152 formation-eligible issuer-years and a governed economic route for every scheduled holding window.

The analysis was registered in two transparent stages. Version 1 required complete score inputs to cover at least 70% of eligible capitalization. An outcome-blind measurement run found 64.26%, so version 1 correctly stopped and left all hypotheses untested. Before any DMPS value, research return, future-outcome association, or predictive coefficient was computed or inspected, a version-2 addendum was approved on the original OSF record. It kept the score, sample construction, hypotheses, estimands, and inference rules fixed but treated the observed coverage shortfall as a reported selection limitation instead of an arbitrary stopping rule. This paper reports that bounded version-2 execution.

The evidence is disciplined but not dramatic. For H1, neither co-primary operating outcome has a statistically detectable DMPS coefficient. A 90th-to-10th percentile DMPS difference implies a fitted -2.00 percentage-point change in next-year gross margin and a +1.39 percentage-point change in next-year operating profitability. Both intervals include zero, both Holm-adjusted p -values are far above 0.05, and the prespecified wild-cluster sensitivity agrees. By the registered decision rule, DMPS is not validated as a durable operating-performance construct in this sample.

For H2, the high-minus-low portfolio earns a monthly intercept of 0.855% after the Fama–French five factors and momentum, equivalent to 10.26 percentage points per year. That magnitude is economically substantial, but its annualized 95% interval ranges from -2.77% to 23.29% and its p -value is 0.121. H3 points in the same direction: the mean monthly Fama–MacBeth DMPS slope is 0.665%, or 7.98% annualized, with a 95% interval from -3.76% to 19.71% and $p = 0.180$. These estimates are compatible with a meaningful positive association, but they do not meet the registered rejection threshold and do not warrant a return-predictability claim.

The paper contributes in four ways. It offers a transparent accounting-based characteristic that does not pretend to be a structural markup. It tests operating validation before interpreting returns. It documents the selection induced by demanding multi-year point-in-time histories. And it supplies a reproducible public-data workflow that treats identity changes, corporate actions, successor securities, terminal cash, and missing prices as first-class measurement problems rather than silently setting disappearances to zero or dropping them.

2. Economic Motivation and Related Literature

2.1 Profitability, quality, and durability

Accounting profitability predicts returns in several influential designs, but the precise construct matters. Novy-Marx (2013) shows that gross profitability has distinct asset-pricing content. The present score uses gross margin rather than gross profit scaled by assets: it asks whether an issuer retains a larger and more stable fraction of revenue after cost of goods sold. This is closer to a pricing-power signal, but it can also reflect product mix, accounting classification, outsourcing, or business-model differences. It should not be interpreted as a markup without a production-function model.

Asness, Frazzini, and Pedersen (2019) organize quality around profitability, growth, safety, and payout. DMPS is narrower. It combines pricing-power and operating-durability summaries with a public-company dominance proxy, excludes valuation from the score, and controls for standard characteristics in the cross-sectional return test. The separation is intentional: a high-quality company need not be cheaply priced, and a score that predicts good operations does not necessarily predict abnormal returns.

Fama and French (2015) add profitability and investment factors to the conventional market, size, and value model. The H2 regression therefore includes all five factors plus momentum. Carhart (1997) motivates the momentum extension, but momentum is not part of DMPS. H3 further conditions security-level returns on book-to-market, preformation momentum, market beta, operating profitability, investment, size, and FF12 indicators. These controls cannot eliminate every omitted-risk explanation, but they make clear that the registered return claim is incremental to common public-data characteristics rather than a raw performance comparison.

2.2 Market power, concentration, and the limits of accounting proxies

The modern market-power literature emphasizes that markups require assumptions and inputs beyond simple margins. De Loecker, Eeckhout, and Unger (2020) estimate markups from production relationships; their framework is precisely why this study avoids calling an accounting ratio a direct markup. Gross margin can move with input classification and product composition. Operating profitability can reflect efficiency, capital intensity, or intangible investment. Neither identifies price relative to marginal cost on its own.

Public-company concentration is also narrower than economic market structure. Grullon, Larkin, and Michaely (2019) document rising concentration among listed firms, while Autor et al. (2017) connect superstar firms to changing industry outcomes. Those arguments motivate asking whether issuer-level profitability and public-company sales share contain persistent information. Yet DMPS's `public_sales_share` denominator omits private businesses, imports, and product-level geographic markets. FF12 is deliberately broad. The measure is best read as public-peer revenue dominance, not as an antitrust market share, Herfindahl index, switching-cost estimate, network-effect measure, or regulatory moat rating.

2.3 Four competing interpretations

The design keeps four explanations separate.

1. **Durable operating performance.** DMPS predicts future gross margin or operating profitability after conditioning on the current outcome.
2. **Full capitalization.** Investors correctly price that durability, leaving no positive abnormal return.

3. **Underreaction.** Investors incorporate the information slowly, producing positive subsequent alpha.
4. **Omitted risk.** A positive spread compensates investors for disruption, regulation, duration, concentration, intangible-capital, or other risks absent from the factor model.

H1 addresses the first proposition. H2 and H3 address predictive return associations after different controls. None alone distinguishes underreaction from omitted risk, and the observational design cannot support a causal statement that market power causes returns.

2.4 Inference and attrition

H1 follows the multiway-clustering logic of Cameron, Gelbach, and Miller (2011), with issuer and formation-year clusters. Because there are only eight H1 formation-year clusters, the preregistration also requires a null-imposed Rademacher wild bootstrap by year. H2 and H3 use the Newey and West (1987) heteroskedasticity-and-autocorrelation-consistent covariance with 12 lags. H3 follows Fama and MacBeth (1973), and Holm's (1979) sequential procedure controls the prespecified test families.

Attrition is not a cosmetic concern. Shumway (1997) shows how omitting delisting outcomes can bias return evidence. This project does not import a CRSP-specific delisting imputation. It uses observed primary transaction terms, contractual conversion ratios, successor prices, and frozen contingent-value conventions. An unresolved economic route remains missing; it is never assigned zero, -100%, or a convenient adverse return.

3. Preregistration and Evidentiary Sequence

3.1 What was known before registration

The registration is an analysis-plan registration after data acquisition, not a pre-data-collection registration. Before the outcome lock, the project had inspected listing identities, exchange mappings, security classes, formation price and liquidity screens, raw daily prices, corporate-action fields, primary transaction documents, successor routes, terminal-event contracts, and a fixed 50-row return-input validation sample. A small SEC sample was used to validate point-in-time fact-selection code.

Before either registered predictive execution, the project had not computed or inspected DMPS values for the research universe, score portfolios, research monthly returns, future-fundamental associations, score-to-return associations, or H1–H3 coefficients, test statistics, intervals, or p -values. Thus the design was not pre-data, but its predictive choices were frozen before the outcomes to which those choices were applied.

3.2 Version 1 and the coverage gate

The original plan required all score inputs to cover at least 70% of eligible issuer-years by usable formation capitalization. The full outcome-blind measurement run covered all 7,152 formation rows and selected 228,278 original annual facts with zero look-ahead violations. It found complete inputs for 2,573 issuer-years and usable complete-score capitalization of \$70.4446 trillion out of \$109.6300 trillion, or 64.2567%. Version 1 therefore failed F0. Its H1, H2, and H3 were untested, not null, and no return panel or predictive model was permitted.

3.3 Registered version 2

The version-2 addendum was approved on the original public OSF record on 27 August 2026 at 21:07:26 UTC. It made one substantive change: the already observed 64.26% coverage became a disclosed complete-case limitation rather than a stopping threshold. The score formula, point-in-time rules, portfolio construction, terminal treatments, hypotheses, estimands, sample dates, model controls, covariance estimators, lag lengths, bootstrap seed, and multiplicity rules remained fixed. The repository lock binds the approval, source commit, bundle, and manifest hashes.

The version-2 execution also preserved two failed attempts caused by implementation defects. Run 1 encountered a null provider raw-close cell and stopped because the constructor treated that ordinary missing observation as fatal. The correction skips it as missing, never as a zero return or endpoint, while retaining the intervening-month invalidation and fatal checks on governed action dates. Run 2 later misclassified three `checkpoint.json` control files as provider responses and stopped on an apparent duplicate request hash. The correction indexes only 64-hex response filenames, collapses one byte-identical duplicate, and remains fatal on conflicting duplicates. Neither failed run estimated a primary model or exposed a primary result. Run 3, from a separately committed source tree, is the sole authoritative run.

Table A2. Preserved version-2 execution history

Run	Status	Stopped/completed (UTC)	Reason/outcome	Result status
1	Failed, preserved	2026-08-27T22:24:27Z	AnalysisReturnError	No primary estimate computed or inspected
2	Failed, preserved	2026-08-27T22:30:05Z	PreformationControlError	No primary estimate computed or inspected
3	Complete, authoritative	2026-08-27T22:43:09.699493Z	All registered models completed	Primary results reported in this paper

Note. Runs 1 and 2 stopped on outcome-independent implementation errors before a primary model estimate was computed or inspected. Run 3 is the sole authoritative result run.

4. Data and Historical Sample

4.1 Point-in-time security universe

The historical market frame comes from contemporaneous Q2 SEC Official Lists of Section 13(f) Securities for 2017–2025. These lists establish dated security presence but do not provide a complete CIK, ticker, or specific-exchange history. The project maps each row through frozen SEC identity evidence, dated symbol/exchange routes, and class-specific filings. CIK is the stable issuer key; an internal hash represents the dated security route. Current ticker lists can route requests or cross-check identity but cannot overwrite historical identity.

A security enters a June formation only when it is verified as common equity of a domestic operating company on an eligible U.S. primary listing, has a raw formation close of at least \$5, has at least \$1 million of median prior-60-trading-day dollar volume in 2025 dollars, has at least 120 valid daily observations in the prior 366 calendar days, and is not stale through unchanged closes combined with zero or missing volume. ADRs, funds, ETFs, preferred shares, warrants, units, and bonds are excluded. When an issuer has multiple eligible common classes, the class with the highest frozen 60-day dollar volume is the primary security. Financial companies, utilities, and REITs do not enter the primary score.

This process yields 7,152 formation issuer-years across nine June formations and 1,425 CIKs. The primary-score sector rule retains 6,137. Each formation uses information public by May 31, forms portfolios on the last trading day of June, and begins the holding interval in July. Holdings normally continue through the next June, with the 2025 formation truncated at 31 December 2025.

4.2 Fundamentals

SEC EDGAR Submissions and XBRL Company Facts provide the accounting inputs. For every May 31 cutoff, only annual 10-K or 10-K/A facts accepted by that date are eligible. Flow facts must span 300–430 days; instant facts must match the fiscal-period end. The frozen precedence favors consolidated, non-dimensional USD or share facts and retains taxonomy, concept, unit, period, filing timestamp, accession, and source hash. Later comparative restatements never move backward into an earlier formation.

Historical SIC is assigned from filing-date SEC Financial Statement Data Set evidence across all 66 quarterly archives from 2009 Q1 through 2025 Q2, then mapped through the fixed FF12 definition. An input can be standardized only when its formation-year FF12 cell has at least 20 usable issuers and nonzero variance. This requirement is applied separately to each input rather than to an already complete intersection.

4.3 Prices, corporate actions, and terminal events

Marketstack raw end-of-day data are the main price source. Provider-adjusted closes are forbidden because their semantics did not meet the project's frozen action-consistency threshold. Instead, ordinary daily gross returns use raw-close ratios, ordinary cash-dividend days add a documented distribution to the current raw close, and nonunit corporate actions use an exact-key frozen ledger. The implementation reproduces all 5 ordinary baselines, 20 dividends, and 25 nonunit actions in the fixed validation sample.

The full historical reconstruction contains 1,687,860 daily observations and a governed economic route for each of the 7,152 holding windows. Across the full universe, primary documents identify 74 actual terminal events, one nonterminal continuation, and three independent exact price repairs. In the eventual 2,573-row score sample, the return constructor uses all 42 required action transitions, seven fixed-cash exits, and one independent exact-price route. It skips one invalid ordinary provider observation and invalidates one cross-month boundary; 28,789 of 28,791 monthly security-return rows remain usable. No Marketstack adjusted field is used.

4.4 Factor and control data

Monthly U.S. Fama–French five factors and momentum come from the Kenneth French Data Library. Raw archives are frozen by retrieval timestamp and SHA-256, and published percentage units are converted to decimals once. The factor interval contains all 102 required months from July 2017 through December 2025.

H3's preformation controls are also point-in-time. Book-to-market uses the latest positive book equity public by May 31 divided by formation capitalization. Momentum compounds months $t - 12$ through $t - 2$. Market beta is estimated over the prior 36 months ending one month before formation, with at least 24 observations. Operating profitability and investment use frozen SEC facts. Of 2,573 score-complete issuer-years, 2,343 have momentum and 2,081 have the required beta; the full monthly model begins only when at least 50 complete observations and full design-matrix rank are available.

5. Durable Market Power Score

For issuer i and the latest eligible fiscal year y known at formation t , the raw accounting measures are:

```
gross_margin[i,y] = gross_profit[i,y] / revenue[i,y]

operating_profitability[i,y] =
    operating_income[i,y] / average(total_assets[i,y], total_assets[i,y-1])

public_sales_share[i,t] =
    revenue[i,y(i,t)] /
    sum(revenue[j,y(j,t)] for eligible public issuers j in the same FF12 industry)
```

Revenue and average assets must be positive. Public sales share is defined only when issuers with usable revenue represent at least 70% of observed eligible-industry capitalization and the industry has at least 20 eligible issuers. It is a public-company denominator, not a product market.

Using up to five qualifying fiscal years known before formation, the score summarizes:

```
gm_level      = latest gross_margin
gm_stability   = -sample_sd(gross_margin over the prior five fiscal years)
op_level      = mean(operating_profitability over the prior five fiscal years)
op_stability   = -sample_sd(operating_profitability over the prior five fiscal years)
dominance     = log(public_sales_share)
```

At least four of five annual observations are required for every historical mean or standard deviation. Within each formation year, the five raw summaries are winsorized at the 1st and 99th percentiles across eligible issuers and standardized within FF12. The components and final score are:

```
pricing_power      = 0.5*z(gm_level) + 0.5*z(gm_stability)
economic_durability = 0.5*z(op_level) + 0.5*z(op_stability)
public_dominance   = z(dominance)
DMPS              = (pricing_power + economic_durability + public_dominance) / 3
```

All three components are required; there is no imputation, PCA, learned weight, or future cross-section. In the complete sample, mean DMPS is 0.0465 and its pooled standard deviation is 0.5768. Its correlations with pricing power, economic durability, and public dominance are 0.547, 0.787, and 0.762, respectively. These are mechanical construct diagnostics rather than validation against future outcomes.

6. Registered Hypotheses and Statistical Methods

6.1 H1: one-year-ahead construct validation

H1 estimates a separate model for next-year gross margin and next-year operating profitability:

```
Y[i,t+1] = alpha + beta1*DMPS[i,t] + beta2*Y[i,t]
          + beta3*log_market_cap[i,t]
          + formation_year fixed effects + FF12 fixed effects + error[i,t]
```

The future observation is the next consecutive fiscal year first filed after formation and by 31 December 2025. Each model uses complete cases and must contain at least 500 observations across at least six formation years. The null is $\beta_1 = 0$, tested two-sided at familywise 5%, with positive expected direction. Primary covariance is two-way cluster-robust by issuer and formation year with CR1 corrections and Student- t reference degrees of freedom equal to the smaller cluster count minus one. Holm adjustment covers the two outcomes. A null-imposed formation-year wild bootstrap uses 9,999 Rademacher draws and seed 20260826.

6.2 H2: value-weighted portfolio alpha

Within each June formation, issuers are sorted ascending by (DMP5, security_id) and assigned to five equal-count portfolios. Membership remains fixed from July through June except for registered terminal conversions and cash exits. July uses formation capitalization as its lagged weight; later months use the preceding valid month-end value. A portfolio-month is valid only when usable returns cover at least 90% of known lagged weight, after which weights are renormalized. Each formation requires at least 50 issuers and at least 10 in each tail. A minimum of 96 valid long–short months is required.

The sole primary return regression is:

$$\begin{aligned} Q5_minus_Q1[m] = & \alpha + b_mkt * MKT_RF[m] + b_smb * SMB[m] \\ & + b_hml * HML[m] + b_rmw * RMW[m] \\ & + b_cma * CMA[m] + b_mom * MOM[m] + error[m] \end{aligned}$$

The null is $\alpha = 0$, tested two-sided at 5%, with a positive expected direction. OLS uses Newey–West covariance with Bartlett weights and exactly 12 lags. The primary estimand is $12 * \alpha$, reported in annual percentage points. The arithmetic annualization makes no compounding claim.

6.3 H3: registered secondary cross-sectional test

For each eligible month, H3 regresses individual total returns on formation DMP5, log market capitalization, book-to-market, momentum, preformation beta, operating profitability, investment, and FF12 indicators. Continuous controls are winsorized within formation year at 1% and 99%. Each cross-section requires at least 50 complete rows and full matrix rank. The monthly DMP5 slopes are averaged by the Fama–MacBeth procedure and evaluated with 12-lag Newey–West covariance. H3 is registered secondary and cannot replace H2.

If H2 and H3 are invoked jointly, Holm adjusts their two p -values. The registered hierarchy is reported regardless of sign or significance. Secondary score variants, alternative weights, transaction costs, horizons, industry definitions, permutations, and placebos cannot replace these results. They were executed only after authoritative Run 3 and are reported with that classification in Appendix D.

7. Sample Coverage and Selection

The 6,137 sector-eligible issuer-years yield 2,573 complete scores, or 41.93% by issuer-year count. Those rows represent 64.26% of eligible usable capitalization. Of the complete-score rows, 2,487 have usable point-in-time capitalization; the C0 quality ratio inside that sample is 99.99996%, comfortably above the unchanged 90% requirement. Every June formation satisfies the registered portfolio cross-section minimum.

Coverage improves materially over time. Capitalization coverage rises from 38.38% in 2017 to 75.09% in 2024 and 70.58% in the truncated 2025 formation. This pattern is consistent with greater XBRL history and concept availability in later years, but it also means the score sample is not a balanced representation of the historical universe.

Table 1. Score coverage by June formation year

Formation	Sector-eligible	Complete DMPS	Issuer coverage	Cap coverage
2017	597	194	32.5%	38.4%
2018	653	228	34.9%	48.1%
2019	683	271	39.7%	54.8%
2020	755	288	38.1%	63.4%
2021	699	263	37.6%	63.6%
2022	674	308	45.7%	64.9%
2023	687	334	48.6%	70.5%
2024	697	347	49.8%	75.1%
2025	692	340	49.1%	70.6%
Total	6,137	2,573	41.9%	64.3%

Note. Eligible rows exclude financials, utilities, and REITs from the primary score. Cap coverage is complete-score usable capitalization divided by eligible usable capitalization. The 70% threshold governed version 1 only; version 2 preregistered coverage as a limitation rather than a stopping rule.

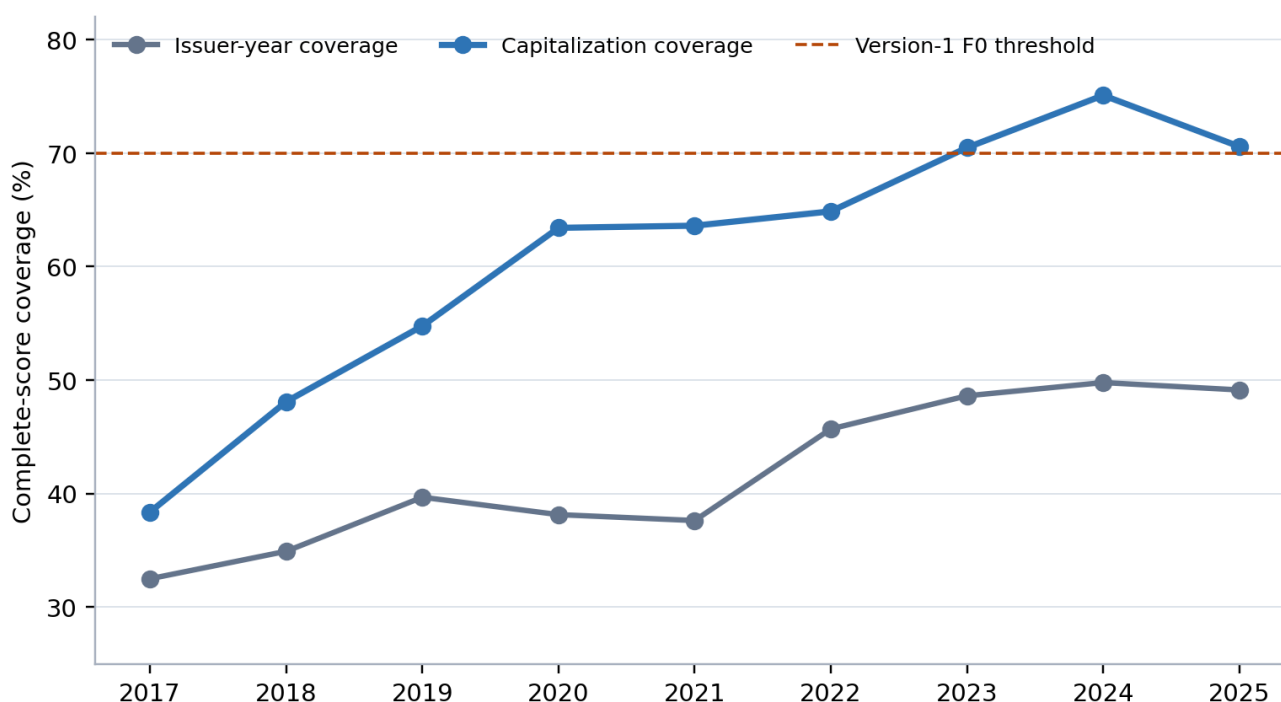


Figure 1. Complete-score coverage by formation year. The horizontal line marks the version-1 70% capitalization gate; it was removed as a stopping rule only in the separately registered version-2 addendum.

Selection favors larger and more liquid issuer-years. The excluded group has a median formation price of \$31.97, median frozen dollar volume of \$20.27 million, median observed capitalization of \$2.08 billion, and median current revenue of \$1.09 billion. The included group has corresponding medians of \$56.14, \$33.60

million, \$3.88 billion, and \$1.75 billion. Missingness is therefore economically patterned, not a benign random loss of rows.

Table 3. Included-versus-excluded formation characteristics

Characteristic	DMPS sample	Rows	Nonmissing	Median	Mean
Formation close	Included	2,573	2,573	\$56.14	\$108.03
Formation close	Excluded	3,564	3,564	\$31.96	\$70.78
60-day dollar volume	Included	2,573	2,573	\$33.60mn	\$443.80mn
60-day dollar volume	Excluded	3,564	3,564	\$20.27mn	\$115.70mn
Observed market cap	Included	2,573	2,488	\$3.88bn	\$28.31bn
Observed market cap	Excluded	3,564	3,466	\$2.08bn	\$161.70bn
Current revenue	Included	2,573	2,573	\$1.75bn	\$9.54bn
Current revenue	Excluded	3,564	2,656	\$1.09bn	\$6.07bn

Note. Included rows have all DMPS inputs; excluded rows are sector-eligible but incomplete. The outcome-blind comparison diagnoses selection and does not reweight the primary sample.

The portfolio characteristics also reveal a pronounced size gradient. Median capitalization rises from roughly \$0.73 billion in Q1 to \$28.85 billion in Q5. The high-score portfolio is not simply a rescaled version of the low-score portfolio, and the factor regression's strongly negative SMB loading is consistent with this sorting pattern. Factor adjustment matters for interpreting the raw spread.

Table 2. Formation characteristics by DMPS quintile

Portfolio	N	Mean DMPS	Median cap	Median \$ volume	Median gross margin	Median mean op. profitability
Q1	518	-0.797	\$729.16mn	\$6.56mn	36.5%	-0.4%
Q2	516	-0.147	\$1.75bn	\$13.49mn	35.7%	6.1%
Q3	513	0.112	\$3.21bn	\$25.21mn	39.1%	7.7%
Q4	516	0.339	\$7.39bn	\$67.22mn	46.1%	8.9%
Q5	510	0.737	\$28.85bn	\$232.53mn	60.9%	12.6%

Note. Statistics pool issuer-years across formations. Dollar amounts are medians. Quintiles are equal-count within formation year, so pooled counts need not be identical. Values are descriptive.

8. Results

8.1 H1: operating-performance validation

Both H1 models satisfy their registered sample requirements. The gross-margin model contains 2,060 issuer-years across eight formation years. The DMPS coefficient is -0.0150 with a two-way clustered standard error of 0.1276, a 95% interval from -0.3168 to 0.2868, and an unadjusted and Holm-adjusted p -value of 0.910. The wild formation-year bootstrap p -value is 0.774. Scaling by the observed DMPS 90th-to-10th percentile spread of 1.3286 gives a fitted -0.0200 change in next-year gross margin, or -2.00 percentage points. Adjusted R-squared is 0.045.

The operating-profitability model contains 2,102 issuer-years across the same eight formations. Its DMPS coefficient is 0.0105 with clustered standard error 0.0128, a 95% interval from -0.0198 to 0.0408, and unadjusted p = 0.440. Holm-adjusted p is 0.881 and wild-bootstrap p is 0.500. The 1.3282 P90-to-P10

spread implies a fitted +0.0139 change in next-year operating profitability, or +1.39 percentage points. Adjusted R-squared is 0.626, largely reflecting the strong persistence of current operating profitability rather than a detectable incremental DMPS association.

Table 4. H1 construct-validation regressions

Outcome	DMPS beta	Cluster SE	95% CI	p	Holm p	Wild p	N	P90-P10 fitted change
Next gross margin	-0.0150	0.1276	[-0.3168, 0.2868]	0.910	0.910	0.773	2,060	-2.00%
Next operating profitability	0.0105	0.0128	[-0.0198, 0.0408]	0.440	0.881	0.499	2,102	1.39%

Note. Separate one-year-ahead models control for the current outcome, log market capitalization, formation-year fixed effects, and FF12 fixed effects. Standard errors are two-way clustered by issuer and formation year. Holm covers both co-primary outcomes; wild p-values use 9,999 null-imposed Rademacher draws by year.

Next-year outcomes are present for most score-complete rows through 2024: 95.48% for gross margin and 97.58% for operating profitability. Model counts are lower because all controls must also be available. The central finding is not that durable operating differences are impossible. It is that this composite does not add statistically reliable one-year-ahead information conditional on its registered controls. Because neither co-primary outcome supports the expected positive association, the preregistered decision rule prohibits describing DMPS as a validated durable operating-performance construct.

8.2 H2: factor-adjusted portfolio returns

All 102 months from July 2017 through December 2025 are valid for Q1 and Q5. Minimum portfolio return-weight coverage is 97.59%, above the 90% requirement. The six-factor regression produces a monthly intercept of 0.00855, or 0.855%, with a 12-lag HAC standard error of 0.00547. The 95% monthly interval is -0.231% to 1.941%, $t = 1.563$, and $p = 0.121$. Multiplying by 12 gives an annualized alpha of 10.26%, standard error 6.56%, and 95% interval from -2.77% to 23.29%. Adjusted R-squared is 0.371.

Table 5. H2 factor model for value-weighted Q5 minus Q1

Regressor	Estimate	HAC SE	95% CI	t	p
Alpha	0.0085	0.0055	[-0.0023, 0.0194]	1.56	0.121
MKT-RF	-0.1491	0.1047	[-0.3570, 0.0589]	-1.42	0.158
SMB	-0.9134	0.2122	[-1.3346, -0.4922]	-4.31	<0.001
HML	-0.2489	0.1652	[-0.5768, 0.0791]	-1.51	0.135
RMW	0.7919	0.1951	[0.4046, 1.1793]	4.06	<0.001
CMA	-0.0431	0.3342	[-0.7065, 0.6204]	-0.13	0.898
MOM	-0.2178	0.1713	[-0.5579, 0.1223]	-1.27	0.207

Note. Monthly OLS coefficients use Newey–West covariance with exactly 12 lags. The zero-cost dependent variable is not reduced by the risk-free rate. The text also reports 12 times the monthly intercept.

The factor exposures are informative. The spread loads negatively on SMB (-0.913, $p < 0.001$) and positively on RMW (0.792, $p < 0.001$). Its market, value, investment, and momentum loadings are individually imprecise. These estimates fit the sample diagnostics: higher-DMPS companies are much larger and exhibit stronger profitability characteristics. The intercept remains positive after those controls, but the interval is wide. H2 therefore does not reject zero at 5%. The correct conclusion is *inconclusive with an economically large positive point estimate*, not evidence of zero alpha and not proof of abnormal performance.

8.3 Descriptive portfolio performance

Before factor adjustment and transaction costs, Q1 earns a 0.354% arithmetic mean per month and Q5 earns 1.520%. Their monthly difference averages 1.167%, or 14.00% on a 12-times annual basis. Compounded over the observed 102 months, one dollar grows by 3.05% in Q1, 304.22% in Q5, and 169.49% under the rebalanced monthly long–short spread convention. Annualized volatility is 28.05% for Q1, 18.50% for Q5, and 21.71% for the spread; maximum drawdowns are -72.22%, -29.79%, and -36.96%, respectively.

Table 7. Descriptive monthly portfolio returns

Series	Monthly mean	Annual mean	Annual volatility	Cumulative growth	Maximum drawdown
Low DMPS (Q1)	0.35%	4.25%	28.05%	3.05%	-72.22%
High DMPS (Q5)	1.52%	18.25%	18.50%	304.22%	-29.79%
Q5 minus Q1	1.17%	14.00%	21.71%	169.49%	-36.96%

Note. Annual mean is 12 times the arithmetic monthly mean; annual volatility is monthly standard deviation times the square root of 12. Cumulative growth compounds the monthly series. These are gross, before-cost descriptions, not substitutes for H2.

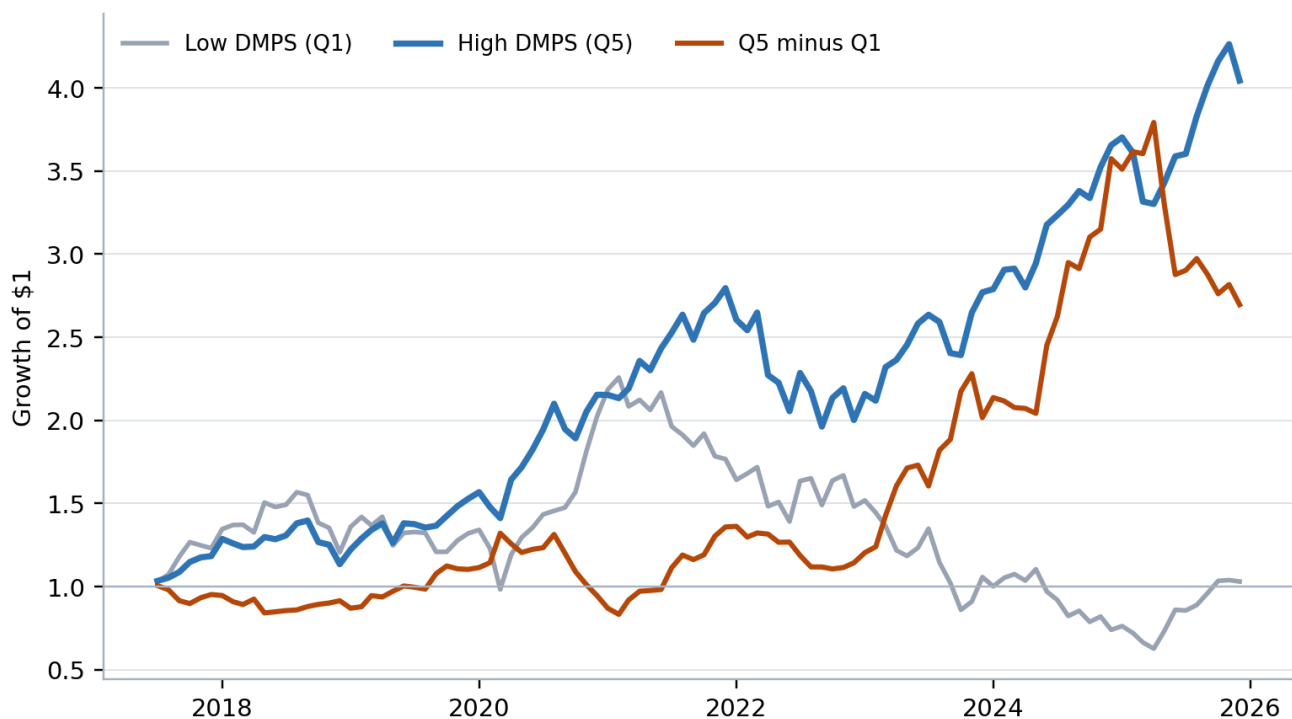


Figure 2. Growth of one dollar in the gross Q1, Q5, and rebalanced Q5 minus Q1 monthly series. The spread curve is descriptive, before costs, and is not the factor-adjusted H2 estimand.

These descriptive numbers are deliberately subordinate to H2. They omit transaction costs, financing and short-borrow frictions, capacity constraints, taxes, and implementation delay. Their cumulative paths are also sensitive to a short sample that includes the pandemic shock and recovery. They help explain the positive alpha estimate but do not change its statistical status.

8.4 H3: cross-sectional corroboration

The full-control Fama–MacBeth model is estimable in 78 months from July 2019 through December 2025. The first 24 return months fall below 50 complete observations because long preformation histories are required; no eligible month is rank deficient. The mean monthly DMPS slope is 0.00665 with HAC standard error 0.00491, 95% interval from -0.00314 to 0.01643, and $p = 0.180$. Its 12-times annualized magnitude is 7.98%, with a 95% interval from -3.76% to 19.71%. Monthly coefficients are positive in 55.13% of estimable months, and their median is 0.00507.

Table 6. H3 registered secondary Fama–MacBeth estimate

Estimand	Estimate	HAC SE	95% CI	p	Months
Monthly mean DMPS slope	0.0066	0.0049	[-0.0031, 0.0164]	0.180	78
12x annualized slope	7.98%	5.90%	[-3.76%, 19.71%]	0.180	78

Note. Valid monthly cross-sections include DMPS, the frozen firm controls, and FF12 indicators. The mean monthly slope uses 12-lag Newey–West covariance. Annualized values equal 12 times monthly values.

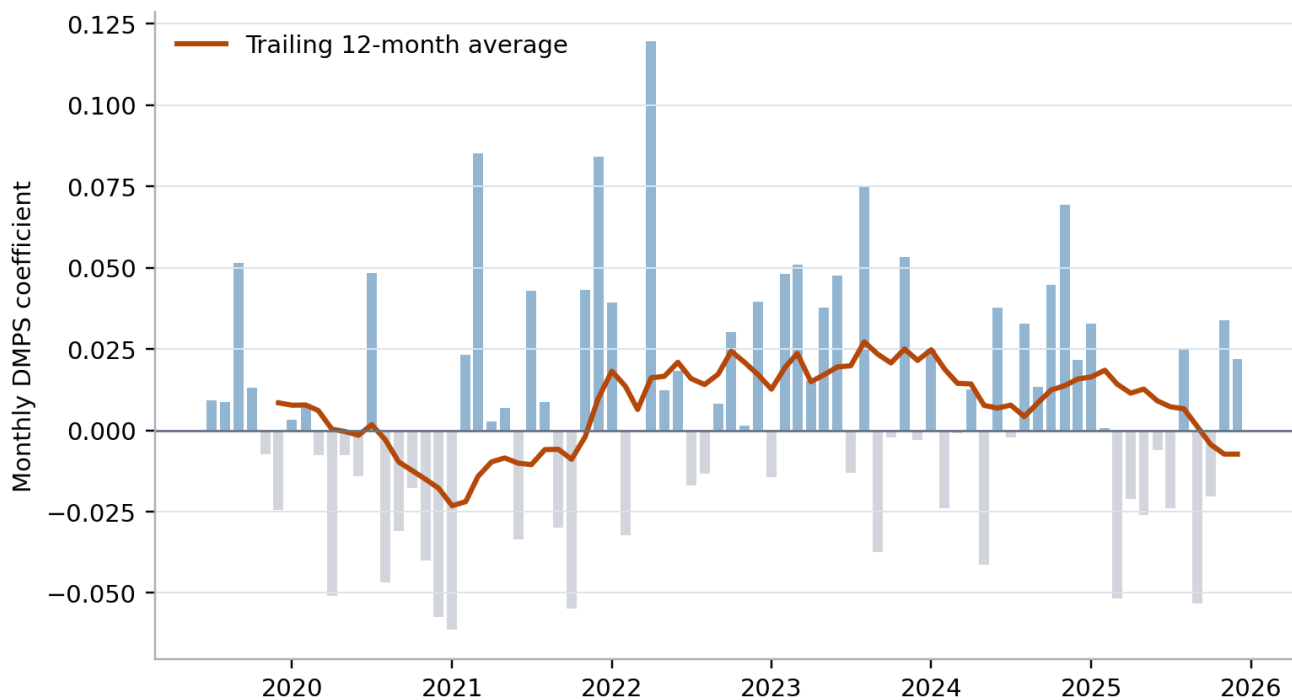


Figure 4. Monthly H3 DMPS coefficients for the 78 estimable cross-sections. The line is a trailing 12-month average; inference uses the preregistered 12-lag Newey–West covariance for the full series.

H3 points in the same positive direction as H2, but it is neither precise nor independently significant. Holm adjustment across H2 and H3 yields 0.243 for both claims. Thus the secondary cross-sectional specification does not provide registered 5% corroboration and cannot rescue the primary return test.

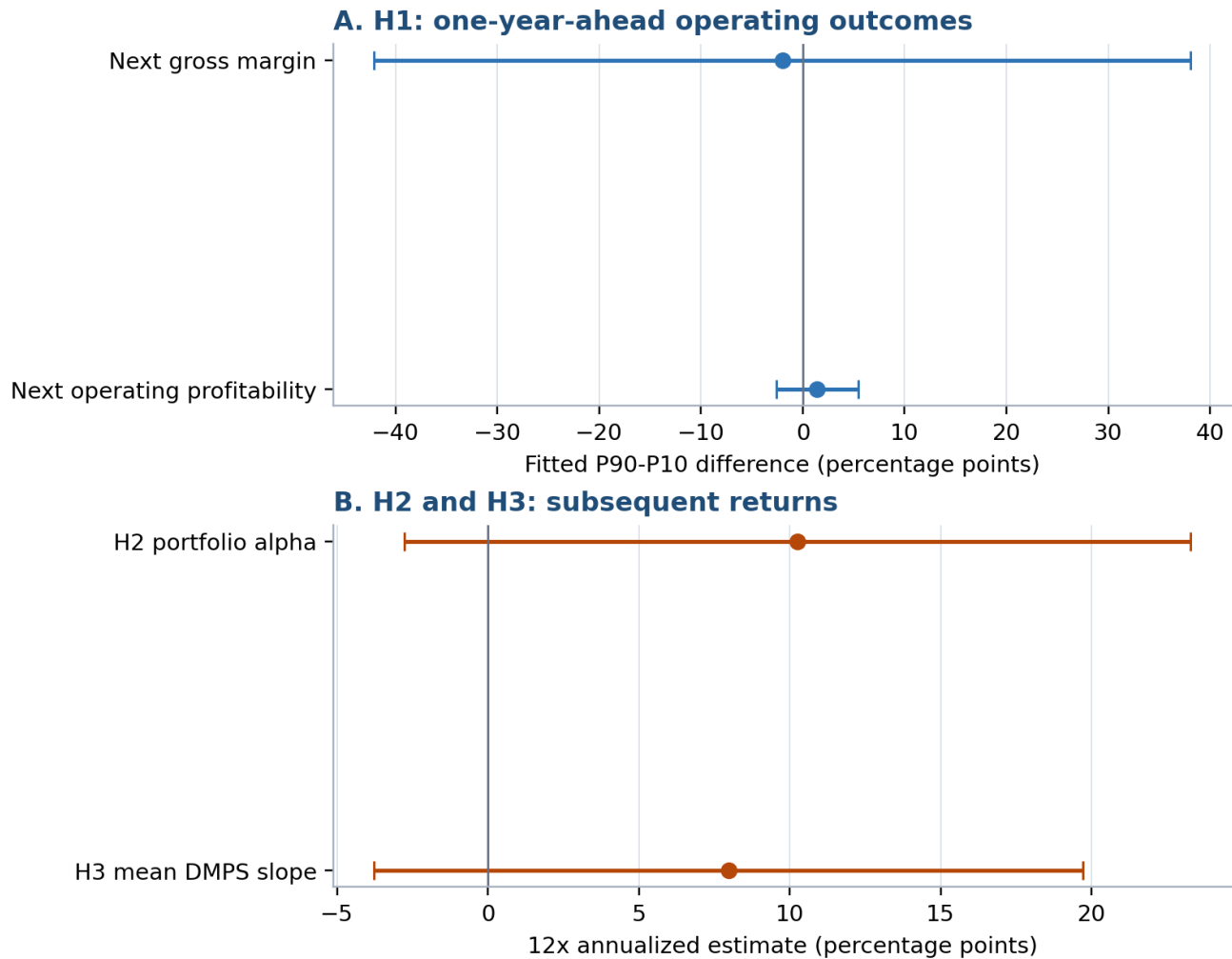


Figure 3. Registered estimates and 95% confidence intervals. Panel A scales each H1 coefficient by the observed DMPS 90th-to-10th percentile spread; Panel B reports 12-times-annualized return coefficients. Units differ across panels.

9. Interpretation

9.1 What the results support

The most defensible interpretation is narrow. Within a selected complete-case sample of U.S. public equities, a higher DMPS is associated with higher subsequent raw and factor-adjusted return point estimates over 2017–2025. Two independently registered return estimators point in the same direction. The magnitudes are large enough that a simple “nothing is there” reading would be unwarranted.

At the same time, neither return test reaches the preregistered 5% level, their intervals cross zero, and the H2 interval allows outcomes ranging from a modest negative annual alpha to a very large positive one. This is weak identification, not a clean return anomaly. A longer out-of-sample interval could narrow uncertainty, but adding future data must be done as a prospective extension rather than selectively changing this completed study.

9.2 What the results do not support

The paper does not establish that market power causes returns. It does not identify a product market, estimate a structural markup, measure private competitors, or separate risk from mispricing. It also does not validate the “durable” label through the registered one-year-ahead operating tests. The positive return point estimates therefore cannot be narrated as investors underreacting to operating durability when that intermediate link is not empirically supported here.

One plausible interpretation is that the composite captures a mix of large-firm quality, public-peer dominance, and accounting availability. The significant negative SMB and positive RMW loadings reinforce that possibility. Another is that the one-year operating outcomes are too noisy or too highly conditioned on current levels for the composite to add detectable information, while returns respond to a different correlated attribute. Those are hypotheses for future designs, not conclusions available from the present models.

9.3 Economic versus statistical significance

The return estimates illustrate why confidence intervals should govern interpretation. An annual H2 point estimate of 10.26% is economically meaningful, and it lies near the outcome-blind minimum-detectable-effect calculation for a 3% monthly residual volatility. Yet the observed standard error produces an interval more than 26 percentage points wide. Non-rejection cannot be converted into evidence that the economically relevant effect is zero; conversely, a large point estimate cannot be treated as established when the interval includes zero.

For H1, operating profitability is estimated more tightly than gross margin, but its interval still includes modest negative and positive incremental effects. Gross margin is especially imprecise under only eight year clusters. The wild-bootstrap results provide no hidden support. The registered inference hierarchy therefore produces a coherent answer: H1 unsupported, H2 inconclusive, H3 noncorroborating.

10. Limitations

First, complete-score coverage is only 41.93% by issuer-year count and 64.26% by usable capitalization. Included firms are larger and more liquid, and coverage improves over time. Complete-case estimates need not generalize to small, young, sparsely tagged, or historically delisted firms. The version-2 addendum transparently permits analysis despite this limitation; it does not make the limitation disappear.

Second, the sample has only 102 return months and eight H1 formation years. HAC estimators and small-cluster sensitivities address dependence but cannot manufacture information. H3 loses the first 24 months to its long-history controls. The resulting intervals are appropriately wide.

Third, public sales share is a deliberately restricted proxy. It ignores private firms, imports, geographic segmentation, multi-product boundaries, and substitution. FF12 standardization is broad. The score can describe public-company revenue dominance without establishing anticompetitive conduct or a defensible antitrust market.

Fourth, accounting comparability is imperfect. XBRL concept precedence and duration rules reduce discretion, but issuer tagging, cost classification, acquisitions, divestitures, fiscal-year changes, and intangible investment can affect margins and profitability. Five-year history requirements trade measurement stability for selection.

Fifth, the historical market layer is not CRSP. The study reconstructs identity and terminal economics from SEC frames, Marketstack raw prices, primary transaction evidence, and narrowly frozen independent routes. The audit is unusually explicit, but public sources can still contain omissions. The return results are gross and do not model bid–ask spreads, market impact, financing, short availability, taxes, or operational latency.

Sixth, the sample period includes unusual macroeconomic regimes and is too short for fine subperiod claims. No confirmatory subperiod test was registered. Pandemic-era movements, rate changes, and the dominance of large technology firms may influence the pooled estimates.

Finally, the study tests a fixed composite. A null H1 does not imply that every component is uninformative, but optimizing weights or choosing components after these outcomes would overfit the same evidence. Any learned score, product-market refinement, antitrust extension, or alternative horizon belongs in a new training/validation/test design or a separately registered extension.

11. Reproducibility, Corrections, and Data Availability

The public OSF record is <https://osf.io/fvgsr/>. The repository binds the original registration and version-2 approval to source commits and artifact hashes. Authoritative Run 3 started at 22:36:40 UTC and completed at 22:43:09 UTC on 27 August 2026 from source commit `403fc01f4c5c1e24cfd1022ee676dbe32e2460b8`. It made zero network requests, verified 17 frozen input hashes, confirmed all 7,152 holding-window routes, verified and parsed 1,357 raw Marketstack responses referenced by the score sample, and wrote 14 hash-recorded result artifacts.

The manuscript builder re-verifies every authoritative output hash before creating tables, figures, Markdown, HTML, DOCX, or PDF. It separately verifies the post-registration completion record and its 14 output hashes before rendering Appendix D. Machine-readable tables are generated directly from the Parquet outputs. The build manifest records both completion hashes, the primary- and post-registration-results hashes, and manuscript artifact hashes. The failed primary and post-registration runs and their corrections remain visible rather than being overwritten.

SEC filings, SEC 13F lists, and Kenneth French factors are public subject to their providers' terms. Marketstack payloads are licensed inputs and are not redistributed by the repository; independent reproduction of the full price layer requires lawful access to equivalent data. Hashes, request lineage, transformation code, evidence contracts, and result tables permit auditing without exposing API credentials. No credential is embedded in this manuscript or its build artifacts.

Author declarations: Affiliation: Independent Researcher. Publisher and ownership: Monopoly MOAT (mplymoat.com) is an independent research publication owned and operated by the author. Funding: This research was self-funded by the author and received no external funding. Competing interests: The author holds diversified investment funds that may include securities represented in the study universe. No issuer, fund sponsor, or index provider funded, commissioned, approved, reviewed, or influenced this research.

12. Conclusion

This preregistered public-data study asks a simple question with deliberately difficult measurement: do observable characteristics associated with durable market power predict future fundamentals and returns? The answer is mixed but clear under the registered hierarchy.

DMPS does not predict either one-year-ahead operating outcome with statistical reliability after conditioning on the current outcome, size, year, and industry. It therefore is not validated here as a durable operating-performance construct. The high-minus-low portfolio and the security-level Fama–MacBeth model both produce economically large positive return point estimates, but both are imprecise, both cross zero, and their joint Holm-adjusted evidence is not significant. The study cannot claim abnormal return predictability at 5%.

That is not a failed research outcome. It narrows the claim, exposes where public-data selection enters, and leaves a falsifiable prospective question. A future extension should add genuinely new return years, retain the locked score or register any revision before observing the extension sample, and treat product-market definitions and private competitors explicitly if the goal shifts from public-company characteristics to economic market power.

References

- Asness, C. S., Frazzini, A., & Pedersen, L. H. (2019). Quality minus junk. *Review of Accounting Studies*, 24, 34–112. <https://doi.org/10.1007/s11142-018-9470-2>
- Autor, D., Dorn, D., Katz, L. F., Patterson, C., & Van Reenen, J. (2017). Concentrating on the fall of the labor share. *American Economic Review*, 107, 180–185. <https://doi.org/10.1257/aer.p20171102>
- Cameron, A. C., Gelbach, J. B., & Miller, D. L. (2011). Robust inference with multiway clustering. *Journal of Business & Economic Statistics*, 29, 238–249. <https://doi.org/10.1198/jbes.2010.07136>
- Carhart, M. M. (1997). On persistence in mutual fund performance. *Journal of Finance*, 52, 57–82. <https://doi.org/10.1111/j.1540-6261.1997.tb03808.x>
- De Loecker, J., Eeckhout, J., & Unger, G. (2020). The rise of market power and the macroeconomic implications. *Quarterly Journal of Economics*, 135, 561–644. <https://doi.org/10.1093/qje/qjz041>
- Fama, E. F., & French, K. R. (2015). A five-factor asset pricing model. *Journal of Financial Economics*, 116, 1–22. <https://doi.org/10.1016/j.jfineco.2014.10.010>
- Fama, E. F., & MacBeth, J. D. (1973). Risk, return, and equilibrium: Empirical tests. *Journal of Political Economy*, 81, 607–636. <https://doi.org/10.1086/260061>
- Grullon, G., Larkin, Y., & Michaely, R. (2019). Are U.S. industries becoming more concentrated? *Review of Finance*, 23, 697–743. <https://doi.org/10.1093/rof/rfz007>
- Holm, S. (1979). A simple sequentially rejective multiple test procedure. *Scandinavian Journal of Statistics*, 6, 65–70. <https://www.jstor.org/stable/4615733>
- Newey, W. K., & West, K. D. (1987). A simple, positive semi-definite, heteroskedasticity and autocorrelation consistent covariance matrix. *Econometrica*, 55, 703–708. <https://doi.org/10.2307/1913610>
- Novy-Marx, R. (2013). The other side of value: The gross profitability premium. *Journal of Financial Economics*, 108, 1–28. <https://doi.org/10.1016/j.jfineco.2013.01.003>
- Shumway, T. (1997). The delisting bias in CRSP data. *Journal of Finance*, 52, 327–340. <https://doi.org/10.1111/j.1540-6261.1997.tb03818.x>

Appendix A. Outcome Coverage

The one-year-ahead outcome panel uses formation years 2017–2024; 2025 has no eligible next fiscal year before the frozen outcome end. Among 2,233 score-complete issuer-years in those formations, gross margin is observed for 2,132 and operating profitability for 2,179. Terminal-event rows have lower future-filing coverage, but only seven score-complete rows through 2024 experience a terminal event during the holding window.

Table A1. One-year-ahead outcome availability

Outcome	Eligible score rows	Observed outcome	Observed share	Model N
Next gross margin	2,233	2,132	95.5%	2,060
Next operating profitability	2,233	2,179	97.6%	2,102

Note. The denominator is score-complete issuer-years from 2017–2024. Regression samples are smaller because H1 also requires every registered covariate.

Appendix B. Return-Construction Audit

The score-sample return panel contains 28,791 monthly security rows, of which 28,789 are usable. One ordinary provider row has an invalid raw close and is skipped as missing; it is neither a zero return nor an endpoint. One boundary return is invalid because valid endpoints cross an intervening calendar month. All 42 governed action transitions required by the score sample are used, and no adjusted-price field enters. Every one of the five portfolios has 102 valid months; the lowest observed monthly weight coverage is 97.59%.

Preformation-control lineage contains 6,869 response references across 2,573 security-years, pointing to 1,357 unique raw responses. The cache audit indexes 5,826 unique request hashes, excludes three control JSON files, collapses one byte-identical duplicate, and rejects conflicting duplicates. These checks completed before Run 3 constructed future outcomes or holding-period results.

Appendix C. Registered Result Hierarchy

H1 is unsupported, so DMPS is not validated durable operating performance. **H2** estimates 10.26% annualized alpha but is inconclusive because its interval includes zero ($p = 0.121$). Secondary **H3** estimates 7.98% annualized but is also imprecise ($p = 0.180$). Their joint Holm-adjusted p is 0.243. Descriptive evidence is subordinate, and no secondary result changes these conclusions.

Appendix D. Post-Registration Extensions

Submission disclosure: every result in this appendix was computed after authoritative preregistered Run 3. The temporal analysis is exploratory; the named robustness and falsification analyses are prespecified secondary analyses; and the independently backfilled H3 model is a post-registration coverage diagnostic. They add information but do not replace, relabel, or rescue the registered H1, H2, or H3 conclusions.

D.1 Did future operating performance change over time?

The available question is narrower than “2017 through 2026.” The frozen outcome source ends on 31 December 2025, so complete one-year-ahead observations exist for the June 2017 through June 2024 formation cohorts. Complete fiscal-year 2026 outcomes do not yet exist in that lock. I therefore estimate an exploratory linear cohort trend over 2017–2024 while controlling for DMPS, the current value of the outcome, log formation capitalization, and FF12 indicators. The model also interacts DMPS with the cohort-year index. Inference uses the same two-way issuer/year clustered covariance convention as H1, supplemented by exact enumeration of all 256 formation-year sign patterns.

The conditional gross-margin trend is -0.00338 per cohort year, or about -0.34 percentage point, with a 95% interval from -1.85 to 1.17 percentage points ($p = 0.613$; exact sign-pattern $p = 0.695$). The operating-profitability trend is -0.00137, or about -0.14 percentage point per year, with a 95% interval from -0.33 to 0.06 percentage point ($p = 0.138$; exact $p = 0.176$). Both point estimates are negative, but neither is distinguishable from zero with only eight year clusters. The DMPS-by-year interactions are also imprecise: 0.01050 for gross margin ($p = 0.570$) and 0.00422 for operating profitability ($p = 0.207$). Thus there is no statistically detectable linear improvement or decline and no evidence that H1's conditional DMPS relationship strengthened or weakened over the observed cohorts.

Table A3. Exploratory H1 calendar-trend estimates

Outcome	Cohorts	Trend/year	95% CI	p	Exact p	DMPS x year p
Next gross margin	2017-2024	-0.0034	[-0.0185, 0.0117]	0.613	0.695	0.570
Next operating profitability	2017-2024	-0.0014	[-0.0033, 0.0006]	0.138	0.176	0.207

Note. These linear 2017-2024 models were run after authoritative Run 3. No confirmatory subperiod test was registered. Covariance is two-way clustered by issuer and formation year; exact p enumerates all 256 formation-year sign patterns.



Figure 5. Exploratory one-year-ahead outcomes by formation cohort. Lines show medians and capitalization-weighted means. The analysis covers 2017-2024 because the frozen outcome lock ends on 31 December 2025. These descriptive cohort series do not constitute registered subperiod tests.

The plotted medians and capitalization-weighted means move across cohorts, particularly around the pandemic period, but they do not form a stable monotonic trend. Those descriptive series also reflect changing sample composition. A future 2025-formation outcome can be appended only after the relevant next annual filing becomes public; a 2026 formation cannot have a complete one-year-ahead fiscal outcome in 2026.

D.2 Does repairing early price history establish H3?

No. A bounded Marketstack attempt made 236 planned requests and received HTTP 422 for every route—207 invalid-exchange and 29 invalid-symbol responses—so it supplied no usable payload. An independent Yahoo chart backfill then validated identities by exact symbol, permitted exchange, equity type, and a scale-normalized raw-close trajectory match against the frozen Marketstack overlap. Adjusted close was neither requested nor used. Of 236 identities, 179 validated and 57 remained unresolved; the valid routes covered 327 of the 422 early score-complete issuer-years.

The splice uses the independent history only through 24 August 2016, before the frozen Marketstack series begins, and reserves the following overlap for identity validation. Early momentum completeness rises from 227 to 376 issuer-years and beta completeness from zero to 327. This makes all 102 monthly cross-sections estimable rather than 78.

The added observations do not make H3 significant. The augmented monthly DMPS slope is 0.00437 with 12-lag HAC standard error 0.00410, 95% interval [-0.00376, 0.01250], and $p = 0.289$. Its 12-times value is 5.24% with a 95% interval from -4.52% to 15.00%. The estimate is smaller and less statistically persuasive than registered H3. This is useful diagnosis: missing early histories were not the reason registered H3 failed to reject zero.

Table A4. Registered H3 and post-registration coverage diagnostic

Specification	Months	Early momentum	Early beta	Annual estimate	95% CI	p
Registered H3	78	227	0	7.98%	[-3.76%, 19.71%]	0.180
Augmented diagnostic	102	376	327	5.24%	[-4.52%, 15.00%]	0.289

Note. The augmented row uses independently validated preformation histories and is a coverage diagnostic, not a replacement for registered H3. Annualized values equal 12 times the monthly mean slope.

D.3 Named secondary operating analyses

DMPS itself is highly persistent in rank. Spearman correlations of within-cohort DMPS percentiles are 0.953 over one year, 0.842 over three years, and 0.746 over five years. This shows that the measured characteristic is persistent; it does not show that it predicts future outcomes.

For the prespecified positive-part decline outcome, higher DMPS predicts a smaller next-year operating-profitability decline. The coefficient is -0.02885, with 95% interval [-0.05229, -0.00541], unadjusted $p = 0.0227$, and Holm-adjusted $p = 0.0453$ across the two decline outcomes. Gross-margin decline is not significant: coefficient -0.11633, 95% interval [-0.31755, 0.08489], Holm-adjusted $p = 0.214$. This narrow secondary result says that high-score firms experience smaller measured positive-part operating-profitability setbacks conditional on controls. It does not overturn the null co-primary H1 level tests or validate the full “durability” claim.

Table A6. Prespecified secondary construct analyses

Analysis	Estimate	N	95% CI	p	Holm p
Gross-margin decline magnitude	-0.1163	2,060	[-0.3176, 0.0849]	0.214	0.214
Operating-profitability decline magnitude	-0.0288	2,102	[-0.0523, -0.0054]	0.023	0.045
DMPS rank persistence, 1y	0.953	1,959	—	<0.001	—
DMPS rank persistence, 3y	0.842	1,112	—	<0.001	—
DMPS rank persistence, 5y	0.746	598	—	<0.001	—

Note. Decline magnitude is max(current minus next, zero), with Holm adjustment across the two decline outcomes. Rank persistence is the Spearman correlation of within-year DMPS percentiles; its p-value is descriptive and not part of a confirmatory H1 family.

D.4 Named secondary return analyses

The return robustness exercises are mixed. Requiring complete five-year histories for all score inputs gives annualized alpha of 8.49% ($p = 0.114$; family Holm $p = 0.229$). Dropping public dominance and using the two accounting components gives -1.16% ($p = 0.876$). Equal-weight quintiles give 4.62% ($p = 0.151$), while

FF49 industry construction gives 3.33% ($p = 0.716$). The value-weighted decile analysis is correctly left untested because its minimum tail contains 19 issuers rather than the required 20.

The estimated annualized alpha remains positive after stylized one-way trading costs: 10.15%, 10.00%, and 9.74% at 10, 25, and 50 basis points, respectively. None is individually significant and their cost-family Holm-adjusted p is 0.376. The public-HHI interaction has an unadjusted monthly slope of 0.00531 ($p = 0.0329$), but its industry-granularity family value is $p = 0.0658$ and therefore misses the prespecified 5% threshold.

The 10,000 within-formation-year-by-FF12 label permutations preserve stratum-specific portfolio sizes. The observed H2 alpha lies at the 97.02nd percentile, but the two-sided empirical p -value is 0.0598. That is suggestive and close to 5%, not statistically significant under the fixed two-sided rule.

Overlapping cumulative Q5-minus-Q1 strategy spreads are 7.92% over six months (Holm $p = 0.0710$), 18.96% over 12 months (Holm $p = 0.0232$), and 65.54% over 36 months (Holm $p < 0.001$). The 60-month test is untestable because its required 59 HAC lags are not fewer than its 43 observations. The significant 12- and 36-month rows are cumulative strategy-return summaries, not factor-model intercepts; overlap makes them highly dependent. They cannot support the claim “established abnormal returns,” whose registered test remains the six-factor H2 alpha.

Table A5. Prespecified secondary return analyses

Analysis	Estimand	N/months	Estimate	p	Holm p	Status
Strict five-of-five score	Alpha	102	8.49%	0.114	0.229	Tested
Two-component score	Alpha	102	-1.16%	0.876	0.876	Tested
Equal-weight quintiles	Alpha	102	4.62%	0.151	0.151	Tested
FF49 construction	Alpha	102	3.33%	0.716	0.716	Tested
Value-weighted deciles	Alpha	tail=19	—	—	—	Untested: tail <20
10 bp one-way costs	Net alpha	102	10.15%	0.125	0.376	Tested
25 bp one-way costs	Net alpha	102	10.00%	0.131	0.376	Tested
50 bp one-way costs	Net alpha	102	9.74%	0.142	0.376	Tested
6-month horizon	Cumulative spread	97	7.92%	0.071	0.071	Tested
12-month horizon	Cumulative spread	91	18.96%	0.012	0.023	Tested
36-month horizon	Cumulative spread	67	65.54%	<0.001	<0.001	Tested
60-month horizon	Cumulative spread	43	—	—	—	Untested: HAC lags $\geq$ N
DMPS x public HHI	Monthly slope	102	0.0053	0.033	0.066	Tested
Within-year x FF12 permutation	Empirical H2	10,000	10.26%	0.060	—	97.02nd percentile

Note. Alpha rows use the six-factor model. Horizon rows are overlapping cumulative annually rebalanced Q5-minus-Q1 spreads, not factor alphas. Holm p -values are within the named preregistered family. A dash denotes an inapplicable or untestable value.

D.5 What evidence would justify a stronger claim?

“Prove” is too strong for an empirical return study. A defensible stronger claim requires a fixed, independently observed sample in which the same score, universe rules, portfolio construction, six-factor model, 12-lag covariance estimator, two-sided threshold, and stopping date are frozen before outcomes are analyzed. The relevant evidence is a positive H2 alpha whose confidence interval excludes zero, ideally accompanied by a directionally consistent H3 estimate, robustness to registered transaction costs, and replication outside the original period. Repeatedly changing the model or stopping when p first crosses 0.05 would not establish the claim.

Rough precision extrapolations underscore the limitation. If the observed effect and dependence-adjusted noise remained exactly constant, H2 would need about 161 total months merely to reach a two-sided 5% threshold and about 328 for 80% power. Registered H3 gives analogous totals of about 164 and 335 months; the augmented H3 diagnostic, whose effect is smaller, gives about 345 and 705 months. These are not promises, deadlines, or guaranteed sample sizes. Effects and volatility can change, so additional data may strengthen, weaken, or reverse the estimates. The current paper is complete now; any genuinely future sample is a separate replication rather than unfinished work on this study.

Finally, both look-ahead falsifications behave correctly. A future-score placebo is rejected because the future score is not public at the current cutoff, and shifting filing availability forward by one year triggers 43,126 violations. The completed analysis therefore gains no apparent precision by relaxing its point-in-time boundary.