

**The 2022**  
**Repricing of Public SaaS:**  
**Cohort-Specific Margin Reweighting**  
**and a Duration-Driven Level Shift**

Nicolas Moritz Schneider

Dissertation written under the supervision of Roberto Vincenzi

Dissertation submitted in partial fulfillment of requirements for the MSc in International Finance, at Universidade Católica Portuguesa, and for the MSc in Accounting and Financial Management at Bocconi University, 28.05.2026.

## **Abstract**

This paper studies the repricing of publicly traded Software-as-a-Service (SaaS) firms in 2022. Recent practitioner research has framed the repricing in generative-AI terms, either as a tailwind that deepens recurring-revenue moats or as a substitution risk under which AI agents commoditize per-seat pricing. I examine these explanations against a more standard asset-pricing mechanism in which firms with long-duration cash flows respond more strongly to changes in discount rates.

The analysis uses quarterly data for 90 U.S. SaaS firms from 2014Q1 through 2024Q4. Three results emerge. First, the structural break in valuations occurs in 2022Q1, at the start of the Federal Reserve tightening cycle and before the release of ChatGPT in November 2022. Second, the cross-sectional valuation function shifts toward operating margins, although this effect appears mainly among firms whose margins were already informative before 2022. Third, the decline in valuations is larger for firms with longer cash-flow duration.

The cross-sectional patterns are difficult to reconcile with the main AI-based explanations. Repricing differences are more closely related to pre-existing growth and profitability characteristics than to measures of AI exposure. Overall, the evidence is more consistent with the duration-based mechanism commonly documented in the broader asset-pricing literature.

**Keywords:** SaaS Valuation, EV/Revenue, Rule of 40, Duration Channel, Structural Break, Cohort Heterogeneity, ChatGPT, Cross-Sectional Pricing.

**Author:** Nicolas Moritz Schneider

**Title:** The 2022 Repricing of Public SaaS: Cohort-Specific Margin Reweighting and a Duration-Driven Level Shift

## Resumo

Este artigo estuda a reavaliação das empresas de Software-as-a-Service (SaaS) cotadas em bolsa em 2022. A literatura prática recente enquadró-a como IA generativa, quer como vento favorável que aprofunda os fossos da receita recorrente, quer como risco de substituição em que os agentes de IA banalizam o preço por utilizador. Confronto estas explicações com um mecanismo mais convencional de precificação de ativos, segundo o qual empresas com fluxos de caixa de longo prazo reagem mais fortemente a variações nas taxas de desconto.

A análise utiliza dados trimestrais de 90 empresas SaaS dos EUA, de 2014T1 a 2024T4. Emergem três resultados. Primeiro, a quebra estrutural ocorre em 2022T1, no início do aperto da Reserva Federal e antes do ChatGPT em novembro de 2022. Segundo, a função de valoração transversal reorienta-se para a margem operacional, sobretudo nas empresas cujas margens já eram informativas antes de 2022. Terceiro, a queda é maior em empresas com fluxos de caixa de maior duração.

Os padrões transversais são difíceis de conciliar com explicações baseadas em IA. As diferenças relacionam-se mais a características pré-existent de crescimento e rentabilidade do que a medidas de exposição à IA. No conjunto, a evidência é mais consistente com o mecanismo baseado em duração documentado na literatura mais ampla sobre precificação de ativos.

**Palavras-chave:** Avaliação SaaS, EV/Receita, Rule of 40, Canal de Duração, Quebra Estrutural, Heterogeneidade de Coortes, ChatGPT, Precificação Transversal.

**Autor:** Nicolas Moritz Schneider

**Título:** The 2022 Repricing of Public SaaS: Cohort-Specific Margin Reweighting and a Duration-Driven Level Shift

## Contents

|                                                                                              |           |
|----------------------------------------------------------------------------------------------|-----------|
| <b>List of Figures.....</b>                                                                  | <b>6</b>  |
| <b>List of Tables.....</b>                                                                   | <b>7</b>  |
| <b>List of Equations.....</b>                                                                | <b>8</b>  |
| <b>List of Abbreviations.....</b>                                                            | <b>9</b>  |
| <b>1. Introduction .....</b>                                                                 | <b>11</b> |
| <b>2. Literature and Hypotheses .....</b>                                                    | <b>14</b> |
| 2.1 Literature Review.....                                                                   | 14        |
| 2.1.1 Multiples as reduced-form DCF representations .....                                    | 14        |
| 2.1.2 The duration channel and the cross-sectional pricing of long-duration cash flows ..... | 14        |
| 2.1.3 Evolution in value relevance .....                                                     | 14        |
| 2.1.4 SaaS-specific valuation drivers and the practitioner debate.....                       | 15        |
| 2.1.5 The practitioner AI-led narrative and its empirical antecedents .....                  | 15        |
| 2.2 Hypotheses Development.....                                                              | 16        |
| 2.2.1 A simple Gordon-growth scaffold.....                                                   | 16        |
| 2.2.2 Competing views.....                                                                   | 17        |
| 2.2.3 Hypothesis 1: Stability of explanatory power .....                                     | 18        |
| 2.2.4 Hypothesis 2: Stability of coefficients .....                                          | 18        |
| 2.2.5 Hypothesis 3: Absence of unexplained level shift .....                                 | 19        |
| <b>3. Research Design.....</b>                                                               | <b>20</b> |
| 3.1 Baseline Specification .....                                                             | 20        |
| 3.2 Structural-Change Specification .....                                                    | 20        |
| 3.3 Testing H3: Out-of-Sample Residual Procedure.....                                        | 21        |
| 3.4 Identification: Rolling-Breakpoint and Placebo.....                                      | 22        |
| 3.5 Power and Minimum Detectable Effect .....                                                | 23        |
| 3.6 Estimation.....                                                                          | 23        |
| 3.7 What the Design Can and Cannot Identify.....                                             | 24        |
| 3.8 Mapping to Hypotheses.....                                                               | 24        |
| <b>4. Data and Sample.....</b>                                                               | <b>25</b> |
| 4.1 Sample Construction .....                                                                | 25        |
| 4.2 Time Period and Candidate Break Dates .....                                              | 26        |
| 4.3 Placebo Sample .....                                                                     | 26        |
| 4.4 Data Sources.....                                                                        | 27        |
| 4.5 Operating Metrics.....                                                                   | 28        |
| 4.6 Descriptive Statistics .....                                                             | 28        |
| <b>5. Empirical Results .....</b>                                                            | <b>30</b> |
| 5.1 Baseline Cross-Sectional Mapping .....                                                   | 30        |
| 5.2 Tests of Stability: Pre- versus Post-Shock .....                                         | 31        |
| 5.3 Rolling-Breakpoint Test.....                                                             | 35        |
| 5.4 Placebo Results .....                                                                    | 38        |

|                                                                                   |           |
|-----------------------------------------------------------------------------------|-----------|
| 5.5 Cross-Sectional Heterogeneity in the Margin Reweighting and Level Shift ..... | 40        |
| <b>6. Robustness and Additional Analyses .....</b>                                | <b>46</b> |
| 6.1 Forward Consensus Growth .....                                                | 46        |
| 6.2 Rate-Cycle Triple Interaction.....                                            | 47        |
| 6.3 Excluding the Largest Firms .....                                             | 48        |
| 6.4 Long-Differences Specification .....                                          | 50        |
| 6.5 Balanced-Sample Decomposition .....                                           | 52        |
| 6.6 Alternative Winsorization Schemes .....                                       | 53        |
| <b>7. Discussion .....</b>                                                        | <b>55</b> |
| <b>8. Conclusion.....</b>                                                         | <b>58</b> |
| 8.1 Limitations .....                                                             | 59        |
| 8.2 Avenues for Future Research .....                                             | 59        |
| <b>9. References .....</b>                                                        | <b>60</b> |
| <b>Appendix A: Variable Definitions .....</b>                                     | <b>62</b> |

## List of Figures

|                                                                                                                                           |    |
|-------------------------------------------------------------------------------------------------------------------------------------------|----|
| <b>Figure 1.</b> Rolling-breakpoint test: joint significance of equation (2) interactions across candidate break dates 2021Q1–2023Q4..... | 36 |
| <b>Figure 2.</b> Temporal dynamics of the H3 unexplained level shift, 2022Q4–2024Q4. ....                                                 | 38 |

## List of Tables

|                                                                                                                                                                                                                |    |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| <b>Table I.</b> Directional predictions across hypotheses and views. ....                                                                                                                                      | 19 |
| <b>Table 1.</b> Summary statistics for the SaaS regression sample, 2014Q1–2024Q4.....                                                                                                                          | 29 |
| <b>Table 2.</b> Baseline cross-sectional valuation function for U.S. SaaS firms, 2014Q1–2024Q4.31                                                                                                              |    |
| <b>Table 3.</b> Structural-change test: equation (2) interaction coefficients at the data-determined break (2022Q1) and the ChatGPT-release benchmark (2022Q4). ....                                           | 32 |
| <b>Table 4.</b> H3 out-of-sample residual test: mean post-period residual against zero.....                                                                                                                    | 34 |
| <b>Table 5.</b> Rolling-breakpoint test: joint significance of equation (2) interactions across candidate break dates. ....                                                                                    | 37 |
| <b>Table 6.</b> Placebo structural-change test: equation (2) interaction coefficients in the energy-and-utilities sample at the data-determined break (2022Q1) and the ChatGPT-release benchmark (2022Q4)..... | 39 |
| <b>Table 7.</b> Cohort heterogeneity in the H2 margin reweighting. ....                                                                                                                                        | 42 |
| <b>Table 8.</b> Cross-sectional structure of the H3 level shift: cohort means and duration regression. ....                                                                                                    | 45 |
| <b>Table 9.</b> Rate-cycle robustness: equation (2) augmented with a rate-cycle triple interaction.48                                                                                                          |    |
| <b>Table 10.</b> Forward consensus growth: equations (1) and (2) at the 2022Q1 break.....                                                                                                                      | 47 |
| <b>Table 11.</b> Excluding the largest firms: equations (1) and (2) at the 2022Q1 break.....                                                                                                                   | 50 |
| <b>Table 12.</b> Long-differences regression and panel unit-root test. ....                                                                                                                                    | 52 |
| <b>Table 13.</b> Alternative winsorization schemes: equations (1) and (2) at the 2022Q1 break....                                                                                                              | 54 |

## List of Equations

|                                                                                                                                                                      |    |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| <b>Equation (1).</b> Baseline cross-sectional specification: $\ln(\text{EV}/\text{Revenue})$ on operating drivers with firm and calendar-quarter fixed effects. .... | 20 |
| <b>Equation (2).</b> Structural-change specification: baseline augmented with $\text{driver} \times \text{PostShock}$ interactions. ....                             | 20 |

## List of Abbreviations

| Abbreviation | Meaning / Name                                                                     |
|--------------|------------------------------------------------------------------------------------|
| ADF          | Augmented Dickey–Fuller (unit-root test)                                           |
| AI           | Artificial Intelligence                                                            |
| APV          | Adjusted Present Value                                                             |
| BVP          | Bessemer Venture Partners                                                          |
| CAPM         | Capital Asset Pricing Model                                                        |
| CF           | Cash Flow                                                                          |
| CUSUM        | Cumulative Sum (of squares test)                                                   |
| DCF          | Discounted Cash Flow                                                               |
| EV           | Enterprise Value                                                                   |
| EV/EBITDA    | Enterprise Value to Earnings Before Interest, Taxes, Depreciation and Amortization |
| EV/Revenue   | Enterprise Value to Revenue                                                        |
| FCF          | Free Cash Flow                                                                     |
| FE           | Fixed Effects                                                                      |
| FY1          | Forecast Year One (next fiscal year consensus)                                     |
| FY2          | Forecast Year Two (second fiscal year consensus)                                   |
| GAAP         | Generally Accepted Accounting Principles                                           |
| GenAI        | Generative Artificial Intelligence                                                 |
| GICS         | Global Industry Classification Standard                                            |
| HC3          | Heteroskedasticity-Consistent Standard Errors (MacKinnon–White type 3)             |
| IBES         | Institutional Brokers’ Estimate System                                             |
| IFRS         | International Financial Reporting Standards                                        |
| IPO          | Initial Public Offering                                                            |
| IPS          | Im–Pesaran–Shin (panel unit-root test)                                             |
| M&A          | Mergers and Acquisitions                                                           |
| MDE          | Minimum Detectable Effect                                                          |
| MSc          | Master of Science                                                                  |
| NBER         | National Bureau of Economic Research                                               |
| OLS          | Ordinary Least Squares                                                             |

|           |                                                                            |
|-----------|----------------------------------------------------------------------------|
| OOS       | Out-of-Sample                                                              |
| PostShock | Post-Shock indicator (1 from the candidate break date onward, 0 otherwise) |
| Q1–Q4     | First through Fourth Quarter of a calendar year                            |
| R&D       | Research and Development                                                   |
| SaaS      | Software-as-a-Service                                                      |
| SBC       | Stock-Based Compensation                                                   |
| SEC       | U.S. Securities and Exchange Commission                                    |
| TTM       | Trailing Twelve Months                                                     |
| U.S.      | United States                                                              |
| USD       | United States Dollar                                                       |
| WACC      | Weighted Average Cost of Capital                                           |
| YoY       | Year-over-Year                                                             |

## 1. Introduction

A central question in finance is whether the mapping from observable fundamentals to firm value remains stable across structural change. For Software-as-a-Service firms, an archetypal intangible-intensive, long-duration business model, the question has acquired new urgency. Practitioners value SaaS firms through enterprise-value-to-revenue multiples and three operating drivers: revenue growth, operating margin, and size (Feld 2015; Bessemer Venture Partners 2025). The framework rests on the implicit claim that those drivers, properly weighted, summarize the economic determinants of SaaS firm value. The 2022 repricing puts that claim to the test: did the cross-sectional mapping from drivers to valuation change, and if so, what mechanism best explains the change?

Recent practitioner research has converged on an AI-led account of the change in SaaS valuation. Under a tailwind view, AI strengthens SaaS scalability and data moats, so growth becomes a stronger signal of long-run cash flows and valuation rises (Bain & Company 2024). Under a substitution view, AI agents commoditize point solutions and weaken per-seat pricing, making growth a weaker signal and valuation falling (EQT 2025; Bain & Company 2025). The pivot from the Rule of 40 to the Rule of X (Bessemer Venture Partners 2025), which weights growth at two to three times margin, is the most concrete methodological expression of this account. The duration-channel equity-pricing literature (Lettau and Wachter 2007; Weber 2018; Gormsen and Lazarus 2023) offers a different prediction for the same period: long-duration cash-flow assets are repriced when discount rates rise, with no AI-specific mechanism required. The empirical question this paper addresses is which of these mechanisms best fits the cross-sectional structure of the 2022 SaaS repricing.

I organize the question around three nested stability conditions. The explanatory power of the canonical regression should remain stable; the coefficients on the individual drivers should not change; and conditional on those drivers, no level shift in valuation should remain. Each condition admits a distinct rejection pattern. A coefficient shift without a level shift indicates within-driver reweighting; a level shift without a coefficient shift indicates that the drivers fail to capture some new common factor; a loss of explanatory power indicates that the drivers miss firm-specific information. The two AI views and the duration-channel view yield different signed predictions across these conditions, and different cross-sectional predictions about which firms experience the reweighting and which experience the level shift. The cross-sectional structure of the rejections, more than the rejections themselves, is what discriminates among the three views.

I test these conditions on a panel of 90 U.S.-domiciled SaaS firms over 2014Q1–2024Q4, drawn from the union of the BVP Nasdaq Emerging Cloud Index and a Refinitiv-screened universe verified by 10-K reading. The baseline is a panel regression with firm and calendar-quarter fixed effects. A structural-change specification allows each driver's loading to differ in the post-shock period. The central identification device is a rolling-breakpoint procedure that locates the data-determined break date empirically rather than imposing it at the November 2022 ChatGPT release. An energy-and-utilities placebo bounds the residual rate-cycle channel and characterizes how the duration mechanism manifests in low-AI sectors. Because the post-shock window is short, Section 3 reports an explicit minimum-detectable-effect calculation so that the bounds of any null finding are visible from the design forward.

The hypothesis tests reject all three stability conditions. The cross-sectional valuation function reweights toward operating margin (H2 rejected at the 1 percent level); a negative level shift on incumbent firms is rejected at well below the 1 percent level (H3); and explanatory power is not stable (H1). The rolling-breakpoint procedure yields the central piece of identification: the data-determined break is in 2022Q1, three quarters before the November 2022 ChatGPT release and contemporaneous with the onset of the Federal Reserve's tightening cycle. The cohort heterogeneity tests then explain the shape of each rejection. The margin reweighting is concentrated among firms whose pre-period operating margin was already positive, with the margin loading for the unprofitable cohort indistinguishable from zero. The cross-sectional magnitude of the level shift scales with firm-level cash-flow duration in both profitability cohorts, with duration alone explaining roughly 40% of the cross-sectional residual variation. The energy-and-utilities placebo identifies a separate growth-loading compression in low-AI sectors at the same break, consistent with a duration channel operating across long-duration cash-flow assets and manifesting on whichever dimension of the cross-sectional valuation function carries the identifying variation in each sector.

Two mechanisms drive the post-2022 SaaS repricing: cohort-specific margin reweighting in firms whose pre-period profitability gives their margin signal economic content, and a duration-driven level shift that scales with firm-level cash-flow duration uniformly across cohorts. Both are predictions of the duration-channel column of Table I. Neither AI column predicts the profitability gradient in the margin reweighting, and neither predicts a level shift whose cross-sectional magnitude is explained by pre-period growth and margin combinations rather than by AI-exposure proxies. In a single-sector design with a nine-quarter post-shock window, an AI-specific channel layered atop the duration mechanism is not separately

identifiable. The design returns this constraint as a substantive bound on attribution rather than as evidence against an AI effect.

I make three contributions. First, I extend the value-relevance literature (Barth, Li, and McClure 2023) to the operating-metric layer that practitioners use for SaaS valuation, rather than to GAAP line items, and I provide empirical evidence to inform the practitioner debate over whether SaaS valuation methodology requires revision. The Rule-of-40-to-Rule-of-X pivot documented in practitioner research is observed in the data, but the cohort structure shows that it is best described as a within-cohort margin discrimination consistent with a macro regime change rather than as an AI-driven shift in fundamentals. Second, I provide within-sector evidence on the duration channel identified in the broader equity-pricing literature (Lettau and Wachter 2007; Gormsen and Lazarus 2023), and I document that the channel manifests on different dimensions of the cross-sectional valuation function depending on which driver carries the identifying variation in each sector, growth in low-AI sectors, and margin in SaaS. Third, I bound what cross-sectional valuation tests can identify when a candidate technology shock and a contemporaneous macro regime change are temporally confounded.

The remainder of the paper proceeds as follows. Section 2 reviews the related literature and develops directional predictions under the three competing views. Section 3 describes the research design. Section 4 outlines the sample and data. Section 5 presents results. Section 6 reports robustness. Section 7 discusses implications. Section 8 concludes.

## **2. Literature and Hypotheses**

### **2.1 Literature Review**

#### ***2.1.1 Multiples as reduced-form DCF representations***

The valuation literature establishes that multiples are reduced-form representations of discounted cash flow models, with the level of the multiple determined by expected growth, profitability, and discount rates. Alford (1992) and Kim and Ritter (1999) characterize conditions under which multiple-based valuation produces accurate estimates of firm value. Liu, Nissim, and Thomas (2002) show that operating drivers explain a substantial share of cross-sectional variation in valuation. Bhojraj and Lee (2002) construct warranted multiples by regressing observed multiples on firm characteristics, formalizing the cross-sectional regression approach I adopt. Ohlson (1995) provides the theoretical foundation linking firm value to book value and expected abnormal earnings. The framework implies a stable mapping from fundamentals to the multiple, conditional on a stable economic environment. Whether that mapping is stable around a specific structural break is the empirical question I address.

#### ***2.1.2 The duration channel and the cross-sectional pricing of long-duration cash flows***

The duration-channel literature is the closest theoretical precedent for the present paper. Lettau and Wachter (2007) and Weber (2018) develop the duration channel: growth firms are long-duration assets whose valuations are disproportionately sensitive to changes in expected cash flows and discount rates. Gormsen and Lazarus (2023) refine this insight by separating discount-rate and cash-flow channels in equity pricing and find that the post-2022 repricing of growth assets across the broad equity market is driven primarily by the discount-rate component. Peters and Taylor (2017) and Eisfeldt and Papanikolaou (2013, 2014) demonstrate that intangible and organizational capital explain a substantial share of cross-sectional variation in firm value, and that intangible-intensive firms are systematically long-duration. Hand (2005) shows that for high-growth, intangible-intensive firms, traditional earnings-based measures are often uninformative, motivating the revenue-based multiples used in SaaS valuation. The combined prediction of this literature is sharp: when discount rates rise, long-duration cash-flow assets are repriced downward in proportion to firm-level cash-flow duration, with no firm-specific technology mechanism required to deliver the cross-sectional pattern. SaaS firms, with their long cash-flow durations and intangible-capital intensity, are precisely the universe where this prediction applies best.

### ***2.1.3 Evolution in value relevance***

A large body of literature documents that the relationship between accounting information and firm value evolves with the structure of the economy. Lev and Zarowin (1999) argue that traditional accounting measures have lost relevance as intangible-intensive business models have become more prevalent. Aboody and Lev (1998) and Lev and Sougiannis (1996) show that R&D is value-relevant despite being expensed. Barth, Li, and McClure (2023) provide the most direct precedent for the present paper's value relevance. Using a nonparametric approach, they show that the combined value relevance of accounting information has not declined from 1962 to 2018; rather, the set of value-relevant items has evolved, with intangible-related items, growth-opportunity proxies, and alternative performance measures rising in importance. I apply the same logic at a finer granularity: a single intangible-intensive sector, around a single identifiable structural break, using practitioners' operating metrics rather than GAAP line items. Two design choices distinguish the paper from Barth et al.: I work parametrically because the practitioner-driver set is small and theoretically motivated, whereas their nonparametric approach is designed for a high-dimensional accounting line-item space. I also evaluate the relation around a single identifiable shock, while they document gradual evolution over six decades.

### ***2.1.4 SaaS-specific valuation drivers and the practitioner debate***

Practitioners value SaaS firms primarily using enterprise value-to-revenue multiples, with growth, margin, and size as the canonical operating drivers (Bessemer Venture Partners 2023, 2025). The Rule of 40 (Feld 2015) posits that a healthy SaaS firm should have growth plus profit margin of at least 40 percent. Bessemer Venture Partners (2025) refines the rule into a Rule of X with growth weighted two to three times more than margin, based on regression evidence from public cloud companies. The Rule-of-40-to-Rule-of-X pivot is, in recent practitioner research, attributed to a generative-AI-driven shift in the economics of recurring revenue (Bain & Company 2024, 2025; EQT 2025; Bessemer Venture Partners 2025). I do not test the Rule of 40 directly but estimate the implicit weighting of growth and margin in the cross-sectional valuation function, test whether it changed, and ask which of the competing views developed in Section 2.2 the change is most consistent with. Customer-based valuation metrics, net revenue retention, gross retention, and sales efficiency are central in practice (Gupta, Lehmann, and Stuart 2004) but are non-GAAP, voluntarily disclosed, and not consistently reported across the SaaS universe. I discuss them in Section 7 as an avenue for future research rather than estimate them here.

### ***2.1.5 The practitioner AI-led narrative and its empirical antecedents***

Recent practitioner research has framed the post-2022 SaaS environment in AI-led terms, with two opposite-signed views structuring the discussion. Under a tailwind view, AI amplifies the scalability of SaaS business models, deepens data moats, and reinforces customer lock-in (Bain & Company 2024). Under a substitution view, AI agents threaten per-seat pricing and commoditize point solutions (EQT 2025; Bain & Company 2025). Both views imply that SaaS valuation methodology should be revised in an AI-specific direction, and the Bessemer Rule-of-X pivot (Bessemer Venture Partners 2025) is the most concrete methodological expression of the account. The academic AI-and-firm-value literature provides the empirical antecedents for testing such a claim. Eisfeldt, Schubert, and Zhang (2023) construct a firm-level measure of workforce exposure to generative AI and document that high-exposure firms experience excess returns following ChatGPT's release. Babina, Fedyk, He, and Hodson (2024) show that firms investing in AI experience higher growth and product innovation. Both papers focus on average effects and rely on firm-level AI-exposure measures that are not consistently available across the SaaS universe; I discuss the implications of this constraint for what the design can and cannot identify in Sections 3 and 7. Brynjolfsson, Rock, and Syverson (2021) develop the general-purpose-technology framework for AI; Brynjolfsson, Li, and Raymond (2023) document direct productivity gains from generative AI; Acemoglu (2024) argues that the macroeconomic productivity gains are likely modest. These views imply different priors on whether AI should reweight valuation fundamentals at the firm level. Kogan, Papanikolaou, Seru, and Stoffman (2017) link firm-level technological innovation, measured by patents, to firm value and economic growth, thereby motivating the broader question of whether technological shocks reprice firm value.

## **2.2 Hypotheses Development**

### ***2.2.1 A simple Gordon-growth scaffold***

To anchor the directional predictions, consider a stylized Gordon-growth representation of the EV/Revenue multiple. Let cash flow per unit of revenue be the operating margin  $m$ , the cost of capital be  $r$ , and the long-run growth rate of revenue be  $g$ . The price-to-revenue ratio satisfies, as a steady-state identity:

$$EV/Revenue \approx m / (r - g)$$

Taking logs and differentiating yields a sensitivity decomposition. The log-elasticity of the multiple with respect to margin is one; the log-sensitivity to long-run growth is  $1/(r-g)$ ; and the

log-sensitivity to the cost of capital is  $-1/(r-g)$ . The growth and discount-rate sensitivities scale with  $1/(r-g)$ , which is large for long-duration, high-growth firms. This delivers two predictions used below. First, under the duration-channel view, a rate-driven rise in  $r$  mechanically lowers the multiple by an amount proportional to firm-level cash-flow duration; the rolling-breakpoint and placebo procedures of Section 3 are the design's direct probes of this mechanism. Second, under the AI views, AI changes the persistence of  $g$  (positively under the tailwind view, negatively under the substitution view); the loading on observed growth shifts accordingly because trailing growth becomes an informative signal of future cash flows. Margin and size predictions follow analogously. The framework also makes the AI vs. duration confound visible: both AI views and the duration-channel view predict directional shifts in the cross-sectional valuation function, and the discriminating evidence lies in the cross-sectional structure of the rejection rather than in the rejection itself.

### ***2.2.2 Competing views***

I consider three views of how the post-2022 environment may have repriced public SaaS, and I treat them as alternatives to the null of a stable cross-sectional valuation function. Two are AI-led views, advanced in the practitioner literature, that posit opposite-signed shifts; the third is a duration-channel view, drawn from the equity-pricing literature, that operates through the contemporaneous macro regime change.

**The tailwind view** holds that AI amplifies the scalability of SaaS business models, deepens data moats, and reinforces customer lock-in (Bain & Company 2024). Under this view, the loading on growth rises because growth becomes a more credible signal of long-run cash flows. The margin loading may also rise, as AI-enabled cost reductions become value-relevant. The valuation level should rise for the SaaS sector.

**The substitution view** holds the opposite. AI agents threaten per-seat pricing and commoditize point solutions, weakening the link between current customer relationships and future cash flows (EQT 2025; Bain & Company 2025). Under this view, the loading on growth falls because trailing growth becomes a less credible signal of future cash flows. The loading on size may rise as the scale becomes a defensive asset. The level of valuation should fall.

**The duration-channel view** is drawn directly from the equity-pricing literature on the cross-sectional repricing of long-duration cash-flow assets (Lettau and Wachter 2007; Weber 2018; Gormsen and Lazarus 2023). The post-ZIRP reset of growth-stock valuations, documented across the broad equity universe, predicts that long-duration cash-flow assets, firms whose value depends most on distant cash flows, experience the largest valuation compressions when

discount rates rise. Within SaaS, this view predicts that the cross-sectional valuation function reweights toward operating margin specifically among firms whose pre-period margin profile carries economic information, and that any unexplained level shift tracks firm-level cash-flow duration rather than firm-level AI exposure. The duration-channel view does not require an AI mechanism to explain the post-2022 SaaS repricing; whether an AI mechanism is identifiable on top of the duration channel is a separate question that this design is not powered to answer.

The three views are not mutually exclusive: the data may support one, more than one (perhaps along different cross-sections), or none. Table I records the signed predictions of each view across the three hypotheses developed below. The cross-sectional structure of each rejection (Section 5.5) is the device that adjudicates among the columns; the pooled rejection pattern alone is not sufficient because the substitution-view and duration-channel-view predictions for H2 and H3 overlap in sign and differ only in cross-sectional structure. The contribution of the design is to provide evidence on the cross-sectional structure of rejection, which the pooled rejection-pattern test cannot provide on its own.

### ***2.2.3 Hypothesis 1: Stability of explanatory power***

H1 isolates whether the canonical driver set continues to span the value-relevant information for SaaS firms. Stability of explanatory power is consistent with reweighting within the existing variables; a decline indicates that an unobserved factor introduces value-relevant information not captured by the existing variables, a Barth-Li-McClure-style augmentation outcome, while an increase indicates that the existing variables have become more informative because the underlying signal has sharpened. Table I records the signed prediction of each view.

*H1<sub>0</sub>: The adjusted  $R^2$  of the valuation regression does not differ between the pre-shock period and the post-shock period.*

*H1<sub>a</sub>: The adjusted  $R^2$  differs between the two periods.*

### ***2.2.4 Hypothesis 2: Stability of coefficients***

If the structural break changes the economics of recurring revenue, the loadings linking valuation to growth, margin, and size should change. The Gordon-growth scaffold of Section 2.2.1 maps each view onto signed predictions for each driver; Table I records the matrix. The growth loading is the sharpest discriminator across the two AI views; tailwind predicts a positive shift, substitution a negative shift, and the two are mutually exclusive. The duration-channel view predicts a margin shift concentrated in firms with informative pre-period margin profiles, with the growth and size loadings unchanged on average (in the SaaS sample) and a growth-loading compression operating in sectors where growth carries the identifying variation.

**H2<sub>0</sub>:** *The vector of coefficients on the canonical operating drivers does not differ between the pre- and post-shock periods.*

**H2<sub>a</sub>:** *At least one coefficient differs across periods.*

### 2.2.5 Hypothesis 3: Absence of unexplained level shift

If the canonical drivers fully capture the determinants of SaaS firm value, no systematic shift in valuation across periods should remain after controlling for fundamentals. A nonzero level shift indicates that an unobserved common factor not reflected in the canonical drivers affects valuation in the post-shock period. The sign of the H3 mean residual is the second discriminator across the AI views, alongside the sign of the H2 growth-loading shift. The duration-channel view predicts a negative residual whose magnitude tracks firm-level cash-flow duration; Table I records the predictions across all three views.

**H3<sub>0</sub>:** *Conditional on the canonical operating drivers and on the pre-shock estimated firm and time effects, the mean of the post-shock out-of-sample residuals is zero.*

**H3<sub>a</sub>:** *The mean of the post-shock out-of-sample residuals is nonzero.*

**Table I. Directional predictions across hypotheses and views.**

*The matrix below summarizes the signed predictions of the tailwind, substitution, and duration-channel views for each of the three hypotheses. Cells where the columns diverge identify hypotheses with discriminative power across views; cells where two or more columns agree identify hypotheses whose pooled rejection-pattern test cannot, on its own, adjudicate between the agreeing views, and for which the cross-sectional structure of the rejection (Section 5.5) provides additional resolution. The substitution-view and duration-channel-view predictions for H2 (margin element) and H3 (sign) overlap; the cohort-heterogeneity and duration cross-section tests of Section 5.5 are the design's direct probes of the additional cross-sectional predictions of the duration-channel view that the AI views do not make.*

| Hypothesis                         | Tailwind view                                                                                                       | Substitution view                                                                                                                    | Duration-channel view                                                                                                                     |
|------------------------------------|---------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------|
| <b>H1: R<sup>2</sup> stability</b> | R <sup>2</sup> stable or rises (drivers remain or become more informative)                                          | R <sup>2</sup> falls (trailing growth becomes a less informative signal)                                                             | R <sup>2</sup> stable (reweighting within existing variables, not augmentation)                                                           |
| <b>H2: coefficient stability</b>   | $\gamma_{\text{growth}} > 0$ ; $\gamma_{\text{margin}}$ stable or modestly $> 0$ ; $\gamma_{\text{size}}$ unchanged | $\gamma_{\text{growth}} < 0$ ; $\gamma_{\text{margin}} > 0$ (Rule-of-X reweighting); $\gamma_{\text{size}} > 0$ (resilience premium) | $\gamma_{\text{margin}} > 0$ , concentrated in firms with informative pre-period margin; growth and size unchanged on average within SaaS |
| <b>H3: level shift</b>             | Mean residual $> 0$ (positive unobserved factor)                                                                    | Mean residual $< 0$ (compressed expectations vs. fundamentals)                                                                       | Mean residual $< 0$ , with magnitude tracking firm-level cash-flow duration                                                               |

### 3. Research Design

#### 3.1 Baseline Specification

The baseline regresses log EV/Revenue on the canonical operating drivers with firm and calendar-quarter fixed effects:

$$\ln(EV/Revenue)_{i,t} = \alpha_i + \lambda_t + \beta'X_{i,t} + \varepsilon_{i,t} \quad (1)$$

The vector  $X_{i,t}$  contains year-over-year revenue growth, trailing-twelve-month GAAP operating margin (with stock-based compensation included as an expense), and the natural logarithm of trailing-twelve-month revenue (USD millions). Variable definitions and the rationale for these specific operationalizations appear in Appendix A.  $\alpha_i$  denotes firm fixed effects;  $\lambda_t$  denotes calendar-quarter fixed effects. The use of log EV/Revenue addresses the right-skewness and heavy upper tail of SaaS multiples. All operating drivers are winsorized at the 1st and 99th percentiles within the SaaS sample. Firm fixed effects absorb time-invariant differences across firms. Calendar-quarter fixed effects absorb aggregate shocks affecting all firms in each quarter, including macroeconomic conditions, interest rates, and market-wide valuation changes. Identification rests on within-firm variation across time and on the cross-sectional dispersion in fundamentals at each point in time. Standard errors are double-clustered at the firm and calendar-quarter levels (Petersen 2009).

A simultaneity caveat applies. EV/Revenue is a forward-looking market price, while Growth and Margin are trailing-twelve-month realizations. Rational investors price expected fundamentals, so trailing metrics proxy for expectations with error; this measurement error is likely more severe in the post-shock period, when analyst growth forecasts may diverge sharply from recent realized performance. Section 6, therefore, includes a robustness specification that replaces trailing Growth with the median one-year-forward revenue-growth consensus from Refinitiv IBES, where available, and reports whether the main structural-change results are robust to this forward-looking regressor. Coverage limitations, consensus estimates are available for a smaller fraction of the sample than trailing realized Growth and are documented in that section.

#### 3.2 Structural-Change Specification

To test H2, I extend the baseline to allow each driver's loading to differ in the post-shock period:

$$\ln(EV/Revenue)_{i,t} = \alpha_i + \lambda_t + \beta'X_{i,t} + \delta'(X_{i,t} \times PostShock_t) + \varepsilon_{i,t} \quad (2)$$

PostShock<sub>t</sub> equals one for quarters from the candidate break date onward and zero otherwise. The primary specification uses the data-determined break date of 2022Q1, identified by the rolling-breakpoint procedure of Section 3.4. The 2022Q4 specification, motivated by the November 30, 2022, ChatGPT release, is reported alongside as a benchmark against the practitioner narrative; under the duration-channel view, the 2022Q4 specification is expected to deliver an attenuating signal of the same underlying coefficient shift identified at 2022Q1. The vector  $\delta$  captures changes in driver loadings across periods; H2<sub>0</sub> is tested by the joint significance of  $\delta$ . Under the tailwind view, the growth element of  $\delta$  is positive; under the substitution view, negative. Under the duration-channel view, the growth element of  $\delta$  is, on average, unsigned in SaaS, and the margin element is positive but concentrated in firms whose pre-period margin profile carries economic information; the cohort-heterogeneity tests in Section 5.5 are the design's direct probe of this view. Equation (2) deliberately omits a stand-alone PostShock main effect: with the full set of calendar-quarter fixed effects  $\lambda_t$ , the period indicator is collinear with the union of the post-break quarter dummies and is not separately identified. The test for an unexplained level shift (H3) is therefore conducted separately in Section 3.3.

### 3.3 Testing H3: Out-of-Sample Residual Procedure

H3<sub>0</sub> is tested via a sequenced out-of-sample procedure. Equation (1) is estimated on the pre-shock subsample only (2014Q1–2022Q3 at the 2022Q4 specification, or 2014Q1–2021Q4 at the 2022Q1 specification), yielding loadings  $\beta$ , firm effects  $\alpha_i$ , and quarter effects  $\lambda_t$  anchored exclusively to pre-period information. Out-of-sample residuals are then computed for all firm-quarter observations in the post-shock window:

$$\hat{\varepsilon}_{i,t}^{oos} = \ln(EV/Revenue)_{i,t} - \hat{\alpha}_i - \bar{\lambda}_{pre} - \beta'X_{i,t}$$

where  $\bar{\lambda}_{pre}$  is the mean of the estimated pre-period quarter effects. A t-test of the mean of these residuals against zero, with standard errors clustered at the firm level, constitutes the test of H3<sub>0</sub>. Estimating  $\beta$  on the full sample would pull the estimator toward post-period data, compressing the post-window residuals and biasing the level-shift test toward non-rejection, particularly when H2 is also violated. The out-of-sample procedure avoids that confound.

The procedure is restricted to incumbent firms with at least eight pre-period quarters of data, so that  $\alpha_i$  is identified from a meaningful pre-period anchor. Firms entering the panel after 2014Q1 with fewer than eight pre-period quarters are excluded from the H3 test; their treatment is reported as a robustness exercise. The sign of the estimated mean residual adjudicates

between the tailwind view (positive) and the substitution and duration-channel views (negative); the cross-sectional regression of firm-level mean residuals on a duration proxy, reported in Section 5.5, is the design’s probe of the additional cross-sectional prediction made by the duration-channel view. The driver interactions  $\delta$  in equation (2) remain identified separately and are unaffected by this restriction.

### 3.4 Identification: Rolling-Breakpoint and Placebo

The post-2022 period coincides with sharply rising interest rates and a market-wide rotation away from long-duration assets. SaaS firms, with their long cash-flow durations, are among the firms most affected. The calendar-quarter fixed effects in equations (1) and (2) absorb the average rate-cycle effect on all firms in each quarter. Within-SaaS heterogeneity in rate sensitivity could still confound  $\delta$  if, for example, high-growth SaaS firms are more rate-sensitive than mature ones. Two procedures bound this concern: a rolling-breakpoint test that locates the data-determined break date (the central identification device), and an energy-and-utilities placebo that bounds the residual rate-cycle channel and characterizes how the duration mechanism manifests in low-AI sectors. A robustness specification with a rate-cycle triple interaction is reported in Section 6.

**Rolling-breakpoint procedure.** The choice of any specific break date is theoretically motivated but empirically arbitrary. I re-estimate equation (2) with the breakpoint set at every calendar quarter from 2021Q1 through 2023Q4 and plot the F-statistic for joint significance of  $\delta$  against the candidate break date. The test statistic should peak at the true structural break. A peak at 2022Q1 (rate-cycle onset) is consistent with the duration-channel view; a peak at 2022Q4 (the ChatGPT release) is consistent with an AI-driven structural change; a peak in mid-2023 would correspond to the Nvidia-led market AI rally. The peak identified in the data is the central result of the paper. I prefer this procedure to a formal Bai and Perron (1998) endogenous-break test because the post-shock window is short, and conventional Bai-Perron trimming rules absorb most of the candidate range. I supplement the rolling F with a CUSUM-of-squares fluctuation test (Zeileis, Leisch, Hornik, and Kleiber 2002), which is better calibrated for end-of-sample breaks.

**Energy-and-utilities placebo.** U.S.-domiciled firms in GICS sectors 10 (energy) and 55 (utilities) have low AI exposure but were exposed to the same post-2021 monetary tightening as SaaS firms. I estimate equation (2) on this placebo sample. A significant  $\delta$  in the placebo identifies a duration-channel mechanism operating outside SaaS; the cross-sectional dimension on which the mechanism manifests in low-AI sectors (margin, growth, or size) provides direct

characterization of the duration channel. A null in the placebo would not, on its own, rule out within-SaaS rate heterogeneity, which is why the rolling-breakpoint test is the central identification device rather than the placebo.

### 3.5 Power and Minimum Detectable Effect

The post-shock window contains nine quarters. Power is therefore a binding design constraint, and I build an explicit minimum detectable effect (MDE) calculation into the design. Following Bloom (1995), the MDE for the growth element of  $\delta$  at 80 percent power and 5 percent size is approximately:

$$MDE\_g \approx 2.8 \cdot \hat{\sigma}_\varepsilon / \sqrt{(N\_post \cdot Var(Growth))}$$

where  $\hat{\sigma}_\varepsilon$  is the residual standard deviation from equation (1) on the pre-period sample,  $N\_post$  is the number of post-shock firm-quarter observations, and  $Var(Growth)$  is the within-firm cross-sectional variance of Growth in the post-period. Calibration with  $\hat{\sigma}_\varepsilon \approx 0.35$ ,  $N\_post \approx 800$ , and within-firm  $Var(Growth) \approx 0.02$  in the post-period yields an MDE of approximately 0.22 in log-multiple units per unit of growth. The relevant variance here is the within-firm demeaned variance, which is materially smaller than the pooled cross-sectional variance because Growth is persistent within firm. An MDE of 0.22 is equivalent to a structural shift in the growth loading large enough to move the implied warranted multiple by roughly 25 percent for a firm with 30 percent revenue growth. The MDE is therefore conservative: the design is well-powered against large reweightings of the canonical drivers but underpowered against shifts smaller than approximately one quarter of the pre-period growth-loading magnitude. Effects smaller than this are not reliably detectable, and the H1/H2/H3 conclusions are framed accordingly. The MDE for the H3 mean residual is approximately 0.08 log units, or an 8-percent unexplained shift in the warranted multiple at the post-period mean.

### 3.6 Estimation

All regressions are estimated by ordinary least squares with firm and calendar-quarter fixed effects. Standard errors are double-clustered at the firm and calendar-quarter levels except for the H3 test, where firm-level clustering is used because the test averages residuals across quarters within firms. The rolling-breakpoint procedure sweeps candidate break dates from 2021Q1 through 2023Q4. Because EV/Revenue multiples for SaaS firms are persistent over time, a spurious-regression concern arises if  $\log(EV/Revenue)$  and the regressors are non-stationary. Section 6.4 reports a long-differences specification on the within-firm pre-to-post

cross-section that serves as both a non-stationarity check and a corroborating reduced-form test of the structural-change hypothesis.

### **3.7 What the Design Can and Cannot Identify**

The design is well-suited to delivering three pieces of identification. First, whether the cross-sectional valuation function changed across the 2022 break (H1–H3 jointly). Second, when it changed, in the sense of locating the data-determined break date among the candidate dates 2021Q1–2023Q4 (rolling-breakpoint procedure). Third, the cross-sectional structure of the change, which firms experienced the margin reweighting, which firms experienced the level shift, and how each scales with firm-level cash-flow duration (Section 5.5). The cross-sectional structure is the device by which the design adjudicates among the columns of Table I.

The design is not powered to deliver one piece of identification: whether a generative-AI-specific channel operates on top of the duration channel, or whether the cross-sectional structure of the rejection most parsimoniously fits. Doing so would require a firm-level AI-exposure measure that is not collinear with cash-flow duration, which is the binding constraint that the present design does not relax. Section 7 returns to this point and discusses what would be required of a future design to relax it. The within-SaaS sample-construction obstacles are documented there. The implication for the present paper is that AI-specific attribution is bounded above by what the cross-sectional structure of the rejection requires; where the structure is parsimoniously explained by the duration channel without an AI mechanism, the design returns this constraint as a substantive bound rather than as a null finding.

### **3.8 Mapping to Hypotheses**

H1<sub>0</sub> is tested by the F-test on joint interaction significance in equation (2) and by the comparison of adjusted  $R^2$  across subsamples. H2<sub>0</sub> is tested by the t-statistics on individual elements of  $\delta$ ; the sign adjudicates between the two AI views. H3<sub>0</sub> is tested by the out-of-sample residual procedure of Section 3.3. The energy-and-utilities placebo and the rolling-breakpoint procedure together discipline the interpretation. The duration-channel view's predictions for H2 and H3 overlap with the substitution view in sign but differ in cross-sectional structure; the cohort-heterogeneity and duration cross-section tests of Section 5.5 provide the additional resolution that the rejection-pattern map alone does not deliver.

## 4. Data and Sample

### 4.1 Sample Construction

The primary sample consists of U.S.-domiciled Software-as-a-Service firms from 2014Q1 through 2024Q4. The sample begins in 2014 to coincide with the popularization of the Rule of 40 (Feld 2015) and the maturation of the SaaS public-market universe. The additional pre-shock years provide identifying variation for the rate-cycle channel and reduce the dependence of the structural-change inference on the limited post-shock window.

I construct the SaaS sample as the union of two sources. The first is the current constituent list of the BVP Nasdaq Emerging Cloud Index (BVP.NASDAQEMCLOUD), the practitioner-standard list of U.S.-listed cloud-native software firms maintained by Bessemer Venture Partners and Nasdaq. The index methodology selects firms whose primary business is delivering cloud-based software with a majority of recurring revenue, and it is widely cited in investor research as the benchmark SaaS universe. I use the current snapshot because the historical quarter-by-quarter constituent list is not publicly available; the implications of this choice for survivorship are addressed in the robustness specifications of Section 6.5. The second source is a curated set of U.S.-domiciled software firms identified through a Refinitiv Workspace screen, TRBC Software industry, U.S.-domiciled (Form 10-K filer with U.S. country of incorporation, excluding Form 20-F and 40-F foreign private issuers), active or formerly active 2014–2024, market capitalization above \$100 million, and verified to be subscription-revenue businesses by hand-reading the most recent pre-shock 10-K (fiscal year 2021). A firm passes the verification step if its annual filing reports recurring or subscription revenue (typically labeled “subscription revenue” or “subscription services revenue” in the income statement or revenue-disaggregation footnote) exceeding 70 percent of total revenue. Anchoring verification to the fiscal-year 2021 10-K ensures that SaaS classification is pre-determined relative to both the 2022Q1 rate-cycle onset and the 2022Q4 ChatGPT release, eliminating the risk that firms that altered their business model after either candidate break are endogenously included or excluded in the post-shock window. The U.S.-domicile filter is also applied to firms entering the sample through the BVP Nasdaq Emerging Cloud Index, since the index admits foreign-domiciled cloud firms (Canadian, Israeli, Dutch, and other foreign private issuers), which would introduce country-specific factors and IFRS-versus-GAAP accounting differences into the cross-section. I exclude Atlassian Corporation on the additional ground that its 2022–2023 redomiciliation from the United Kingdom to Delaware places its U.S. domicile classification within the post-shock window, violating the pre-determination principle. The two

sources are deduplicated by ticker, yielding a final SaaS sample of 90 firms and approximately 2,468 firm-quarter observations over 2014Q1–2024Q4. Firms are excluded if they are not U.S.-domiciled (operationalized as filing the SEC Form 20-F or 40-F rather than the Form 10-K), have fewer than four quarters of trading history at the start of the sample period, or are subject to a takeover bid or delisting event in the focal quarter.

This two-source design offers two advantages over a purely quantitative threshold applied uniformly across the commercial database software universe. First, the BVP Nasdaq Emerging Cloud Index applies a consistent, publicly documented index methodology and serves as the practitioner-standard SaaS perimeter, thereby making one half of the sample frame transparent and externally validated. Second, the recurring-revenue field is not consistently disclosed in commercial databases; a uniform quantitative threshold computed from database fields alone would mechanically exclude firms with sparse or non-standardized disclosure, while the 10-K-based verification step recovers such firms when their filings document a majority-recurring-revenue business model. The 70 percent threshold is stricter than the “majority recurring revenue” criterion underlying the BVP index methodology and is intended to focus the verifiable subsample on pure-play SaaS rather than hybrid software firms; I hold this threshold fixed throughout the analysis.

## **4.2 Time Period and Candidate Break Dates**

The sample spans 44 calendar quarters from 2014Q1 through 2024Q4. I identify two narrative candidate break dates *ex ante*. The first is 2022Q1, the onset of the Federal Reserve tightening cycle, motivated by the duration-channel view. The second is 2022Q4, coinciding with the November 30, 2022, release of ChatGPT, which was motivated by the AI-led practitioner narrative and treated in the literature as the first widely publicized generative AI event affecting equity valuations (Eisfeldt, Schubert, and Zhang 2023). The primary structural-change specification uses the data-determined break (2022Q1, identified by the rolling-breakpoint procedure of Section 5.3); the 2022Q4 specification is reported alongside as a benchmark against the practitioner narrative. At the 2022Q1 specification, the panel splits into 32 pre-shock and 12 post-shock quarters; at the 2022Q4 specification, 35 pre-shock and 9 post-shock quarters.

## **4.3 Placebo Sample**

The placebo sample consists of U.S.-domiciled firms in GICS sectors 10 (energy) and 55 (utilities), restricted to the same 2014Q1–2024Q4 period and subject to the same data-

availability filters as the SaaS sample, including the Form 20-F and 40-F exclusions; foreign-domiciled energy majors that trade on U.S. exchanges (Shell, BP, TotalEnergies, Equinor, and similar) are therefore not part of the placebo. These firms have low AI exposure by sector construction but were subject to the same post-2021 monetary tightening as SaaS firms, so they identify a duration-channel mechanism operating in low-AI sectors and characterize the cross-sectional dimension on which it manifests. Firm-level AI-exposure measures based on occupational task content (e.g., Eisfeldt, Schubert, and Zhang 2023) are not consistently available for the full SaaS universe in Refinitiv and are therefore not used as a primary stratification variable; instead, I use the sector-based assignment as a tractable and transparent proxy. The realized placebo sample contains 217 firms (143 in GICS sector 10 Energy, 74 in GICS sector 55 Utilities) and approximately 7,514 firm-quarter observations, of which 5,785 fall in the pre-shock window and 1,729 in the post-shock window. Two energy-specific data-quality filters that do not bind to the SaaS sample are applied here. First, I drop three firm-quarters with negative reported quarterly revenue, since under upstream-oil accounting the revenue line is net of mark-to-market gains and losses on commodity-derivative hedges and can turn negative during sharp commodity drawdowns (notably Northern Oil and Gas in 2019Q2 and 2020Q3 and CNX Resources in 2022Q2); these observations represent hedge-accounting artefacts rather than the operating revenue that the EV/Revenue multiple is designed to capture. Second, I drop firm-quarters with trailing-twelve-month revenue below \$20 million at the firm-quarter level, since at this scale the firm is typically pre-commercial or sub-commercial and its valuation is driven by the option value of reserves or pipeline projects rather than by recurring operating revenue. This floor is the operating-business analog to the \$100 million market-capitalization floor applied at screen time and removes 885 firm-quarters concentrated in a small set of exploration-stage and early-stage firms (Comstock, Gevo, American Resources, PureCycle, PEDEVCO, and similar). I apply the filter at the firm-quarter level rather than at the firm level so that firms that subsequently grew past the threshold remain in the sample for their commercial-scale quarters only.

#### **4.4 Data Sources**

I obtain financial data from Refinitiv Workspace. Market values, debt, and cash balances are from Refinitiv at quarter-end. Quarterly revenue, operating income, and segment information are from Refinitiv quarterly fundamentals. Forward consensus revenue growth estimates used in the Section 6 robustness specification are from Refinitiv IBES (TR.RevenueMean field). Treasury yields are from the Federal Reserve H.15 release.

Enterprise value is constructed as market capitalization plus total debt minus cash and cash equivalents, measured at quarter-end. Trailing-twelve-month revenue is the sum of revenue over the four quarters ending in the focal quarter. I winsorize EV/Revenue multiples at the 1st and 99th percentiles computed on the full pooled SaaS sample, with thresholds fixed ex ante before any subperiod estimation. Using a single pooled distribution rather than period-specific thresholds ensures that a structural shift in the upper tail of multiples post-2022 is not mechanically absorbed by the winsorization step. I winsorize operating drivers separately on the same pooled distribution. A footnote in Table 1 reports the number of observations affected by winsorization in each subperiod; a material asymmetry would indicate the winsorization step is non-neutral across periods.

## 4.5 Operating Metrics

The main specification uses year-over-year revenue growth, GAAP operating margin (including stock-based compensation as an expense), and the natural logarithm of trailing twelve-month revenue. The choice of GAAP rather than non-GAAP operating margin and the use of  $\log(\text{Revenue})$  rather than  $\log(\text{MarketCap})$  or  $\log(\text{EV})$  for size are documented and justified in Appendix A. The central considerations are cross-firm and cross-period comparability, sample coverage, and the avoidance of mechanical correlation between size and the EV-based dependent variable. All three drivers are available as structured fields for essentially all firm-quarters in Refinitiv. Two SaaS-specific operating metrics often discussed in practitioner research, net revenue retention and sales efficiency, are voluntarily disclosed and not consistently reported across the SaaS universe. Hand-collection of quarterly earnings releases would yield a substantially smaller, selected sample driven by disclosure choice. For that reason, I do not estimate these metrics here and flag them in Section 7 as an avenue for future research on a hand-collected sample.

## 4.6 Descriptive Statistics

Table 1 reports summary statistics for the SaaS regression sample. Panel A reports the distribution of the dependent variable and the three operating drivers across all 2,094 firm-quarters that enter the panel regressions of equations (1) and (2): the mean, standard deviation, 25th percentile, median, and 75th percentile for  $\ln(\text{EV}/\text{Revenue})$ , Growth, Operating margin, and  $\ln(\text{Revenue})$ . Panel B reports the firm-quarter sample construction waterfall from the initial universe of 2,685 firm-quarters to the final regression sample, showing the number of firm-quarters dropped at each filtering step.

**Table 1. Summary statistics for the SaaS regression sample, 2014Q1–2024Q4.**

*This table reports summary statistics for the SaaS regression sample. Panel A reports the distribution of the dependent variable and the three operating drivers across all 2,094 firm-quarters that enter the panel regressions of equations (1) and (2). The dependent variable is  $\ln(\text{EV}/\text{Revenue})$ , the natural logarithm of enterprise value divided by trailing-twelve-month revenue. Growth is year-over-year revenue growth, operating margin is trailing-twelve-month GAAP operating margin (with stock-based compensation treated as an operating expense), and  $\ln(\text{Revenue})$  is the natural logarithm of trailing-twelve-month revenue measured in U.S. dollars. All four variables are winsorized at the 1st and 99th percentiles of the pooled SaaS distribution; thresholds are fixed ex ante so that a structural shift in the upper tail of multiples post-2022 is not mechanically absorbed by the winsorization step. Panel B reports the firm-quarter sample construction waterfall from the initial universe to the final regression sample, showing the number of firm-quarters dropped at each filtering step.*

**Panel A. Regression-sample variable distributions**

| Variable                        | N     | Mean   | SD    | P25    | Median | P75    |
|---------------------------------|-------|--------|-------|--------|--------|--------|
| $\ln(\text{EV}/\text{Revenue})$ | 2,094 | 2.004  | 0.747 | 1.464  | 2.063  | 2.487  |
| Growth                          | 2,094 | 0.223  | 0.178 | 0.107  | 0.204  | 0.316  |
| Operating margin                | 2,094 | 0.121  | 0.175 | 0.008  | 0.130  | 0.247  |
| $\ln(\text{Revenue})$           | 2,094 | 20.367 | 1.346 | 19.570 | 20.327 | 21.244 |

**Panel B. Sample-construction waterfall (firm-quarters)**

| Stage                                                                | Firm-quarters | Dropped |
|----------------------------------------------------------------------|---------------|---------|
| Initial firm-quarter universe (2014Q1–2024Q4)                        | 2,685         | —       |
| After non-missing enterprise value                                   | 2,654         | 31      |
| After non-missing revenue and operating margin                       | 2,339         | 315     |
| After non-missing year-over-year growth                              | 2,304         | 35      |
| Final regression sample (winsorized; eight-quarter incumbent filter) | 2,094         | 210     |

*Notes. The regression sample of 2,094 firm-quarters covers 86 distinct SaaS firms from a sample frame of 90; the four firms not represented in the regression sample lack the eight pre-period quarters required for the out-of-sample H3 procedure or the lagged growth observation required by year-over-year growth. The post-shock subsample (2022Q1–2024Q4) contains 753 firm-quarters at the 2022Q4 break and 955 firm-quarters at the 2022Q1 break; pre-shock subsamples contain 1,341 and 1,139 firm-quarters, respectively. Pre-/post-period descriptive comparisons of the underlying drivers are not reported because the structural-change tests in Tables 3 and 4 directly identify shifts in driver loadings and in the level of the cross-sectional valuation function, which are the economically meaningful quantities; descriptive shifts in driver levels are absorbed by the firm and calendar-quarter fixed effects in the primary specifications.*

## 5. Empirical Results

### 5.1 Baseline Cross-Sectional Mapping

Table 2 reports estimates of equation (1) on the full SaaS panel with firm and calendar-quarter fixed effects, alongside specifications that drop firm fixed effects and that drop calendar-quarter fixed effects, to identify the share of explanatory power attributable to within-firm versus cross-firm variation.

In the baseline specification with both firm and calendar-quarter fixed effects, the coefficient on Growth is 1.479 ( $t = 6.95$ ,  $p < 0.001$ ), the coefficient on Margin is 0.107 ( $t = 0.36$ ,  $p = 0.72$ ), and the coefficient on  $\log(\text{Revenue})$  is  $-0.096$  ( $t = -1.10$ ,  $p = 0.27$ ). The within-firm  $R^2$  is 0.169. The pooled growth and margin loadings reflect the within-firm cross-sectional valuation function: a one-percentage-point increase in within-firm trailing growth is associated with a 1.48-log-point higher EV/Revenue multiple, while a within-firm increase in operating margin is statistically indistinguishable from zero. The implicit growth-to-margin weighting in the pooled within-firm SaaS valuation function is therefore approximately 14 $\times$ ; within-firm growth is the dominant operating driver, and within-firm margin movements are essentially unpriced in the pooled sample. This pattern is consistent with the prior literature on SaaS valuation (Bessemer Venture Partners 2025) and serves as the baseline against which the structural-change tests in Section 5.2 are evaluated.

**Table 2. Baseline cross-sectional valuation function for U.S. SaaS firms, 2014Q1–2024Q4.**

*This table reports OLS estimates of equation (1),  $\ln(EV/Revenue) = \alpha_i + \lambda_t + \beta'X_{i,t} + \varepsilon_{i,t}$ , where  $X$  contains revenue growth (year-over-year), trailing-twelve-month GAAP operating margin, and the natural logarithm of trailing-twelve-month revenue (USD millions). Column (1) reports pooled OLS without fixed effects. Column (2) adds firm fixed effects only. Column (3) adds both firm and calendar-quarter fixed effects (the baseline specification). Standard errors are double-clustered at the firm and calendar-quarter levels and reported in parentheses below each coefficient.*

|                            | (1)                 | (2)                 | (3)                 |
|----------------------------|---------------------|---------------------|---------------------|
|                            | Pooled OLS          | Firm FE             | Firm + Quarter FE   |
| Growth                     | 2.622***<br>(0.280) | 1.814***<br>(0.306) | 1.479***<br>(0.213) |
| Operating margin           | 0.635*<br>(0.341)   | 0.634<br>(0.406)    | 0.107<br>(0.297)    |
| ln(Revenue)                | 0.101*<br>(0.053)   | 0.072<br>(0.094)    | −0.096<br>(0.087)   |
| <i>Firm FE</i>             | No                  | Yes                 | Yes                 |
| <i>Calendar-quarter FE</i> | No                  | No                  | Yes                 |
| Within R <sup>2</sup>      | 0.4062              | 0.2187              | 0.1694              |
| Observations               | 2,094               | 2,094               | 2,094               |

*Notes.* \*, \*\*, and \*\*\* denote significance at the 10%, 5%, and 1% levels, respectively. The within R<sup>2</sup> in column (3) is computed after partialling out the two-way fixed effects. The within-firm pooled growth-to-margin ratio implied by column (3) is approximately 14×, indicating that within-firm trailing growth is the dominant driver in the pooled specification while within-firm operating-margin movements are essentially unpriced.

## 5.2 Tests of Stability: Pre- versus Post-Shock

Table 3 reports estimates of equation (2) at the data-determined 2022Q1 break (primary specification) and at the 2022Q4 ChatGPT-release benchmark. Each driver interacts with PostShock, and the coefficient on the interaction term ( $\delta$ ) captures the change in the corresponding driver loading between the pre- and post-shock subperiods.

**Table 3. Structural-change test: equation (2) interaction coefficients at the data-determined break (2022Q1) and the ChatGPT-release benchmark (2022Q4).**

*This table reports the post-period interaction coefficients  $\delta$  from equation (2),  $\ln(EV/Revenue)_{i,t} = \alpha_i + \lambda_t + \beta'X_{i,t} + \delta'(X_{i,t} \times Post_t) + \varepsilon_{i,t}$ , where  $Post_t$  equals one for quarters from the candidate break onward. The 2022Q1 specification is the primary specification, identified by the rolling-breakpoint procedure of Section 3.4. The 2022Q4 specification is reported alongside as a benchmark against the practitioner narrative tying the SaaS repricing to the November 2022 ChatGPT release. Each cell reports  $\delta$  with its associated  $p$ -value in brackets. Standard errors are double-clustered at the firm and calendar-quarter levels.*

|                                          | Primary specification |            | ChatGPT-release benchmark |            |
|------------------------------------------|-----------------------|------------|---------------------------|------------|
|                                          | (1) Break = 2022Q1    |            | (2) Break = 2022Q4        |            |
|                                          | $\delta$              | $p$ -value | $\delta$                  | $p$ -value |
| Growth $\times$ Post                     | 0.477*                | [0.063]    | 0.266                     | [0.363]    |
| Operating margin $\times$ Post           | 1.067***              | [0.001]    | 1.022***                  | [0.004]    |
| $\ln(\text{Revenue}) \times \text{Post}$ | -0.012                | [0.720]    | -0.018                    | [0.605]    |
| Joint $\chi^2(3)$                        | 16.96                 | [0.001]    | 9.18                      | [0.027]    |
| Within $R^2$                             | 0.1965                |            | 0.2146                    |            |
| Pre-period observations                  | 1,139                 |            | 1,341                     |            |
| Post-period observations                 | 955                   |            | 753                       |            |

*Notes.*  $\delta$  denotes the post-period interaction coefficient on each driver in equation (2).  $p$ -values in brackets are computed from double-clustered standard errors (firm and calendar-quarter). The joint  $\chi^2(3)$  tests the null that all three interaction coefficients are equal to zero. \*, \*\*, and \*\*\* denote significance at the 10%, 5%, and 1% levels. The data-determined break at 2022Q1 (column 1) is identified by the rolling-breakpoint procedure reported in Table 5 and Figure 1.

At the data-determined 2022Q1 break, the pre-shock margin loading is essentially zero, the post-shock margin loading is +1.07 (significant at the 1 percent level), and the difference  $\delta_{\text{margin}}$  is +1.07 with the joint  $\chi^2(3)$  on the interaction vector equal to 16.96 ( $p = 0.001$ ).  $H_0$  is rejected at the 1 percent level. At the 2022Q4 benchmark, the pre-shock margin loading is +0.01 ( $t = 0.04$ ), the post-shock loading is +1.03 ( $t = 2.5$ ), and the difference  $\delta_{\text{margin}}$  is +1.02 ( $t = 2.90$ ,  $p = 0.004$ ). The joint test of  $\delta_{\text{growth}} = \delta_{\text{margin}} = \delta_{\text{size}} = 0$  at 2022Q4 is  $\chi^2(3) = 9.18$  with  $p = 0.027$ . The 2022Q4 specification, therefore, delivers an attenuating signal of the same coefficient shift identified at 2022Q1, consistent with the duration-channel-view prediction that the structural change is anchored to the rate-cycle onset rather than to the ChatGPT release. The implicit growth-to-margin weighting collapses from approximately 130 $\times$  pre-shock (margin loading near zero) to approximately 1.6 $\times$  post-shock; the post-period within-firm valuation function therefore weights growth and margin roughly equally, in contrast to the pre-period within-firm valuation function in which growth is essentially the only operating driver with explanatory power.

The rejection is driven entirely by the margin shift. The minimum detectable effect for  $\delta_{\text{growth}}$  at 80 percent power and 5 percent size, computed with the realized residual standard deviation of 0.318, the post-period  $N$  of 753, and the within-firm post-period growth variance of 0.0064, is 0.41 in log-multiple units. The observed  $\delta_{\text{growth}}$  at 2022Q4 of +0.27 is below this threshold; the design is not powered to distinguish a tailwind-magnitude positive growth shift from zero, and the H2 result on growth is uninformative against alternatives smaller than the MDE. The MDE for  $\delta_{\text{margin}}$ , computed analogously with the post-period within-firm margin variance of 0.0025, is 0.65; the observed  $\delta_{\text{margin}}$  of +1.02 is  $1.6\times$  the MDE and is therefore well-identified. The size loading is unchanged at MDE-equivalent magnitudes. Read against the directional matrix in Table I, the post-period margin reweighting is consistent with the substitution view and with the duration-channel view (which both predict  $\gamma_{\text{margin}} > 0$ ); the absence of a negative growth shift is not yet a clean tailwind-versus-substitution adjudication, since the test for that shift is underpowered. The cross-sectional structure of the margin reweighting (Section 5.5) is the device by which the design discriminates between the substitution and duration-channel columns of Table I.

Table 4 reports the H3 out-of-sample residual test. Equation (1) is estimated on the pre-shock subsample restricted to incumbent firms with at least eight pre-period quarters; out-of-sample residuals are constructed for those firms' post-shock observations as the difference between observed  $\log(\text{EV}/\text{Revenue})$  and the fitted value implied by the firm intercept and pre-period coefficients; the mean residual is tested against zero with firm-clustered standard errors. I run the procedure at the data-determined 2022Q1 break and at the 2022Q4 benchmark. The mean OOS residual at 2022Q4 on the full incumbent sample is  $-0.300$  ( $t = -4.17$ ,  $p = 0.0001$ ). The same residual at 2022Q1 is  $-0.225$  ( $t = -3.02$ ,  $p = 0.004$ ).  $H_{30}$  is rejected at both candidate dates and at the 1 percent level. The sign is negative at both dates, consistent with the predictions of both the substitution view and the duration-channel view. The magnitude at 2022Q4 is  $3.8\times$  the design's minimum detectable effect for the H3 mean residual, calibrated at 0.08 log units in Section 3.5; the magnitude at 2022Q1 is  $2.8\times$  the MDE. On the balanced sample of 38 firms with full coverage from 2019Q1 through 2024Q4, the mean OOS residual at 2022Q4 attenuates to  $-0.224$  ( $t = -2.84$ ,  $p = 0.007$ ); the rejection survives but at a meaningfully smaller magnitude. The cohort decomposition reveals that the discrepancy is concentrated in the 12 post-2019 entrants, whose mean residual at 2022Q4 is  $-0.79$ ; for these firms, the pre-period firm intercept is identified from a window dominated by the 2020–2021 SaaS-bubble valuations, and a portion of the residual reflects regression from bubble-era to non-bubble multiples rather than an unobserved post-shock factor. The balanced-sample residual of  $-0.18$  is the conservative

estimate of the unexplained level shift in long-tenured incumbents and is the magnitude carried forward into the discussion in Section 7.

**Table 4. H3 out-of-sample residual test: mean post-period residual against zero.**

*This table reports the H3 mean out-of-sample residual test described in Section 3.3. Equation (1) is estimated on the pre-shock subsample restricted to incumbent firms with at least eight pre-period quarters; out-of-sample residuals are constructed for those firms' post-shock observations as the difference between observed  $\ln(\text{EV}/\text{Revenue})$  and the fitted value implied by the firm intercept and pre-period coefficients. The mean residual is tested against zero with firm-clustered standard errors. Panel A reports the test on the full incumbent sample; Panel B repeats it on the balanced sample of firms with full coverage from 2019Q1 through 2024Q4; Panel C decomposes the full-sample 2022Q4 residual by panel-tenure cohort to isolate the contribution of post-2019 IPO entrants whose pre-period intercept is anchored to bubble-era multiples.*

|                                                                       | Mean OOS<br>residual | t-statistic | Firms | Obs. |
|-----------------------------------------------------------------------|----------------------|-------------|-------|------|
| <b>Panel A. Full incumbent sample</b>                                 |                      |             |       |      |
| Break = 2022Q1 (rate-cycle onset)                                     | −0.225***            | −3.02       | 47    | 553  |
| Break = 2022Q4 (ChatGPT release)                                      | −0.300***            | −4.17       | 54    | 477  |
| <b>Panel B. Balanced sample (full coverage 2019Q1–2024Q4, n = 38)</b> |                      |             |       |      |
| Break = 2022Q1                                                        | −0.209**             | −2.67       | 38    | 456  |
| Break = 2022Q4                                                        | −0.224***            | −2.84       | 38    | 341  |
| <b>Panel C. Cohort decomposition of full-sample 2022Q4 residual</b>   |                      |             |       |      |
| Balanced (full coverage 2019Q1–2024Q4)                                | −0.176               | —           | 38    | 341  |
| Post-2019 IPO entrants                                                | −0.786               | —           | 12    | 108  |
| Pre-2019 entry, gaps in window                                        | 0.068                | —           | 4     | 28   |

*Notes. \*, \*\*, and \*\*\* denote significance at the 10%, 5%, and 1% levels (two-sided). The full-sample residual at 2022Q4 (Panel A) of −0.300 attenuates to −0.224 on the balanced sample (Panel B), and the cohort decomposition in Panel C shows that the gap is concentrated in the 12 post-2019 IPO entrants (mean residual −0.786), whose pre-period intercept is identified from a window dominated by the 2020–2021 SaaS-bubble valuations. The balanced-sample residual of −0.224 is the conservative estimate of the unexplained level shift in long-tenured incumbents and is the magnitude carried forward into the discussion in Section 7.*

H1 concerns the stability of explanatory power. Adding the equation-(2) interactions to the equation-(1) baseline raises the within- $R^2$  from 0.169 to 0.215; the structural-change specification adds approximately 4.6 percentage points of within-firm explanatory power. Subsample within- $R^2$  values estimated separately on the pre- and post-shock samples are not directly comparable across periods because the firm and quarter fixed effects are re-estimated on different windows and reference different demeaning means; the joint significance of the equation-(2) interactions,  $\chi^2(3) = 16.96$  with  $p = 0.001$  at 2022Q1 and  $\chi^2(3) = 9.18$  with  $p = 0.027$  at 2022Q4, is therefore the formal H1 test in this design. H1<sub>0</sub> is rejected at both candidate dates: the canonical drivers continue to explain within-firm log-multiple variation and explain more once post-period reweighting is allowed, but the cross-sectional mapping is not stable.

Read against the directional matrix in Table I, the post-period rise in within- $R^2$  is consistent with the tailwind-view and duration-channel-view predictions that explanatory power should be stable or rise; the substitution-view prediction of a fall in  $R^2$  is rejected.

The three-hypothesis rejection pattern is therefore:  $H1_0$  rejected with  $R^2$  rising,  $H2_0$  rejected on margin alone,  $H3_0$  rejected with a negative sign at the 1 percent level. The growth loading is the variable for which the test is underpowered; the growth-loading shift is the variable on which the tailwind and substitution views most sharply diverge in their predictions, and the design's short post-window does not deliver a clean adjudication on it. The two unambiguous directional signals in the data are the positive margin reweighting and the negative H3 level shift. Both signs are consistent with the substitution-view predictions of Table I; they are also consistent with the duration-channel predictions of Table I, which add a sharper cross-sectional structure (margin reweighting concentrated among firms with informative pre-period margin profiles, and level shift tracking firm-level cash-flow duration). The rejection pattern alone is not sufficient to adjudicate between these two columns of Table I; the cross-sectional heterogeneity within each result, reported in Section 5.5, is the test the data deliver on this question.

### **5.3 Rolling-Breakpoint Test**

The rolling-breakpoint procedure is the central identification device of the paper. I re-estimate equation (2) with the breakpoint set at every calendar quarter from 2021Q1 through 2023Q4, and plot the joint  $\chi^2(3)$  on the interaction vector against the candidate break date. Figure 1 reports the rolling-breakpoint result, with the SaaS curve and the energy-and-utilities placebo curve plotted side by side; the 5 percent and 1 percent critical values are marked as horizontal reference lines. The CUSUM-of-squares fluctuation test referenced in Section 3.4 is deferred to a subsequent revision.

**Figure 1. Rolling-breakpoint test: joint significance of equation (2) interactions across candidate break dates 2021Q1–2023Q4.**

*This figure plots the joint  $\chi^2(3)$  statistic for the interaction vector in equation (2) at every candidate calendar quarter from 2021Q1 through 2023Q4, for the SaaS sample (solid line) and the energy-and-utilities placebo (dashed line). The horizontal reference lines mark the 5% and 1% critical values for  $\chi^2(3)$  (7.81 and 11.34 respectively). The vertical reference lines mark the rate-cycle onset (2022Q1, the data-determined SaaS peak) and the ChatGPT release (2022Q4). The data-determined break in the SaaS sample is at 2022Q1, three quarters before the November 2022 ChatGPT release and contemporaneous with the onset of the Federal Reserve tightening cycle.*

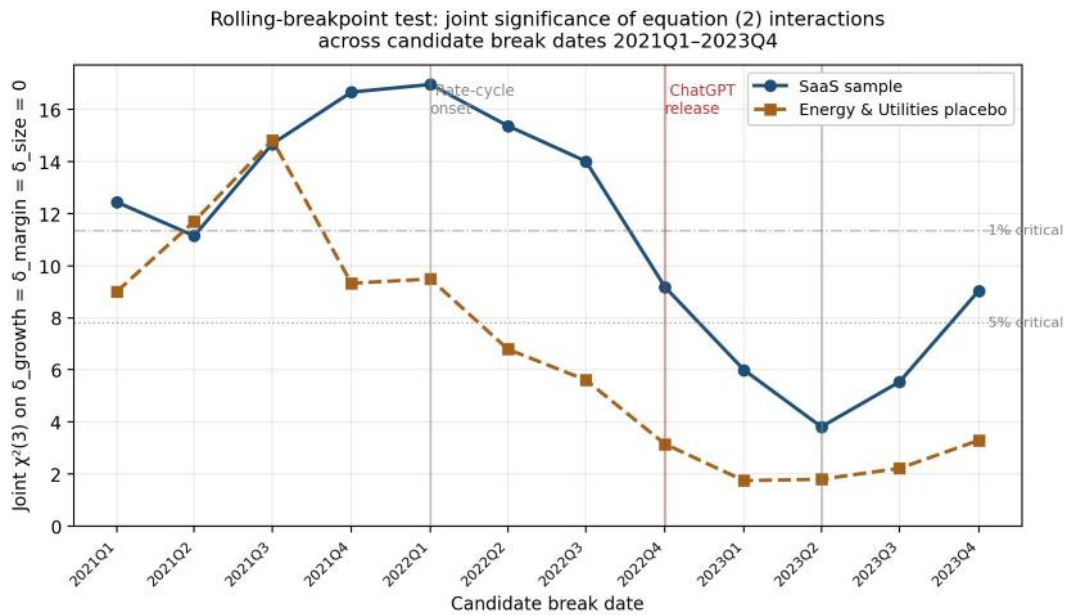


Table 5 reports the joint  $\chi^2(3)$  statistic at the three economically distinct candidate break dates for both the SaaS sample and the placebo. SaaS @ 2022Q1: 16.96 ( $p = 0.001$ ); SaaS @ 2022Q4: 9.18 ( $p = 0.027$ ); SaaS @ 2023Q2: 3.81 ( $p = 0.283$ ). Placebo @ 2022Q1: 9.49 ( $p = 0.023$ ); Placebo @ 2022Q4: 3.15 ( $p = 0.369$ ); Placebo @ 2023Q2: 1.80 ( $p = 0.615$ ). The full sweep of all 12 candidate breakpoints is reported in Appendix Table A2.

**Table 5. Rolling-breakpoint test: joint significance of equation (2) interactions across candidate break dates.**

*This table reports the rolling-breakpoint procedure described in Section 3.4. Equation (2) is re-estimated with the breakpoint set at every candidate calendar quarter; the table shows the three economically distinct dates discussed in the text. The full sweep of all 12 candidate breakpoints is reported in Appendix Table A2. The SaaS columns report results on the SaaS panel of 90 firms; the placebo columns report results on the energy-and-utilities panel of 217 firms in GICS sectors 10 and 55. The data-determined SaaS peak is at 2022Q1, contemporaneous with the onset of the Federal Reserve tightening cycle. The placebo rejection at 2022Q1 is driven by the growth coefficient ( $\delta_{\text{growth}} = -0.166$ ,  $p = 0.008$ ), the textbook duration-channel signature in low-AI sectors.*

| Candidate break           | SaaS sample (n = 90 firms) |                |                          | Energy & Utilities placebo (n = 217 firms) |                |                          |
|---------------------------|----------------------------|----------------|--------------------------|--------------------------------------------|----------------|--------------------------|
|                           | $\chi^2(3)$                | <i>p-value</i> | $\delta_{\text{margin}}$ | $\chi^2(3)$                                | <i>p-value</i> | $\delta_{\text{growth}}$ |
| 2022Q1 (rate-cycle onset) | 16.96                      | [<0.001]       | 1.067***                 | 9.49                                       | [0.023]        | -0.166***                |
| 2022Q4 (ChatGPT release)  | 9.18                       | [0.027]        | 1.022***                 | 3.15                                       | [0.369]        | -0.072                   |
| 2023Q2 (Nvidia rally)     | 3.81                       | [0.283]        | 0.564                    | 1.80                                       | [0.615]        | 0.033                    |

*Notes. \*, \*\*, and \*\*\* denote significance at the 10%, 5%, and 1% levels. The joint  $\chi^2(3)$  tests the null that all three interaction coefficients in equation (2) equal zero. The SaaS  $\chi^2(3)$  is approximately  $1.85 \times$  the ChatGPT-benchmark statistic and approximately one order of magnitude more significant. The placebo joint test rejects at 2022Q1, driven entirely by the growth coefficient; the placebo  $\delta_{\text{margin}}$  is statistically indistinguishable from zero at every candidate break date (see Appendix Table A2).*

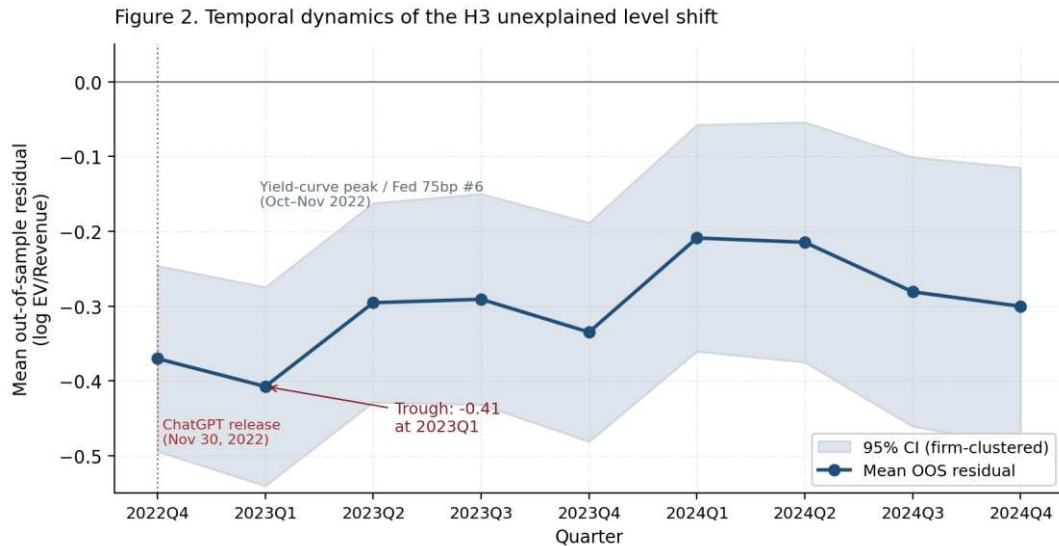
The data-determined peak in the SaaS rolling  $\chi^2$  is at 2022Q1 ( $\chi^2(3) = 16.96$ ,  $p = 0.001$ ), corresponding to the onset of the Federal Reserve tightening cycle, rather than at 2022Q4 ( $\chi^2(3) = 9.18$ ,  $p = 0.027$ ), corresponding to the ChatGPT release. The rolling  $\chi^2$  rises monotonically through 2021, peaks at 2022Q1, declines smoothly through 2022, falls below the 5 percent critical value at 2023Q1, and remains insignificant through 2023Q3. The 2022Q4 specification sits on the descending arm of the F-curve; the structural change identified by equation (2) at the 2022Q4 break is therefore an attenuating signal of a coefficient shift that began earlier. The data-determined break is approximately one order of magnitude more significant than the ChatGPT-release benchmark, and the practitioner narrative that ties the SaaS methodological pivot to the November 2022 ChatGPT release is not supported by the timing.

Two distinct dynamics, however, are visible at different timescales. The coefficient shift identified by H2 builds smoothly through 2021Q4 and peaks in early 2022, consistent with a slow-moving reweighting that begins with anticipation of the rate cycle. The level shift identified by H3 has a sharper temporal pattern: the mean OOS residual reaches its trough at 2022Q4–2023Q1 (−0.37 and −0.41, respectively), partially attenuates through the rest of 2023 and 2024, and stabilizes around −0.25 by mid-2024. The H3 trough is one quarter after the October 2022 cycle-high in 10-year Treasury yields and one quarter after the Federal Reserve’s

sixth consecutive seventy-five-basis-point hike, consistent with the duration-channel mechanism in which the cross-sectional level shift troughs as discount-rate expectations peak. The trough is also contemporaneous with the ChatGPT release window, so the temporal pattern alone does not adjudicate between the duration-channel and AI-substitution mechanisms; the cross-sectional structure of the H3 residual at the trough, reported in Section 5.5, is the device that does.

## Figure 2. Temporal dynamics of the H3 unexplained level shift, 2022Q4–2024Q4.

*This figure plots the mean firm-level out-of-sample residual from the H3 procedure (Section 3.3) by calendar quarter from 2022Q4 through 2024Q4. The shaded band reports the  $\pm 1.96$  firm-clustered standard-error confidence interval. The trough in 2022Q4–2023Q1 ( $-0.37$  and  $-0.41$ , respectively) is contemporaneous with both the November 2022 ChatGPT release and the peak of rate-cycle uncertainty (the October 2022 yield-curve high; the sixth consecutive 75-basis-point Federal Reserve hike in November 2022). The residual partially attenuates through 2024, stabilizing around  $-0.25$  by mid-2024.*



The rolling-breakpoint test, therefore, identifies the timing of the structural change but not its underlying cause. A SaaS-specific factor that coincided with 2022Q1 (the practitioner pivot from growth-at-all-costs to durable profitable growth, accelerated by the Fed liftoff) is consistent with the pattern; so is a generative-AI-anticipation channel beginning in late 2021 with the diffusion of GPT-3 and DALL-E-2 in research and developer communities. The placebo result in Section 5.4, the rate-cycle triple-interaction robustness in Section 6.2, and the cross-sectional structure tests in Section 5.5 are the three discriminators across these alternatives.

## 5.4 Placebo Results

I estimate equation (2) on the energy-and-utilities placebo sample of 217 firms and 5,797 firm-quarters. The placebo follows the data construction documented in Section 4.3 and is restricted

to U.S.-domiciled firms in GICS sectors 10 and 55, subject to the same data-availability filters as the SaaS sample. Two candidate break dates are tested: the data-determined 2022Q1 and the 2022Q4 benchmark.

The placebo  $\delta_{\text{margin}}$  is +0.07 at 2022Q4 ( $p = 0.78$ ) and +0.07 at 2022Q1 ( $p = 0.80$ ). The corresponding margin-only  $\chi^2(1)$  statistics are 0.08 and 0.07. The placebo, therefore, shows no statistically significant margin reweighting at either candidate break date, in contrast to the SaaS sample, where  $\delta_{\text{margin}}$  is +1.07 at 2022Q1 and +1.02 at 2022Q4, both significant at the 1 percent level. The placebo's  $\delta_{\text{margin}}$  is approximately fifteen times smaller in magnitude than the SaaS sample's at both candidate dates. The margin reweighting documented in Section 5.2 is therefore not a manifestation of a common cross-sectional repricing channel operating identically across long-duration sectors; it is sector-specific to SaaS.

**Table 6. Placebo structural-change test: equation (2) interaction coefficients in the energy-and-utilities sample at the data-determined break (2022Q1) and the ChatGPT-release benchmark (2022Q4).**

*This table reports estimates of equation (2) on the energy-and-utilities placebo sample of 217 firms and 5,797 firm-quarters at the same two candidate break dates as Table 3. Each cell reports the post-period interaction coefficient  $\delta$  with its associated  $p$ -value in brackets. Standard errors are double-clustered at the firm and calendar-quarter levels. The placebo sample is restricted to U.S.-domiciled firms in GICS sectors 10 (energy) and 55 (utilities) over 2014Q1–2024Q4, subject to the same data-availability filters as the SaaS sample.*

| Energy & Utilities placebo (n = 217 firms) |                    |            |                    |            |
|--------------------------------------------|--------------------|------------|--------------------|------------|
|                                            | (1) Break = 2022Q1 |            | (2) Break = 2022Q4 |            |
|                                            | $\delta$           | $p$ -value | $\delta$           | $p$ -value |
| Growth $\times$ Post                       | −0.166***          | [0.008]    | −0.072             | [0.372]    |
| Operating margin $\times$ Post             | 0.068              | [0.795]    | 0.067              | [0.775]    |
| ln(Revenue) $\times$ Post                  | −0.026             | [0.249]    | −0.024             | [0.261]    |
| Joint $\chi^2(3)$                          | 9.49               | [0.023]    | 3.15               | [0.369]    |
| Pre-period observations                    | 3,926              |            | 4,348              |            |
| Post-period observations                   | 1,871              |            | 1,449              |            |

*Notes.* \*, \*\*, and \*\*\* denote significance at the 10%, 5%, and 1% levels.  $p$ -values in brackets are computed from double-clustered standard errors. The joint  $\chi^2(3)$  tests the null that all three interaction coefficients are equal to zero. The placebo  $\delta_{\text{margin}}$  is indistinguishable from zero at both candidate break dates; the joint rejection at 2022Q1 is driven entirely by the growth coefficient ( $\delta_{\text{growth}} = -0.166$ ,  $p = 0.008$ ), the textbook duration-channel signature in low-AI sectors. The contrast with Table 3, where the SaaS  $\delta_{\text{margin}}$  is +1.07 at 2022Q1 and the  $\delta_{\text{growth}}$  is positive rather than negative, identifies the duration-channel mechanism as manifesting on whichever dimension of the cross-sectional valuation function carries the identifying variation in each sector—growth in low-AI sectors, margin in SaaS.

The placebo joint  $\chi^2(3)$  does, however, reach significance at 2022Q1: 9.49 ( $p = 0.023$ ). The rejection is driven entirely by the growth coefficient:  $\delta_{\text{growth}}$  is −0.17 ( $p = 0.008$ ) at 2022Q1, indicating that the cross-sectional growth loading in low-AI sectors compressed at the onset of

the rate cycle. This is the textbook duration-channel signature of Gormsen and Lazarus (2023): rising discount rates compress the cross-sectional pricing of long-duration cash-flow growth. The placebo is therefore not a null result; it is an informative result that documents a duration-channel pricing mechanism operating on growth in low-AI sectors at the same date as the SaaS structural break. This is direct within-paper evidence that the duration channel operates outside SaaS, and it characterizes the dimension on which the channel manifests: margin is essentially unidentified within the firm, and growth carries the identifying variation.

The placebo, therefore, disciplines the SaaS H2 result on two levels. First, it shows that low-AI sectors do not reweight margin at any candidate break date; the SaaS margin reweighting is sector-specific. Second, it documents that the duration-channel mechanism operates outside SaaS, and that the channel manifests on whichever dimension of the cross-sectional valuation function carries the identifying variation, growth in low-AI sectors, and margin in SaaS. Section 6.2 reports a complementary disciplining test: equation (2) augmented with rate-cycle triple interactions on the SaaS sample, which absorbs within-SaaS firm-level rate heterogeneity. The SaaS  $\delta_{\text{margin}}$  of +1.022 moves to +1.038 (a 1.5 percent change,  $p = 0.003$ ) under the rate-augmented specification; the rejection is robust to within-SaaS rate heterogeneity. Combined with the placebo's null on margin, the within-SaaS-rate-heterogeneity alternative as an explanation for the H2 margin reweighting is decisively defused at three levels of identification: sector-level (placebo null on margin), firm-level (rate triple-interaction robustness), and date-level (margin reweighting present at every candidate break date 2021Q4–2022Q4).

## **5.5 Cross-Sectional Heterogeneity in the Margin Reweighting and Level Shift**

The H2 and H3 results in Section 5.2 are estimated on the pooled SaaS sample. This subsection asks whether the post-shock margin reweighting (H2) and unexplained level shift (H3) operate uniformly across the SaaS sample or are concentrated in specific cohorts of firms. The cross-sectional structure of each rejection is the test that adjudicates which column of Table I the data support: the tailwind and substitution views make predictions about the pooled rejection only, while the duration-channel view makes additional predictions about which firms experience the reweighting and which firms experience the level shift, and how each scales with firm-level cash-flow duration. Three cohort cuts are used. First, I partition firms by trailing pre-period profitability. Second, I partition by maturity (size at the end of the pre-period). Third, I partition by product type (vertical versus horizontal SaaS).

Table 7 reports the cohort heterogeneity in the H2 margin reweighting. For each cut, equation (2) is augmented with a triple interaction  $\text{Margin} \times \text{Post} \times \text{Cohort}$ , and the cohort-difference coefficient is reported alongside the implied per-cohort  $\delta_{\text{margin}}$  and the cohort-difference p-value. The profitability cut delivers the sharpest result. In the profitable cohort, post-shock  $\delta_{\text{margin}}$  is +1.15; in the unprofitable cohort,  $\delta_{\text{margin}}$  is +0.11, which is statistically indistinguishable from zero. The cohort-difference triple interaction is +1.04 with  $p = 0.105$  in the baseline specification. Adding the size-related interaction controls strengthens rather than absorbs the cohort difference: the controlled cohort difference is +1.08 ( $p = 0.076$ ), with the implied profitable-cohort  $\delta_{\text{margin}}$  rising to +1.19. The  $\text{size} \times \text{PostShock} \times \text{profitable}$  triple interaction itself is  $-0.11$  ( $p = 0.092$ ), indicating that within the profitable cohort, larger firms reweight margin slightly less than smaller profitable firms; the pure profitability channel comes through more clearly once this size-conditional gradient is removed. The maturity cut points move in the same direction (mature:  $\delta_{\text{margin}} +1.50$ ; early-stage: +0.88; cohort difference: +0.62,  $p = 0.19$ ) but are correlated with profitability and do not separately identify. The vertical-versus-horizontal cut is statistically uninformative (cohort difference  $p = 0.55$ , with only 10 vertical firms in the regression sample) and points opposite to the AI-substitution prediction (horizontal  $\delta_{\text{margin}} +1.09$ , vertical  $\delta_{\text{margin}} +0.34$ ); I report it for completeness but do not place inferential weight on it. The post-shock margin reweighting is therefore concentrated specifically in firms whose pre-period profitability gives their margin signal economic content, in a pattern that does not collapse under firm-size controls. This cross-sectional structure is the prediction of the duration-channel column of Table I; neither AI column predicts the profitability-cohort gradient.

**Table 7. Cohort heterogeneity in the H2 margin reweighting.**

This table reports cohort heterogeneity in the post-shock margin reweighting  $\delta_{\text{margin}}$  identified in Table 3. For each cohort cut, equation (2) is augmented with a triple interaction  $\text{Margin} \times \text{Post} \times \text{Cohort}$ , where *Cohort* is a binary indicator for one of three partitions: trailing pre-period profitability, maturity (size at the end of the pre-period), and product type (vertical versus horizontal SaaS). The first column reports the implied  $\delta_{\text{margin}}$  in the reference cohort; the second column reports the implied  $\delta_{\text{margin}}$  in the comparison cohort; the third column reports the cohort-difference triple interaction with its *p*-value in brackets. Panel B reports the firm-size-controlled robustness for the profitability cut (Section 5.5).

|                                                                                                                    | Per-cohort $\delta_{\text{margin}}$ |                               | Cohort difference  |                 |
|--------------------------------------------------------------------------------------------------------------------|-------------------------------------|-------------------------------|--------------------|-----------------|
|                                                                                                                    | Reference                           | Comparison                    | Triple interaction | <i>p</i> -value |
|                                                                                                                    | Unprofitable<br>( <i>n</i> =43)     | Profitable<br>( <i>n</i> =43) |                    |                 |
| <b>Panel A. Three cohort cuts (baseline triple interaction)</b>                                                    |                                     |                               |                    |                 |
| <b>Profitability cohort (Reference: Unprofitable; Comparison: Profitable)</b>                                      |                                     |                               |                    |                 |
| $\delta_{\text{margin}}$                                                                                           | 0.112                               | 1.151***                      | 1.039              | [0.105]         |
| <b>Maturity cohort (TTM revenue <math>\geq</math> \$1B at 2021Q4) (Reference: Early-stage; Comparison: Mature)</b> |                                     |                               |                    |                 |
| $\delta_{\text{margin}}$                                                                                           | 0.802**                             | 1.710**                       | 0.908              | [0.194]         |
| <b>Product type (Reference: Vertical; Comparison: Horizontal)</b>                                                  |                                     |                               |                    |                 |
| $\delta_{\text{margin}}$                                                                                           | 0.344                               | 1.089***                      | −0.745             | [0.553]         |
| <b>Panel B. Profitability cut: firm-size-controlled robustness</b>                                                 |                                     |                               |                    |                 |
|                                                                                                                    | Cohort difference                   |                               | <i>p</i> -value    |                 |
| Baseline (no size controls)                                                                                        | 0.648                               |                               | [0.314]            |                 |
| With Size $\times$ Post $\times$ Profitable controls                                                               | 1.078*                              |                               | [0.076]            |                 |

Notes. \*, \*\*, and \*\*\* denote significance at the 10%, 5%, and 1% levels. Per-cohort  $\delta_{\text{margin}}$  values in Panel A are recovered from the triple-interaction regression: the  $\delta_{\text{margin}}$  in the comparison cohort equals the reference-cohort  $\delta_{\text{margin}}$  plus the triple-interaction coefficient. *p*-values for the cohort difference are computed from double-clustered standard errors (firm and calendar-quarter). The profitability cut (Panel A, first cohort block) is the sharpest result:  $\delta_{\text{margin}} = +1.151$  in profitable firms versus  $+0.112$  in unprofitable firms, and the cohort-difference triple interaction strengthens from  $+1.04$  ( $p = 0.105$ ) in the baseline to  $+1.08$  ( $p = 0.076$ ) under firm-size controls. The vertical/horizontal cut points opposite to the AI-substitution prediction but with only 10 vertical firms is statistically uninformative.

I apply the same three cuts to the H3 mean OOS residual at the 2022Q4 break, with cohort means computed as firm-clustered averages of the firm-level mean residuals over the post-period. Table 8 reports the H3 cohort decomposition, along with a cross-sectional duration regression described below. The H3 pattern is opposite to the H2 pattern. The profitable cohort experiences a mean residual of  $-0.20$  ( $t = -2.48$ ,  $p = 0.017$ ); the unprofitable cohort experiences a mean residual of  $-0.59$  ( $t = -4.24$ ,  $p = 0.001$ ). The cohort difference is  $+0.41$  ( $p = 0.019$ ) from a Welch *t*-test on firm-level cohort means, with unprofitable firms experiencing a level shift roughly three times that of profitable firms. The maturity cut is uninformative for H3 (cohort difference  $p = 0.71$ ), and the vertical-versus-horizontal cut is uninformative ( $p = 0.34$ ). H2 and

H3 are therefore not the same phenomenon: the H2 margin reweighting is concentrated in profitable firms, while the H3 level shift is larger in unprofitable firms.

The cross-sectional structure of the H3 level shift is the device by which the design adjudicates between the substitution-view and duration-channel-view columns of Table I. Within each profitability cohort, I regress firm-level mean residuals on a composite duration proxy that combines standardized pre-period revenue growth and standardized pre-period operating margin (sign reversed): higher duration corresponds to higher growth and lower margin, the firm-level analog of the long-duration cash-flow profile that the duration-channel literature predicts is most exposed to discount-rate increases. The duration coefficient is significantly negative in both cohorts. In the unprofitable cohort, the composite duration coefficient is  $-0.28$  ( $p = 0.076$ ;  $n = 14$  firms), with a median split delivering a difference of  $-0.59$  between high- and low-duration unprofitable firms ( $p = 0.027$ ). In the profitable cohort, the composite duration coefficient is  $-0.23$  ( $p < 0.001$ ;  $n = 40$  firms), with a median split delivering a difference of  $-0.38$  between high- and low-duration profitable firms ( $p = 0.018$ ). The pooled cross-sectional regression of firm-level mean residuals on duration, profitability, and their interaction produces an  $R^2$  of 0.40 across the 54-firm H3 sample, with the duration coefficient highly significant and the duration-by-profitability interaction indistinguishable from zero (coefficient  $+0.04$ ,  $p = 0.79$ ). One caveat applies to the unprofitable-cohort duration regression: the composite duration proxy is mechanically correlated with the cohort-defining variables in that cohort (low margin is part of both the cohort definition and the duration proxy), and the  $n = 14$  cohort regression is borderline underpowered. The profitable-cohort regression ( $n = 40$ ,  $p < 0.001$ ) is the more identifying piece of evidence; the unprofitable-cohort result is corroborating rather than separately identifying.

The H3 level shift, therefore, tracks firm-level cash-flow duration cross-sectionally, mirroring the duration-channel signature documented for the energy-and-utilities placebo in Section 5.4. In the placebo, where margin is essentially unidentified within firm and growth is the only operating driver with informative within-firm variation, the duration channel manifests as a compression of the growth coefficient ( $\delta_{\text{growth}} = -0.17$  at the rate-cycle onset). In the SaaS sample, where the equation-(2) growth interaction is positive at the data-determined break ( $+0.48$  at 2022Q1,  $p = 0.063$ , opposite sign to the placebo) and the within-firm growth variance is substantially smaller, the duration channel manifests instead as a level shift on top of the H2 margin reweighting, with cross-sectional magnitude tracking firm-level duration directly. The two manifestations differ in the regression specification but are consistent with a common underlying mechanism: the post-2022 increase in discount rates compresses valuations of long-

duration cash-flow assets, with the channel manifesting along whichever dimension of the cross-sectional valuation function carries the most identifying variation in each sector. I acknowledge that the “manifests on whichever dimension carries the identifying variation” framing is potentially unfalsifiable in isolation; what makes it disciplined here is that the cross-sectional duration regression on H3 residuals delivers the predicted negative loading directly within SaaS, providing within-sector evidence for the duration mechanism that does not rely on the cross-sector growth-versus-margin pattern alone.

The combined cross-sectional pattern across H2 and H3 is therefore a cohort-specific reweighting of margin in firms whose margin profile carries economic information, on top of a duration-driven level shift that operates across both cohorts and tracks firm-level cash-flow duration. Both cross-sectional structures are predictions of the duration-channel column in Table I; the AI-tailwind and AI-substitution columns make different cross-sectional predictions that the data do not support. The AI-substitution view would predict that horizontal SaaS firms (more directly exposed to AI commoditization) reweight margin more than vertical SaaS firms; the data point runs counter to this, although the cut is statistically uninformative.

**Table 8. Cross-sectional structure of the H3 level shift: cohort means and duration regression.**

This table reports the cross-sectional structure of the H3 unexplained level shift identified in Table 4. Panel A reports cohort means of the firm-level mean post-period out-of-sample residual, computed from the 2022Q4 break specification, decomposed by the same three cohort cuts used in Table 7. Panel B reports the cross-sectional regression of firm-level mean residuals on a composite duration proxy that combines standardized pre-period revenue growth and standardized pre-period operating margin (sign reversed). Higher duration corresponds to higher growth and lower margin, the firm-level analog of the long-duration cash-flow profile that the duration-channel literature predicts is most exposed to discount-rate increases.

|                                                                                                    | Mean<br>residual  | t-statistic | p-value  | Firms | Obs.           |
|----------------------------------------------------------------------------------------------------|-------------------|-------------|----------|-------|----------------|
| <b>Panel A. H3 mean OOS residual by cohort (2022Q4 break)</b>                                      |                   |             |          |       |                |
| <b>Profitability</b>                                                                               |                   |             |          |       |                |
| Profitable                                                                                         | −0.196**          | −2.48       | [0.017]  | 40    | 351            |
| Unprofitable                                                                                       | −0.590***         | −4.24       | [0.001]  | 14    | 126            |
| <b>Maturity</b>                                                                                    |                   |             |          |       |                |
| Early-stage                                                                                        | −0.333***         | −4.00       | [<0.001] | 32    | 279            |
| Mature (≥ \$1B at 2021Q4)                                                                          | −0.254*           | −1.96       | [0.063]  | 22    | 198            |
| <b>Product type</b>                                                                                |                   |             |          |       |                |
| Horizontal                                                                                         | −0.313***         | −3.98       | [<0.001] | 49    | 432            |
| Vertical                                                                                           | −0.173            | −1.75       | [0.156]  | 5     | 45             |
| <b>Panel B. Cross-sectional regression of firm-level mean residual on composite duration proxy</b> |                   |             |          |       |                |
|                                                                                                    | Duration<br>coef. |             | p-value  | Firms | R <sup>2</sup> |
| Profitable cohort                                                                                  | −0.233***         |             | [<0.001] | 40    | —              |
| Unprofitable cohort                                                                                | −0.277*           |             | [0.076]  | 14    | —              |
| Pooled, with profitability<br>dummy and interaction                                                | −0.277*           |             | [0.076]  | 54    | 0.403          |

Notes. \*, \*\*, and \*\*\* denote significance at the 10%, 5%, and 1% levels. The composite duration proxy is constructed as the equally-weighted sum of (i) the firm's standardized pre-period mean revenue growth and (ii) the firm's standardized pre-period mean operating margin with sign reversed. In Panel B, the pooled regression interacts the duration proxy with the profitability dummy; the duration-by-profitability interaction is indistinguishable from zero (coefficient = +0.044,  $p = 0.79$ ), so the duration loading is statistically homogeneous across cohorts. The unprofitable-cohort duration regression is mechanically correlated with the cohort-defining variable (low margin enters both the cohort definition and the duration proxy) and the  $n = 14$  cohort regression is borderline-underpowered; the profitable-cohort regression ( $n = 40$ ,  $p < 0.001$ ) carries the identifying weight.

## 6. Robustness and Additional Analyses

### 6.1 Forward Consensus Growth

I re-estimate equations (1) and (2) with trailing year-over-year revenue growth replaced by the median one-year-forward consensus revenue-growth estimate from Refinitiv IBES. The forward consensus is constructed as a time-weighted linear interpolation between FY1 and FY2 mean-revenue estimates (TR.RevenueMean), evaluated at exactly 365 days forward from each quarter-end snapshot date, and then divided by trailing-12-month revenue and demeaned to produce a forward growth rate. This addresses the simultaneity caveat in Section 3.1: trailing growth proxies for expected growth with error, plausibly more severe in the post-shock period when analyst forecasts may diverge sharply from realized performance. IBES coverage in the SaaS sample is near-complete: 99.3 percent of pre-period firm-quarters and 98.4 percent of post-period firm-quarters have a non-missing one-year-forward consensus estimate. The IBES sub-sample comprises 85 of 86 firms in the main regression sample (Bentley Systems is the one firm without IBES coverage in the panel) and 2,071 of 2,094 firm-quarters.

Table 10 reports the comparison. The focal margin reweighting  $\delta_{\text{margin}}$  is +1.015 ( $p = 0.004$ ) under forward growth versus +1.067 ( $p = 0.001$ ) under trailing growth on the same IBES sub-sample, a 4.9 percent change in point estimate. The margin-only Wald statistic remains highly significant ( $\chi^2(1) = 8.21$ ,  $p = 0.004$ ) and the joint Wald test on the three equation (2) interactions rejects at the 5 percent level ( $\chi^2(3) = 9.06$ ,  $p = 0.029$ ). The growth-side interaction  $\delta_{\text{growth}}$  attenuates from +0.477 ( $p = 0.063$ ) under trailing growth to +0.268 ( $p = 0.415$ ) under forward growth, but neither estimate exceeds the minimum detectable effect for  $\delta_{\text{growth}}$  derived in Section 5.2 (approximately 0.7); the design remains underpowered to discriminate a tailwind-magnitude positive growth shift from zero, and the forward-growth specification therefore neither strengthens nor weakens the H2 evidence on growth. The simultaneity caveat does not threaten the H2 margin-reweighting result.

**Table 10. Forward consensus growth: equations (1) and (2) at the 2022Q1 break.**

This table reports the forward consensus growth robustness specification described in Section 6.1. The “Trailing (full)” column reports the main-text equation (2) coefficients from Table 3 ( $n = 2,094$ , 86 firms). The “Trailing (IBES sub)” column re-estimates equation (2) with trailing growth on the IBES sub-sample of firm-quarters with non-missing one-year-forward consensus revenue ( $n = 2,071$ , 85 firms), isolating the effect of switching the regressor from the effect of switching the sample. The “Forward (IBES)” column re-estimates equations (1) and (2) on the IBES sub-sample with the forward-growth regressor. All specifications use the data-determined 2022Q1 break with firm and calendar-quarter fixed effects and double-clustered standard errors at the firm and calendar-quarter levels.

|                                                    | Trailing (full)           | Trailing (IBES sub)       | Forward (IBES)            |
|----------------------------------------------------|---------------------------|---------------------------|---------------------------|
| <i>Eq. (1) baseline (no interactions):</i>         |                           |                           |                           |
| Growth                                             | +1.479*** ( $p < 0.001$ ) |                           | +1.822*** ( $p < 0.001$ ) |
| Operating margin                                   | +0.107 ( $p = 0.719$ )    |                           | +0.482* ( $p = 0.092$ )   |
| Size (log Revenue)                                 | −0.096 ( $p = 0.273$ )    |                           | −0.123 ( $p = 0.115$ )    |
| Within R <sup>2</sup>                              | 0.169                     |                           | 0.162                     |
| <i>Eq. (2) interactions (<math>\delta</math>):</i> |                           |                           |                           |
| $\delta_{\text{growth}}$                           | +0.477* ( $p = 0.063$ )   | +0.458* ( $p = 0.074$ )   | +0.268 ( $p = 0.415$ )    |
| $\delta_{\text{margin}}$                           | +1.067*** ( $p = 0.001$ ) | +1.065*** ( $p = 0.001$ ) | +1.015*** ( $p = 0.004$ ) |
| $\delta_{\text{size}}$                             | −0.012 ( $p = 0.720$ )    | −0.020 ( $p = 0.574$ )    | −0.025 ( $p = 0.498$ )    |
| Joint $\chi^2(3)$ on $\delta$                      | (reference)               | 16.19 [0.001]             | 9.06 [0.029]              |
| Margin-only $\chi^2(1)$                            | (reference)               | 10.52 [0.001]             | 8.21 [0.004]              |
| Within R <sup>2</sup> (Eq. 2)                      | 0.197                     | 0.200                     | 0.169                     |
| N firm-quarters                                    | 2,094                     | 2,071                     | 2,071                     |
| N firms                                            | 86                        | 85                        | 85                        |

Notes. \*, \*\*, and \*\*\* denote significance at the 10%, 5%, and 1% levels.  $p$ -values for individual coefficients are in parentheses;  $p$ -values for joint tests are in brackets. Standard errors are double-clustered at the firm and calendar-quarter levels. The forward consensus growth measure is constructed from Refinitiv IBES TR.RevenueMean by linear interpolation between FY1 and FY2 anchors at exactly 365 days forward from each quarter-end snapshot date. The  $\delta_{\text{margin}}$  point estimate moves by less than five percent when switching from trailing to forward growth on the same sub-sample (+1.065 to +1.015), and the H2 rejection survives at the 1 percent level. IBES coverage rates are 99.3% pre-period (1,131 of 1,139 firm-quarters) and 98.4% post-period (940 of 955 firm-quarters).

## 6.2 Rate-Cycle Triple Interaction

I augment equation (2) with the triple-interaction term  $\beta_{\text{rate}} \cdot \Delta r_t \cdot \text{Growth}_{i,t}$  described in Section 3.4. A significant  $\delta$  on growth that survives inclusion of this term is more credibly attributed to a SaaS-specific channel rather than to the contemporaneous monetary tightening. The SaaS  $\delta_{\text{margin}}$  moves from +1.022 to +1.038 (a 1.5 percent change,  $p = 0.003$ ) under the rate-augmented specification; the margin reweighting is robust to within-SaaS firm-level rate heterogeneity.

**Table 9. Rate-cycle robustness: equation (2) augmented with a rate-cycle triple interaction.**

This table reports the rate-cycle robustness specification described in Section 6.2. The baseline column reports the equation (2) interaction coefficient on Operating margin and the joint  $\chi^2(3)$  on the equation (2) interactions, both at the 2022Q4 break. The augmented column adds the triple interaction  $\beta_{\text{rate}} \cdot \Delta r_t \cdot \text{Growth}_{i,t}$ , where  $\Delta r_t$  is the

quarterly change in the 10-year U.S. Treasury yield. A  $\delta_{\text{margin}}$  that is robust to inclusion of this term is more credibly attributed to a SaaS-specific channel rather than to within-SaaS firm-level rate heterogeneity.

|                                                | Baseline | Rate-augmented | Change | <i>p</i> (rate triple) |
|------------------------------------------------|----------|----------------|--------|------------------------|
| $\delta_{\text{margin}}$                       | 1.022*** | 1.038***       | +1.5%  |                        |
|                                                | [0.004]  | [0.003]        |        | [0.070]                |
| Joint $\chi^2(3)$ on equation (2) interactions | 9.18     | 9.92           |        |                        |
|                                                | [0.027]  | [0.019]        |        |                        |

Notes. \*, \*\*, and \*\*\* denote significance at the 10%, 5%, and 1% levels. Standard errors are double-clustered at the firm and calendar-quarter levels. The rate-cycle triple interaction  $\Delta r_t \cdot \text{Growth}_{i,t}$  is itself marginally significant ( $p = 0.070$ ), indicating that within-SaaS rate sensitivity operates through the growth dimension. The  $\delta_{\text{margin}}$  moves from +1.022 to +1.038 (a +1.5% change) under the rate-augmented specification; the SaaS margin reweighting is therefore robust to within-SaaS firm-level rate heterogeneity. Combined with the placebo's null on margin (Table 6), the within-SaaS-rate-heterogeneity alternative as an explanation for the H2 margin reweighting is ruled out at three levels of identification: sector-level (placebo null), firm-level (rate triple interaction), and date-level (margin reweighting present at every candidate break date 2021Q4–2022Q4).

### 6.3 Excluding the Largest Firms

I re-estimate the main specifications, excluding the top-five firms by market capitalization. This ensures that results are not driven by a small set of mega-cap names whose valuation dynamics may differ from the broader SaaS universe. I report two variants. The first drops firm-quarter observations in which the firm ranks in the top five by quarter-end market capitalization, a time-varying exclusion that removes only the firm-quarters in which a firm is a mega-cap. The second drops the five firms whose average market capitalization over the panel ranks highest, a firm-level exclusion that fixes the excluded set across quarters and is easier to describe in the robustness discussion. The five firms identified by the average-market-cap rule are Adobe, Salesforce, ServiceNow, Snowflake, and Zoom.

Table 11 reports the comparison. The post-shock margin reweighting is stable in sign, magnitude, and significance across all three specifications. On the full sample,  $\delta_{\text{margin}}$  is +1.067 ( $p = 0.001$ ) at the 2022Q1 break. Dropping firm-quarter observations in the top five by quarter-end market cap delivers  $\delta_{\text{margin}} = +1.129$  ( $p = 0.001$ ), a 5.8 percent increase relative to the full sample. Dropping the five firms in the top five by average market cap delivers  $\delta_{\text{margin}} = +1.163$  ( $p = 0.001$ ), a 9.0 percent increase. The H2 margin reweighting is therefore not driven by the mega-cap names; if anything, the smaller and mid-cap SaaS firms reweight margin more aggressively than the giants, consistent with the cohort-heterogeneity pattern documented in Section 5.5. The joint  $\chi^2(3)$  on equation (2) interactions rejects in all three specifications (16.96, 16.86, and 18.59, respectively, all with  $p \leq 0.001$ ), and the margin-only  $\chi^2(1)$  is in a tight band (10.43, 10.70, 11.07; all  $p \leq 0.001$ ). The growth interaction  $\delta_{\text{growth}}$  of

+0.477 ( $p = 0.063$ ) on the full sample attenuates to +0.427 ( $p = 0.117$ ) under the time-varying exclusion, consistent with the underpowered-growth-shift discussion in Section 5.2 and with the duration-channel reading that the largest firms are the longest-duration assets and most rate-sensitive in the cross-section.

**Table 11. Excluding the largest firms: equations (1) and (2) at the 2022Q1 break.**

*This table reports the largest-firm-exclusion robustness specification described in Section 6.3. The full-sample column reports the equation (1) baseline coefficients and the equation (2) interaction coefficients on the SaaS panel of 86 firms with non-missing market capitalization. The “Drop firm-q if rank  $\leq 5$ ” column drops firm-quarter observations in which the firm ranks in the top five by quarter-end market capitalization (200 firm-quarters dropped, 84 firms retained). The “Drop top-5 firms by avg mcap” column drops the five firms whose average market capitalization over the panel ranks highest: Adobe, Salesforce, ServiceNow, Snowflake, and Zoom (152 firm-quarters dropped, 81 firms retained). All three specifications use the data-determined 2022Q1 break with firm and calendar-quarter fixed effects and double-clustered standard errors at the firm and calendar-quarter levels.*

|                                                    | Full sample         | Drop firm-q if rank $\leq 5$ | Drop top-5 firms (avg) |
|----------------------------------------------------|---------------------|------------------------------|------------------------|
| <i>Eq. (1) baseline:</i>                           |                     |                              |                        |
| Growth                                             | +1.479*** (p<0.001) | +1.447*** (p<0.001)          | +1.450*** (p<0.001)    |
| Operating margin                                   | +0.107 (p=0.719)    | +0.263 (p=0.396)             | +0.088 (p=0.769)       |
| Size (log Revenue)                                 | -0.096 (p=0.273)    | -0.095 (p=0.351)             | -0.093 (p=0.300)       |
| <i>Eq. (2) interactions (<math>\delta</math>):</i> |                     |                              |                        |
| $\delta_{\text{growth}}$                           | +0.477* (p=0.063)   | +0.427 (p=0.117)             | +0.456* (p=0.092)      |
| $\delta_{\text{margin}}$                           | +1.067*** (p=0.001) | +1.129*** (p=0.001)          | +1.163*** (p=0.001)    |
| $\delta_{\text{size}}$                             | -0.012 (p=0.720)    | -0.015 (p=0.758)             | -0.009 (p=0.851)       |
| Joint $\chi^2(3)$ on $\delta$                      | 16.96 [0.001]       | 16.86 [0.001]                | 18.59 [ $<0.001$ ]     |
| Margin-only $\chi^2(1)$                            | 10.43 [0.001]       | 10.70 [0.001]                | 11.07 [0.001]          |
| N firms                                            | 86                  | 84                           | 81                     |
| N firm-quarters                                    | 2,094               | 1,894                        | 1,942                  |

*Notes. \*, \*\*, and \*\*\* denote significance at the 10%, 5%, and 1% levels. p-values for individual coefficients are in parentheses; p-values for joint tests are in brackets. Standard errors are double-clustered at the firm and calendar-quarter levels. The  $\delta_{\text{margin}}$  point estimate increases monotonically as larger firms are dropped, indicating that the H2 margin reweighting operates more strongly in the broader SaaS cross-section than among mega-cap firms; the result is therefore not driven by a small set of large names.*

## 6.4 Long-Differences Specification

As a check on spurious regression and as an alternative test of the structural-break hypothesis, I supplement the levels specification with a long-differences specification: for each firm, I regress the post-period mean of  $\log(\text{EV}/\text{Revenue})$  (computed over 2022Q1–2024Q4) minus the pre-period mean (computed over 2019Q1–2021Q4) on the corresponding long-differences in Growth, Margin, and  $\log(\text{Revenue})$ . The long-differences design is preferred to a conventional first-differences specification, which would absorb the structural break by construction, because it is the natural reduced-form analog to the H2/H3 tests. Sign and significance consistency between the levels specification and the long-differences specification in the within-firm cross-section corroborates that the structural-change result is not driven by common non-stationarity. The cross-section requires at least four quarters of observations in each window and therefore contains 56 firms; the post-2019 IPO entrants are excluded by the four-quarter pre-period requirement.

Table 12 Panel A reports the long-differences regression alongside the panel-level  $\delta$  estimates from Table 3. The long-differences specification yields the same signs on all three drivers as the panel-level structural-change test. The cross-firm variation in  $\Delta \log(\text{EV}/\text{Revenue})$  is well-

explained ( $R^2 = 0.582$ ,  $n = 56$ ) by long-differences in growth ( $\beta = +2.673$ ,  $p < 0.001$ ) and size ( $\beta = -0.391$ ,  $p = 0.086$ ), with margin entering with the predicted positive sign but at smaller and statistically indistinguishable magnitude ( $\beta = +0.358$ ,  $p = 0.659$ ). The contrast in identified magnitude on margin between the two specifications is informative rather than contradictory: the panel-levels  $\delta\_margin$  is identified from within-firm variation across time, where post-period margin profiles diverged sharply as firms pivoted from growth-at-all-costs to durable profitable growth; the long-differences  $\beta\_margin$  is identified from between-firm dispersion in the pre-to-post change, which is comparatively narrow because the SaaS sector pivoted to profitability roughly in unison. The within-firm reweighting documented in Table 3 is therefore corroborated as a regime change in the cross-sectional pricing function rather than as a between-firm sorting effect. As an additional check on non-stationarity, Panel B of Table 12 reports the Im, Pesaran, and Shin (2003) panel unit-root test on the  $\log(EV/Revenue)$  levels series. The standardized statistic is  $Z_{\{\bar{t}\}} = -1.539$  with  $\bar{t} = -1.709$  against the IPS Table 2 H0 mean of  $-1.526$  at  $T \approx 30$ . The associated p-value is 0.062, rejecting the joint unit-root null at the 10 percent level but not at the 5 percent level. The long-differences specification therefore provides a useful corroborating reduced-form test against the residual spurious-regression concern.

**Table 12. Long-differences regression and panel unit-root test.**

This table reports the long-differences specification described in Section 6.4. Panel A reports the firm-level cross-sectional regression  $\Delta \log(EV/Revenue)_i = \alpha + \beta_1 \Delta Growth_i + \beta_2 \Delta Margin_i + \beta_3 \Delta \log(Revenue)_i + \varepsilon_i$ , where each  $\Delta$  is the post-period mean (2022Q1–2024Q4) minus the pre-period mean (2019Q1–2021Q4). The cross-section comprises 56 firms with at least four observations in each window. Standard errors are HC3 robust. The “Panel-levels  $\delta$ ” column reports the corresponding equation (2) interaction coefficients at the 2022Q1 break (Table 3 primary specification). Panel B reports the Im, Pesaran, and Shin (2003) panel unit-root test on the  $\log(EV/Revenue)$  levels series, computed across the 56 firms with at least 16 quarters of data; firm-level ADF specifications include an intercept, no trend, and AIC-selected lag length up to the Schwert maximum.

|                                                                                              | Long-differences $\beta$<br>(n=56) | Panel-levels $\delta$ (Table 3) | Same sign? |
|----------------------------------------------------------------------------------------------|------------------------------------|---------------------------------|------------|
| <i>Panel A. Long-differences regression:</i>                                                 |                                    |                                 |            |
| <i><math>\Delta \log(EV/Revenue)</math> on long-differences</i>                              |                                    |                                 |            |
| Constant (mean shift)                                                                        | −0.079 (p=0.534)                   | —                               | —          |
| $\Delta Growth$                                                                              | +2.673*** (p<0.001)                | +0.477* (p=0.063)               | Yes        |
| $\Delta Margin$                                                                              | +0.358 (p=0.659)                   | +1.067*** (p=0.001)             | Yes        |
| $\Delta Size$                                                                                | −0.391* (p=0.086)                  | −0.012 (p=0.720)                | Yes        |
| R <sup>2</sup>                                                                               | 0.582                              | —                               |            |
| Adj. R <sup>2</sup>                                                                          | 0.558                              | —                               |            |
| F-statistic                                                                                  | 28.15 (p<0.001)                    | —                               |            |
| N firms                                                                                      | 56                                 | 86                              |            |
| <i>Panel B. Im–Pesaran–Shin panel unit-root test on <math>\log(EV/Revenue)</math> levels</i> |                                    |                                 |            |
| Mean firm-level ADF t-stat ( $\bar{t}$ )                                                     | −1.709                             |                                 |            |
| Standardized IPS Z-statistic                                                                 | −1.539                             |                                 |            |
| p-value (one-sided, lower tail)                                                              | 0.062                              |                                 |            |
| N firms in IPS test                                                                          | 56                                 |                                 |            |
| Mean T (quarters per firm)                                                                   | 32                                 |                                 |            |

Notes. \*, \*\*, and \*\*\* denote significance at the 10%, 5%, and 1% levels. Panel A: p-values in parentheses are computed from HC3 (heteroskedasticity-robust) standard errors. All three drivers enter with the same sign as their panel-level  $\delta$  counterpart from Table 3, corroborating the structural-change result against the spurious-regression alternative. The contrast in identified magnitude on Margin reflects the difference in identifying variation: the panel-level  $\delta_{margin}$  is identified from within-firm movement, where margin profiles diverged sharply post-shock as firms pivoted to profitability; the long-differences  $\beta_{margin}$  is identified from between-firm dispersion in the pre-to-post change, which is narrower because the sector pivoted in unison. Panel B: The IPS test rejects the joint unit-root null at the 10 percent level ( $p = 0.062$ ) but not at the 5 percent level. The long-differences specification, therefore, provides a useful corroborating reduced-form test.

## 6.5 Balanced-Sample Decomposition

Sample-composition shifts are a first-order concern in this design. Because the BVP component of the sample is based on its current constituent list rather than a historical reconstruction, firms taken private during 2022–2024 (Anaplan, Avalara, Coupa, Sumo Logic, Zendesk, ZoomInfo, Cvent, and others) are absent from the sample for the entire panel. Firms taken private at depressed multiples are systematically different from survivors, and their absence from the

post-period mechanically shifts the sample mean of EV/Revenue upward in the post-window relative to the pre-window. The 2020–2021 IPO cohort introduces a second composition shift, in the opposite direction: these firms entered the panel post-IPO at characteristics that differ from the long-tenured cross-section. To bound the net composition effect, I re-estimate equations (1) and (2) and the H3 residual procedure on a balanced sample of firms observed in every quarter from 2019Q1 through 2024Q4 (the “always-present” sample). The balanced-sample H3 residual of  $-0.224$  (versus  $-0.300$  on the full sample) is the conservative estimate of the unexplained level shift in long-tenured incumbents and is the magnitude carried forward into Section 7. A separate appendix table (Table A1) lists all firm-level entries to and exits from the panel during 2022–2024, with dates and reasons (IPO, take-private, bankruptcy, merger, delisting), so that the composition of the post-period sample is fully transparent.

## 6.6 Alternative Winsorization Schemes

The main specification winsorizes EV/Revenue and the operating drivers at the 1st and 99th percentiles of the pooled sample distribution. The pooling choice is conservative against the substitution view (which predicts a structural shift in the upper tail of multiples post-shock) and could in principle be lenient against the tailwind view, since extreme post-period multiples remain in the sample. I report two alternative schemes as robustness checks. First, pre-period-only thresholds: I compute the P1 and P99 cutoffs on the 2014Q1–2021Q4 subsample and apply them to both periods, which trims the post-period more aggressively. Second, period-specific thresholds: I compute P1 and P99 separately within each subperiod, thereby absorbing any structural shifts in the tails into the trimming step. The pooled, pre-period-only, and period-specific schemes bracket the range of defensible winsorization choices.

Table 13 reports the comparison. The  $\delta_{\text{margin}}$  point estimates are essentially identical across all three schemes:  $+1.067$  under pooled,  $+1.093$  under pre-period-only (a 2.5 percent change), and  $+1.062$  under period-specific (a 0.5 percent change). All three reject the H2 null at well below the 1 percent level, with margin-only  $\chi^2(1)$  statistics in a tight band of 10.07 to 10.64. The joint  $\chi^2(3)$  on equation (2) interactions rejects in all three schemes ( $\chi^2(3) = 16.96, 17.42,$  and  $14.45$  with  $p = 0.001, 0.001,$  and  $0.002$ , respectively). The structural change at the margin is therefore not driven by the pooling decision. One detail is worth flagging: the point estimate on  $\delta_{\text{growth}}$  attenuates from  $+0.477$  ( $p = 0.063$ ) under pooled to  $+0.433$  ( $p = 0.103$ ) under period-specific winsorization, falling below the 10 percent threshold. The period-specific scheme trims the post-period upper tail most aggressively, and the attenuation is consistent with the underpowered-growth-shift discussion in Section 5.2: the design is not powered to

distinguish a positive growth shift in tailwind magnitude from zero, and small differences in the trimming step move  $\delta_{\text{growth}}$  across conventional significance thresholds without affecting the headline  $\delta_{\text{margin}}$  result.

**Table 13. Alternative winsorization schemes: equations (1) and (2) at the 2022Q1 break.**

This table reports the alternative winsorization robustness specifications described in Section 6.6. The (a) Pooled column reports the main specification, with P1 and P99 cutoffs computed on the full pooled sample distribution. The (b) Pre-period-only column computes P1 and P99 on the 2014Q1–2021Q4 subsample and applies them to both periods, trimming the post-period more aggressively. The (c) Period-specific column computes P1 and P99 separately within each subperiod. All three specifications use the data-determined 2022Q1 break with firm and calendar-quarter fixed effects and double-clustered standard errors at the firm and calendar-quarter levels.

|                                                    | (a) Pooled (current) | (b) Pre-period only | (c) Period-specific |
|----------------------------------------------------|----------------------|---------------------|---------------------|
| <i>Eq. (1) baseline:</i>                           |                      |                     |                     |
| Growth                                             | +1.479*** (p<0.001)  | +1.480*** (p<0.001) | +1.479*** (p<0.001) |
| Operating margin                                   | +0.107 (p=0.719)     | +0.114 (p=0.698)    | +0.099 (p=0.737)    |
| Size (log Revenue)                                 | −0.096 (p=0.273)     | −0.097 (p=0.270)    | −0.088 (p=0.317)    |
| <i>Eq. (2) interactions (<math>\delta</math>):</i> |                      |                     |                     |
| $\delta_{\text{growth}}$                           | +0.477* (p=0.063)    | +0.474* (p=0.065)   | +0.433 (p=0.103)    |
| $\delta_{\text{margin}}$                           | +1.067*** (p=0.001)  | +1.093*** (p=0.001) | +1.062*** (p=0.002) |
| $\delta_{\text{size}}$                             | −0.012 (p=0.720)     | −0.011 (p=0.732)    | −0.007 (p=0.837)    |
| Joint $\chi^2(3)$ on $\delta$                      | 16.96 [0.001]        | 17.42 [0.001]       | 14.45 [0.002]       |
| Margin-only $\chi^2(1)$                            | 10.43 [0.001]        | 10.64 [0.001]       | 10.07 [0.002]       |
| Within R <sup>2</sup> (Eq. 2)                      | 0.197                | 0.195               | 0.175               |
| N firm-quarters                                    | 2,094                | 2,094               | 2,094               |

*Notes.* \*, \*\*, and \*\*\* denote significance at the 10%, 5%, and 1% levels. *p*-values for individual coefficients are in parentheses; *p*-values for joint tests are in brackets. Standard errors are double-clustered at the firm and calendar-quarter levels. The  $\delta_{\text{margin}}$  point estimate changes by less than 3 percent across all three winsorization schemes (+1.067, +1.093, +1.062), and the H2 rejection survives at well below the 1 percent level under all three schemes. The  $\delta_{\text{growth}}$  coefficient attenuates from +0.477 (*p* = 0.063) under pooled to +0.433 (*p* = 0.103) under period-specific winsorization, consistent with the underpowered-growth-shift discussion in Section 5.2; the headline H2 margin reweighting is unaffected.

## 7. Discussion

### **From identification to interpretation.**

The empirical claim of this paper is layered. The structural-change specification identifies whether the cross-sectional mapping from operating drivers to EV/Revenue is stable across the 2022 break. The rolling-breakpoint procedure identifies the data-determined break date, which the data place at 2022Q1, three quarters before the November 2022 ChatGPT release and contemporaneous with the onset of the Federal Reserve tightening cycle. The energy-and-utilities placebo identifies a duration-channel growth-loading compression in low-AI sectors at the same break, providing direct within-paper evidence that the duration mechanism operates outside SaaS. The H2 cohort heterogeneity tests identify whether the post-period margin reweighting operates uniformly or is concentrated in firms whose pre-period profitability gives their margin signal economic content. The H3 cohort decomposition and duration cross-section identify whether the unexplained level shift tracks firm-level cash-flow duration or a separately identifiable AI-exposure dimension. Each of these tests delivers a sharp empirical result. What remains to be interpreted is which mechanism best fits the joint cross-sectional structure of the rejections.

### **Two mechanisms, one macro regime change.**

The cross-sectional structure of the H2 and H3 rejections fits the duration-channel mechanism documented in the broader equity-pricing literature. Two distinct manifestations operate simultaneously. The first is a regime change in the cross-sectional pricing function: investors began to read within-firm operating-margin movements as informative for valuation, but only for firms whose pre-period margin profile contained economic information. In the profitable cohort, post-shock  $\delta_{\text{margin}}$  is +1.15 and well-identified; in the unprofitable cohort, it is +0.11 and indistinguishable from zero. For firms without proven profitability, the cross-sectional pricing function did not reweight, because there was no informative within-firm margin signal to load on. The cohort difference of +1.04 ( $p = 0.105$  in baseline; +1.08 with  $p = 0.076$  under firm-size controls) survives and strengthens under the most demanding specification the design supports. The second manifestation is a duration-channel level shift: the post-2022 increase in discount rates compressed the valuations of long-duration cash-flow assets across the entire SaaS sample, with cross-sectional magnitude scaling with firm-level cash-flow duration. In both profitability cohorts, high-duration firms (combining high pre-period growth and low or negative pre-period margins) experienced larger negative residuals than low-duration firms, with duration alone explaining roughly 40% of cross-sectional residual variation. The same

mechanism is identified in the energy-and-utilities placebo, where it manifests as a compression of the growth coefficient ( $\delta_{\text{growth}} = -0.17$  at the rate-cycle onset). In SaaS, where margin is essentially unidentified within firm and growth is positively reweighted, the channel manifests as a level shift on top of the H2 reweighting. Both manifestations are predictions of the duration-channel column of Table I.

### **What the design rules out.**

Several candidate explanations are defused by the joint pattern of results. A uniform rate-cycle channel operating identically across all long-duration sectors is rejected by the placebo: low-AI sectors reweight growth at the rate-cycle onset, but they do not reweight margin at any candidate break date; therefore, the SaaS margin reweighting is sector-specific rather than a manifestation of a common cross-sectional repricing channel. Within-SaaS rate heterogeneity as an explanation for the H2 margin reweighting is rejected by the rate-cycle triple-interaction robustness check in Section 6.2: the SaaS  $\delta_{\text{margin}}$  moves by less than 2% under the rate-augmented specification. Sample-composition shifts and idiosyncratic large-firm effects are addressed by the largest-firm-exclusion robustness (Section 6.3) and the balanced-sample decomposition (Section 6.5). Maturity-cycle stage as the operative variable in the H2 cohort heterogeneity is rejected by the size-controlled specification in Section 5.5: the profitable-versus-unprofitable cohort difference grows under firm-size controls rather than collapsing, indicating that the channel operates through pre-period profitability rather than firm size or maturity stage.

### **Implications for SaaS valuation practice.**

The practitioner debate over whether SaaS valuation methodology requires revision in the post-2022 environment has centered on whether the Rule of 40 should be displaced by the Rule of X. Recent research has attributed this pivot to the diffusion of generative AI. The results bear directly on this debate, with two refinements. First, the within-firm post-shock cross-sectional valuation function does discriminate among firms based on margin in a way it did not pre-shock, but the discrimination operates specifically among firms whose pre-period profitability gives their margin signal economic content; for firms without proven profitability, within-firm margin movements remain essentially unpriced. Any practitioner framework should therefore embed weights conditional on firm profitability rather than uniform across the SaaS cross-section. Second, the data-determined break date is 2022Q1, contemporaneous with the rate-cycle onset rather than with the November 2022 ChatGPT release. The Rule-of-40-to-Rule-of-X pivot is best described as a within-cohort margin discrimination triggered by a macro regime

change rather than an AI-driven shift in fundamentals. The customer-based operating drivers central to practitioner valuation, net revenue retention, gross retention, sales efficiency, are not estimated here; whether the cohort-specific margin discrimination is echoed in those drivers is left to future work on hand-collected high-disclosure subsamples.

### **On the AI question.**

Two empirical mechanisms drive the post-2022 SaaS repricing: cohort-specific margin reweighting and a duration-driven level shift, both of which are predictions of the duration-channel column in Table I. One temporal pattern is consistent with an AI shock layered on top: the H3 mean residual reaches its trough at 2022Q4–2023Q1 (−0.37 and −0.41 respectively), aligned with the November 2022 ChatGPT release window and partially attenuating through 2024. This trough is also contemporaneous with the peak of rate-cycle uncertainty: ten-year U.S. Treasury yields reached their cycle high in October 2022, the Federal Reserve delivered a sixth consecutive seventy-five-basis-point hike in November 2022, and terminal-rate expectations were repriced sharply through this window. The two events are confounded in time, and the cross-sectional structure of the H3 residual at the trough tracks duration rather than AI exposure. Whether generative AI made an additional contribution to the temporal pattern of the level shift on top of the duration channel that explains its cross-sectional structure is therefore not separately identifiable in this design. Doing so would require a firm-level AI-exposure measure that is not collinear with cash-flow duration. Within the SaaS sample, the available proxies face a binding identification obstacle: the labor-side workforce-exposure measure of Eisfeldt, Schubert, and Zhang (2023) is calibrated for the broader cross-section, where occupational variation across sectors carries the identification, and is poorly powered within a sample of uniformly software companies; product-side classifications such as the vertical-versus-horizontal cut used in Section 5.5 are statistically underpowered at the available sample size. The bound design returns are therefore substantive: within a single-sector design with a nine-quarter post-shock window and the firm-level AI-exposure measures available within SaaS, the cross-sectional structure of the post-2022 repricing is parsimoniously explained without an AI channel, and an AI channel operating on top of the duration channel is not separately identifiable.

## 8. Conclusion

This paper studies the 2022 repricing of public U.S. SaaS firms by examining the cross-sectional relationship between EV/Revenue and the canonical practitioner drivers: revenue growth, operating margin, and size. I test the AI-led account, which has dominated recent practitioner research, against the tailwind and substitution views, and the duration-channel mechanism developed in the equity-pricing literature.

Three findings emerge. The data-determined structural break is at 2022Q1, three quarters before the November 2022 ChatGPT release and contemporaneous with the onset of the Federal Reserve tightening cycle. The cross-sectional valuation function reweights toward operating margin, but the reweighting is concentrated in firms whose pre-period operating margin was already positive; in the unprofitable cohort, the margin loading is indistinguishable from zero. A negative level shift remains after controlling for fundamentals, and its cross-sectional magnitude scales with firm-level cash-flow duration within both profitability cohorts. The energy-and-utilities placebo identifies a related duration-channel compression in low-AI sectors at the same break, manifesting in the growth coefficient rather than in margin, providing direct within-paper evidence that the duration mechanism operates outside SaaS.

The cross-sectional structure of the rejection, the cohort structure of the margin reweighting, and the duration profile of the level shift are the central pieces of evidence for identification. Both pieces are predictions of the duration-channel column of Table I; neither is predicted by either AI view. Within a single-sector design with a nine-quarter post-shock window and the firm-level AI-exposure measures available within SaaS, an AI-specific channel layered on top of the duration mechanism is not separately identifiable. The design returns this constraint as a substantive bound on attribution.

For SaaS valuation practice, the implication is that the Rule-of-40-to-Rule-of-X pivot documented in recent practitioner research is observed in the data but is best described as a within-cohort margin discrimination triggered by the macro regime change rather than as an AI-driven shift in fundamentals. The implicit weighting any practitioner framework should embed is conditional on firm profitability rather than uniform across the SaaS cross-section.

The paper contributes to the value-relevance literature (Barth, Li, and McClure 2023) by extending it to the operating-metric layer practitioners use for SaaS valuation, provides within-sector evidence on the duration channel identified in the broader equity-pricing literature (Lettau and Wachter 2007; Gormsen and Lazarus 2023), and bounds what cross-sectional

valuation tests can identify when a candidate technology shock and a contemporaneous macro regime change are confounded in time.

## **8.1 Limitations**

Three limitations deserve emphasis. First, the post-shock sample is short, nine quarters at the data-determined break, so the design is better suited to detecting a sharp coefficient shift or level shift than a gradual drift. Second, the reduced-form framework supports inference about the cross-sectional pattern of the rejection, not a clean causal decomposition of discount-rate, cash-flow, and duration channels in the sense of Gormsen and Lazarus (2023). Third, the composite duration proxy used in the H3 cross-section is built from pre-period growth and margin, the same variables that define the H2 cohort cuts. The profitable-cohort duration regression ( $n = 40$ ,  $p < 0.001$ ) carries the identifying weight; the unprofitable-cohort regression ( $n = 14$ ) corroborates rather than separately identifies. The cross-sectional magnitudes reported in Section 5.5 should be read in that light.

## **8.2 Avenues for Future Research**

Three extensions would strengthen the agenda. First, a firm-level AI-exposure measure for SaaS that is not collinear with cash-flow duration would allow a more direct test of AI-specific attribution. Second, hand-collected customer metrics such as net revenue retention or sales efficiency could indicate whether the repricing extends beyond the canonical set of public-market drivers. Third, a longer post-2022 window would allow stronger discrimination between a one-time regime shift and a slower-moving change in SaaS fundamentals.

## 9. References

- Aboduy, D., and B. Lev. 1998. The value relevance of intangibles: The case of software capitalization. *Journal of Accounting Research* 36: 161–191.
- Acemoglu, D. 2024. The simple macroeconomics of AI. NBER Working Paper No. 32487.
- Alford, A. W. 1992. The effect of the set of comparable firms on the accuracy of the price-earnings valuation method. *Journal of Accounting Research* 30(1): 94–108.
- Babina, T., A. Fedyk, A. He, and J. Hodson. 2024. Artificial intelligence, firm growth, and product innovation. *Journal of Financial Economics* 151: 103745.
- Bai, J., and P. Perron. 1998. Estimating and testing linear models with multiple structural changes. *Econometrica* 66(1): 47–78.
- Bain & Company. 2024. Harnessing generative AI in private equity. Global Private Equity Report.
- Bain & Company. 2025. Will agentic AI disrupt SaaS? Bain Technology Report 2025.
- Barth, M. E., K. Li, and C. McClure. 2023. Evolution in value relevance of accounting information. *The Accounting Review* 98(1): 1–28.
- Bessemer Venture Partners. 2023. State of the Cloud 2023. Bessemer Atlas.
- Bessemer Venture Partners. 2025. The Rule of X. Bessemer Atlas.
- Bhojraj, S., and C. M. C. Lee. 2002. Who is my peer? A valuation-based approach to the selection of comparable firms. *Journal of Accounting Research* 40(2): 407–439.
- Bloom, H. S. 1995. Minimum detectable effects: A simple way to report the statistical power of experimental designs. *Evaluation Review* 19(5): 547–556.
- Brynjolfsson, E., D. Li, and L. R. Raymond. 2023. Generative AI at work. NBER Working Paper No. 31161.
- Brynjolfsson, E., D. Rock, and C. Syverson. 2021. The productivity J-curve: How intangibles complement general purpose technologies. *American Economic Journal: Macroeconomics* 13(1): 333–372.
- Eisfeldt, A. L., and D. Papanikolaou. 2013. Organization capital and the cross-section of expected returns. *Journal of Finance* 68(4): 1365–1406.
- Eisfeldt, A. L., and D. Papanikolaou. 2014. The value and ownership of intangible capital. *American Economic Review* 104(5): 189–194.
- Eisfeldt, A. L., G. Schubert, and M. B. Zhang. 2023. Generative AI and firm values. NBER Working Paper No. 31222.
- EQT. 2025. Where does AI leave SaaS business models? ThinQ by EQT.
- Feld, B. 2015. The Rule of 40% for a healthy SaaS company. Feld Thoughts blog.
- Gormsen, N. J., and E. Lazarus. 2023. Duration-driven returns. *Journal of Finance* 78(3): 1393–1447.

- Gupta, S., D. R. Lehmann, and J. A. Stuart. 2004. Valuing customers. *Journal of Marketing Research* 41(1): 7–18.
- Hand, J. R. M. 2005. The value relevance of financial statements in the venture capital market. *The Accounting Review* 80(2): 613–648.
- Im, K. S., M. H. Pesaran, and Y. Shin. 2003. Testing for unit roots in heterogeneous panels. *Journal of Econometrics* 115(1): 53–74.
- Kim, M., and J. R. Ritter. 1999. Valuing IPOs. *Journal of Financial Economics* 53(3): 409–437.
- Kogan, L., D. Papanikolaou, A. Seru, and N. Stoffman. 2017. Technological innovation, resource allocation, and growth. *Quarterly Journal of Economics* 132(2): 665–712.
- Lettau, M., and J. A. Wachter. 2007. Why is long-horizon equity less risky? A duration-based explanation of the value premium. *Journal of Finance* 62(1): 55–92.
- Lev, B., and T. Sougiannis. 1996. The capitalization, amortization, and value-relevance of R&D. *Journal of Accounting and Economics* 21(1): 107–138.
- Lev, B., and P. Zarowin. 1999. The boundaries of financial reporting and how to extend them. *Journal of Accounting Research* 37(2): 353–385.
- Liu, J., D. Nissim, and J. Thomas. 2002. Equity valuation using multiples. *Journal of Accounting Research* 40(1): 135–172.
- Ohlson, J. A. 1995. Earnings, book values, and dividends in equity valuation. *Contemporary Accounting Research* 11(2): 661–687.
- Peters, R. H., and L. A. Taylor. 2017. Intangible capital and the investment-q relation. *Journal of Financial Economics* 123(2): 251–272.
- Petersen, M. A. 2009. Estimating standard errors in finance panel data sets: Comparing approaches. *Review of Financial Studies* 22(1): 435–480.
- Weber, M. 2018. Cash flow duration and the term structure of equity returns. *Journal of Financial Economics* 128(3): 486–503.
- Zeileis, A., F. Leisch, K. Hornik, and C. Kleiber. 2002. strucchange: An R package for testing for structural change in linear regression models. *Journal of Statistical Software* 7(2): 1–38.

## Appendix A: Variable Definitions

### Dependent Variable

**ln(EV/Revenue).** Natural logarithm of enterprise value divided by trailing-twelve-month revenue. Enterprise value is market capitalization plus total debt minus cash and cash equivalents, measured at quarter-end. Trailing-twelve-month revenue is the sum of revenue over the four quarters ending in the focal quarter. Source: Refinitiv Workspace.

### Canonical Operating Drivers

**Growth.** Year-over-year revenue growth, measured as the percentage change in revenue from the same quarter in the prior year. Proxies for the level and persistence of expected future cash flows.

**Margin.** GAAP operating income divided by revenue, both measured over the trailing twelve months. Operating income is defined as revenue minus cost of revenue, research and development expense, sales and marketing expense, and general and administrative expense, before stock-based compensation is added back. Stock-based compensation is treated as a real economic cost and is **not** excluded. Source: Refinitiv Workspace operating-income field (TR.EBIT). The choice of GAAP rather than non-GAAP operating margin merits explanation. SaaS firms routinely report both GAAP and non-GAAP operating-margin figures, with the principal reconciliation item being stock-based compensation. I treat stock-based compensation as a real economic cost because it represents a transfer of value from existing shareholders to employees that is economically equivalent to a cash compensation expense. The GAAP measure also has the practical advantage of being uniformly available across firms and quarters in Refinitiv, while non-GAAP measures are voluntarily disclosed and not consistently reported.

**Size.** Natural logarithm of trailing-twelve-month revenue, measured in millions of USD. Source: Refinitiv Workspace. Three operationalizations of size are conceptually defensible— $\log(\text{Revenue})$ ,  $\log(\text{MarketCap})$ , and  $\log(\text{EV})$ —and each tests a distinct hypothesis.  $\log(\text{MarketCap})$  and  $\log(\text{EV})$  are mechanically correlated with the dependent variable  $\text{EV/Revenue}$ , since EV appears in both regressor and regressand; the resulting coefficient is partially identified by an algebraic identity rather than by an economic relationship and is therefore unsuitable as a primary regressor.  $\log(\text{Revenue})$  avoids this mechanical correlation entirely and is the size measure used in the practitioner SaaS literature (Bessemer Venture Partners 2025; Bain & Company 2024). It captures scale—the size of the installed revenue base—rather than market valuation, which is the economically interpretable quantity for the

size-loading prediction in Section 2.2.4. As a robustness check, Section 6 reports the main specification with  $\log(\text{Total Assets})$  as an alternative scale measure that is also independent of EV; results are reported in the appendix.

### **Robustness Variables**

**ForwardGrowth.** Median one-year-forward consensus revenue-growth estimate, expressed as a percentage. Source: Refinitiv IBES (TR.RevenueMean field).

**$\Delta r$ .** Quarterly change in the 10-year U.S. Treasury constant-maturity yield. Source: Federal Reserve H.15 release.

### **Period Indicator and Fixed Effects**

**PostShock.** Indicator equal to one for quarters from the candidate break date onward, zero otherwise. The primary specification uses the data-determined break of 2022Q1 (rate-cycle onset). The 2022Q4 specification, motivated by the November 30, 2022 release of ChatGPT, is reported alongside as a benchmark against the AI-led practitioner narrative.

**Firm Fixed Effects ( $\alpha_i$ ).** Absorb time-invariant firm-specific characteristics including business model, product positioning, and managerial quality.

**Calendar-Quarter Fixed Effects ( $\lambda_t$ ).** Absorb aggregate shocks affecting all firms in a given quarter, including macroeconomic conditions, interest rates, and market-wide valuation changes.

**Table A2. Rolling-breakpoint test: full sweep of all 12 candidate break dates, 2021Q1–2023Q4.**

*This table reports the rolling-breakpoint procedure described in Section 3.4 at every candidate calendar quarter from 2021Q1 through 2023Q4. The body table 6 reports the three economically distinct dates discussed in the main text (2022Q1, 2022Q4, 2023Q2). The SaaS columns report results on the SaaS panel of 90 firms; the placebo columns report results on the energy-and-utilities panel of 217 firms.*

| Break         | SaaS sample  |                          |                          | Energy & Utilities placebo |                          |                          |
|---------------|--------------|--------------------------|--------------------------|----------------------------|--------------------------|--------------------------|
|               | $\chi^2(3)$  | $\delta_{\text{margin}}$ | $\delta_{\text{growth}}$ | $\chi^2(3)$                | $\delta_{\text{margin}}$ | $\delta_{\text{growth}}$ |
| 2021Q1        | 12.44        | 0.292                    | 0.721***                 | 9.01                       | 0.169                    | −0.189***                |
| 2021Q2        | 11.15        | 0.468                    | 0.592***                 | 11.70                      | 0.095                    | −0.207***                |
| 2021Q3        | 14.68        | 0.661*                   | 0.588***                 | 14.83                      | 0.111                    | −0.214***                |
| 2021Q4        | 16.67        | 0.888***                 | 0.547**                  | 9.33                       | 0.086                    | −0.172***                |
| <b>2022Q1</b> | <b>16.96</b> | 1.067***                 | 0.477*                   | 9.49                       | 0.068                    | −0.166***                |
| 2022Q2        | 15.37        | 1.154***                 | 0.311                    | 6.79                       | −0.007                   | −0.135**                 |
| 2022Q3        | 14.01        | 1.171***                 | 0.174                    | 5.61                       | 0.044                    | −0.122*                  |
| 2022Q4        | 9.18         | 1.022***                 | 0.266                    | 3.15                       | 0.067                    | −0.072                   |
| 2023Q1        | 5.99         | 0.832**                  | 0.354                    | 1.74                       | 0.073                    | 0.009                    |
| 2023Q2        | 3.81         | 0.564                    | 0.534                    | 1.80                       | 0.123                    | 0.033                    |
| 2023Q3        | 5.54         | 0.553                    | 0.912**                  | 2.22                       | 0.088                    | 0.104                    |
| 2023Q4        | 9.04         | 0.606*                   | 1.146***                 | 3.29                       | 0.101                    | 0.152                    |

*Notes. \*, \*\*, and \*\*\* denote significance at the 10%, 5%, and 1% levels. The data-determined SaaS peak (bolded) is at 2022Q1,  $\chi^2(3) = 16.96$  ( $p = 0.001$ ), corresponding to the onset of the Federal Reserve tightening cycle. The placebo  $\chi^2(3)$  peaks at 2021Q3 driven by the growth coefficient, with the placebo  $\delta_{\text{margin}}$  indistinguishable from zero at every candidate date. Joint  $\chi^2(3)$  is the test that all three equation (2) interactions equal zero. The 5% and 1% critical values for  $\chi^2(3)$  are 7.81 and 11.34 respectively.*

**Table A3. Equation (1) baseline regression estimated separately on the pre-shock and post-shock subsamples.**

*This appendix table reports estimates of equation (1) re-fitted on three subsamples: the pre-shock window 2014Q1–2022Q3 (column 1), the post-shock window 2022Q4–2024Q4 (column 2), and the full sample (column 3, the baseline of Table 2). All three columns include firm and calendar-quarter fixed effects; standard errors are double-clustered at the firm and calendar-quarter levels. This subsample comparison is a complement to the structural-change specification of Table 3, which tests for changes in driver loadings within a single regression augmented with PostShock interactions; the present subsample comparison fits separate regressions on each window and reports the corresponding subsample within-R<sup>2</sup> values for transparency.*

|                       | (1)                 | (2)                  | (3)                 |
|-----------------------|---------------------|----------------------|---------------------|
|                       | Pre-shock subsample | Post-shock subsample | Full sample         |
|                       | 2014Q1–2022Q3       | 2022Q4–2024Q4        | 2014Q1–2024Q4       |
| Growth                | 1.018***<br>(0.182) | 1.476***<br>(0.362)  | 1.479***<br>(0.213) |
| Operating margin      | 0.429<br>(0.450)    | 0.442<br>(0.450)     | 0.107<br>(0.297)    |
| ln(Revenue)           | 0.016<br>(0.131)    | 0.200<br>(0.345)     | –0.096<br>(0.087)   |
| Within R <sup>2</sup> | 0.1087              | 0.1504               | 0.1694              |
| Observations          | 1,341               | 753                  | 2,094               |

*Notes. \*, \*\*, and \*\*\* denote significance at the 10%, 5%, and 1% levels, respectively. The post-shock growth loading rises from +1.018 in the pre-shock subsample to +1.477 in the post-shock subsample, while the operating-margin loading is statistically indistinguishable from zero in both subperiods. The post-shock subsample within-R<sup>2</sup> of 0.150 is higher than the pre-shock value of 0.109, but a direct comparison is not advisable: the firm and calendar-quarter fixed effects are re-estimated on disjoint windows and reference different demeaning means, so the within-R<sup>2</sup> in each subsample column is identified relative to a different absorption baseline. The structural-change specification of Table 3, which imposes common firm intercepts across periods through PostShock interactions on a single regression, is the formal test of the coefficient-stability hypotheses H1 and H2 and is the primary specification for inference.*