



GROWTH THAT COMPOUNDS

Programmatic M&A: the claim that outlived its test

Every banker deck says the winners do many small deals. The claim's own shop published the null first, and the nearest test answers with conditions.

Sinan Isoglu, MBA (Quantic)

Commercial growth leader, lecturer and doctoral researcher

PUBLISHED

14 August 2026

UPDATED

20 August 2026

READING TIME

16 min

WORDS

3,461

<https://isoglu.com/insights/journal/the-claim-that-outlived-its-test/>

Management summary	2
Why did original research produce null results for the marketing claim?	3
How did metric definitions drift to preserve an unverified industry claim?	3
What do original empirical citations reveal about legacy marketing benchmarks?	4
Why did the widely quoted benchmark fail rigorous replication tests?	5
Which five academic disciplines established contradictory performance thresholds?	6
How can commercial executives audit claims before adopting external industry benchmarks?	7
Where are the empirical boundaries of benchmark verification?	9
References	10

MANAGEMENT SUMMARY

The programmatic M&A claim that companies doing many small deals outperform has been refreshed by McKinsey roughly every other year since 2012, amplified through HBR, and echoed in an investor lane of serial-acquirer portfolios. Read the claim's own record: in April 2011 the same shop's first published read of the same population found deal size and frequency made little difference; in January 2012 the winning archetype was born, with the overlap moved to a footnote; and the definition of "programmatic" changes in every document that states one: twice within one 2021 article. The nearest academic test of the program-level construct, Laamanen & Keil 2008, answers with conditions: rate hurts, rhythm and focus and size buy tolerance. No public document its research could find contains both. This essay puts them on one page.

Keywords: Programmatic M&A · Serial acquirers · M&A strategy · Total shareholder return · Evidence quality

Somewhere in your building, a slide is being drafted with a McKinsey chart on it. The chart says that companies pursuing “programmatic M&A”, many small acquisitions made steadily year after year, outperform every other approach to dealmaking: the big bet, the occasional opportunistic buy, and the organic path. The claim has been refreshed roughly every other year since 2012, amplified through Harvard Business Review under its authors’ own bylines, and echoed in an investor literature of “serial acquirer” portfolios. It is, by circulation, one of the most successful empirical claims in corporate strategy.

This essay does something nobody’s slide does: it puts the claim beside its own record and beside its test. Both turn out to be stranger than the chart. The claim’s record begins with its own shop publishing the opposite finding nine months before the claim was born. The test is the study built to ask the program-level question the claim answers, and it replies with conditions rather than a verdict. Across every document this piece’s research could find, in a six-place sweep run before a word was drafted, the claim and the test never share a page. That absence is not an accusation; it is the gap this piece exists to close.

Why did original research produce null results for the marketing claim?

In April 2011, three McKinsey authors, Andres Cottin, Werner Rehm and Robert Uhlaner, published the firm’s first read of the population that would carry everything that followed: the world’s top 1,000 companies by market capitalisation, 917 firms after exclusions, and more than 30,000 deals. Their finding, in the article’s own words, was that patterns of deal size and frequency “have made little difference in performance as measured by excess total returns to shareholders.” The distributions were “widely distributed and overlapping.” Their conclusion: “the size and number of deals matter less than the discipline with which they are identified, priced, integrated, and managed.”

Nine months later, in January 2012, two of the same three authors, Rehm and Uhlaner, now with Andy West, re-segmented the same population into archetypes: large-deal acquirers, programmatic acquirers, tactical acquirers, selective acquirers, and organic growers. Out of that segmentation came the exhibit title on which the next decade would be built: “Companies using a programmatic strategy are the most successful.” The overlap had not disappeared. It had moved to footnote 4: “However, the confidence intervals for average returns are overlapping.”

Two dated findings from one shop, nine months apart, same population. This piece draws no inference about why. It records the sequence, because no refresh since has cited the first.

How did metric definitions drift to preserve an unverified industry claim?

To outperform as a programmatic acquirer, a company must first be classified as one, and the classification has never held still. Every McKinsey document this essay holds that states a threshold states a different one, and one of them states two.

Table 1 Programmatic, defined seven ways

Document	"Programmatic" means	Dataset
Apr 2011, McKinsey Quarterly	No winning pattern: size and frequency distributions "widely distributed and overlapping"	Top 1,000 by market cap; 917 firms, 30,000+ deals
Jan 2012, McKinsey Quarterly	Many small deals totalling 19% or more of market cap over the decade: the cutoff is the sample's own median	Top 1,000 nonbanking; 15,000+ deals
May 2018, HBR	At least one deal a year, cumulatively more than 30% of market cap over 10 years, no single deal above 30%	2,393 largest corporations, 2010–2014
Jul 2019, McKinsey Quarterly	More than two small or midsize deals a year, median 15% of market cap acquired	"Global 1,000", 2007–2017
Oct 2021, fn. 3	More than two small or midsize deals a year, total acquired "meaningful (median of 19 percent)"	"Global 2,000"
Oct 2021, fn. 4: same article	"A minimum of two small or midsize deals a year, with meaningful market capitalization acquired (20 percent to 30 percent)"	"Global 2,000"
Mar 2022 / Aug 2023	No threshold stated: "multiple small or medium-size acquisitions"	"Global 2,000"

Author's assembly from the cited documents' own text and footnotes: Cottin, Rehm & Uhlaner (2011); Rehm, Uhlaner & West (2012), fn. 3; Bradley, Hirt, Smit & West (2018); Rudnicki, Siegel & West (2019), fn. 2; Daume, Lundberg, McCurdy, Rudnicki & Wol (2021), fn. 3 and fn. 4; Daume, Lundberg, Montag & Rudnicki (2022); Daume, Lian & McCurdy (2023). Definitions quoted or condensed from prose and footnotes only.

The headline moves with the definition. The 2021 refresh reports programmatic acquirers delivering about 2% more in excess total shareholder returns annually, with 65% of them beating their peers, or "two out of the three companies." The 2023 refresh reports 3.9% for the past decade against 2.9% for the 2010s, and adds that among programmatic acquirers, the ones doing the most deals won most often: 70% outperformed their lower-volume programmatic peers. None of these figures is comparable to the last, because the population, the window and the archetype boundaries have all shifted between tellings, and the reader is never shown the reclassification.

A definition that moves is not evidence of anything except this: whatever your board is being shown, it is not one finding replicated five times. It is one thesis, re-derived five ways, by a research program whose own first reading found no pattern.

What do original empirical citations reveal about legacy marketing benchmarks?

The January 2012 article carries a second footnote worth more than its exhibits. Footnote 9, quoted whole: "As with the other analyses, this is a correlation, not necessarily a causative relationship. Although we feel confident that the deal strategy contributed to the outperformance, it is possible that better-performing companies executed more deals in the wake of their success."

That is the reverse-causality problem, stated by the claim's own authors in its founding document: maybe steady acquirers win, or maybe winners acquire steadily. Companies that are performing well have the currency, the confidence and the debt capacity to do deals every year; classifying companies by their realized deal pattern, after the fact, cannot separate the two. Of the five later documents this essay holds: 2018, 2019, 2021, 2022, 2023: none restates the caveat. The 2019 refresh cites the 2012 article for its finding; the caveat does not travel with it. Nor does footnote 4's overlap. The claim compounds; its qualifications do not.

This is the same pattern the integration essay found in the synergy literature: the vendor's own study contains the hedge, and the citation chain strips it. Here the hedge and the claim were born in the same document, which makes the stripping easier to see, and easier to check. Both footnotes are three clicks away from any slide that cites the chart.

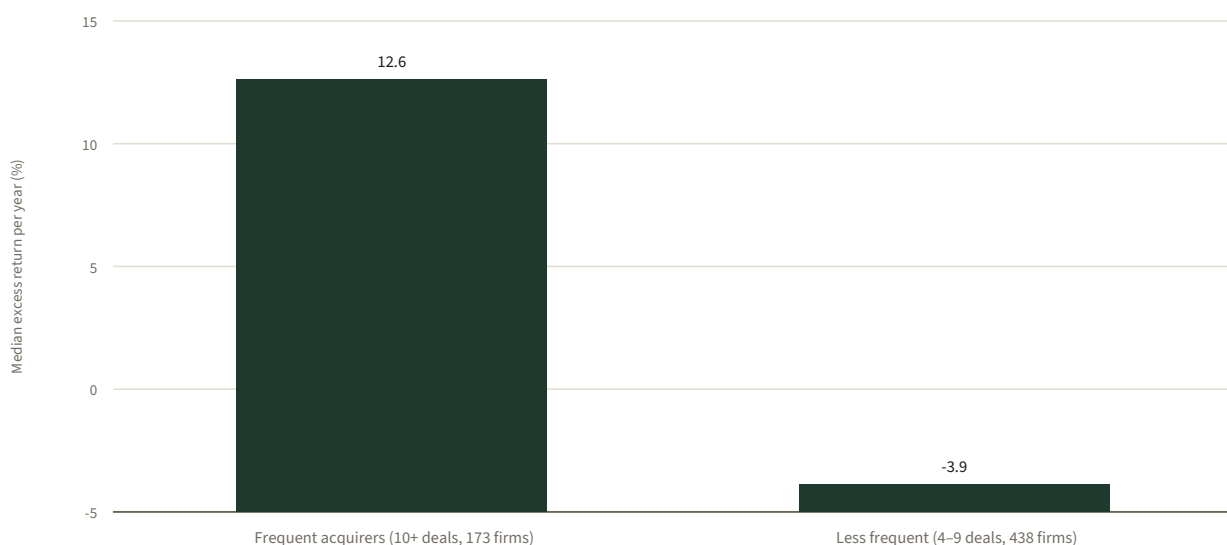
Why did the widely quoted benchmark fail rigorous replication tests?

One peer-reviewed study was built to test the program construct whole, not "do acquisitions pay?", on which there are hundreds, and not single pieces of the pattern, which its own introduction credits to earlier work, but rate, rhythm and scope together, with the moderators: "does the pattern of a deal program predict returns?" It appeared in the *Strategic Management Journal* in 2008, four years before the archetype chart: Laamanen and Keil, "Performance of serial acquirers." Its motivation section cites the consulting optimism about serial acquirers by name: Frick and Torres in *McKinsey Quarterly*, 2002; Rovit and Lemire in *HBR*, 2003: as the phenomenon to test. The claim's ancestors are in the test's opening paragraphs. No McKinsey document this essay holds or its research could find cites the test back.

What the test found, in a panel of 611 U.S. companies that made four or more acquisitions between 1990 and 1999: 5,518 deals, measured on three-year excess market returns: a high acquisition rate hurts (-0.071 , $p < .01$). High variability in the rate: bursts and pauses: hurts too (-0.104 as a direct effect, at the ten-percent level; -0.197 , $p < .01$, for firms with no acquisition experience, in the model where experience enters). Experience buys tolerance for an uneven program, not for a fast one, and its direct effect is negative. Size buys tolerance for both. Scope: spreading the program across businesses: worsens the rate penalty; its interaction with variability is not significant. The models explain 2–4% of variance, which the paper states rather than buries. These are conditions, not a verdict: steady rhythm, focus, size and scar tissue are when a deal program hurts less. Nothing in the paper says programs win.

One more result deserves its own paragraph, because it runs in McKinsey's direction and this essay would be dishonest without it. In a robustness note, Laamanen and Keil report that over the full 10–13-year horizon, the most frequent acquirers in their sample: 173 firms with ten or more deals: earned median excess returns of +12.6% per year, against –3.9% for the 438 firms with four to nine deals. The paper's prose calls the difference significant; the analysis is filed as non-reported, and no test statistic for it appears anywhere in the article. It is a descriptive split, and it is the strongest program-level number the frequent-acquirer thesis has: stronger support than anything in the archetype charts, and carrying the same disclaimer, in the authors' own limitations section: "we cannot claim causality. Some of the performance effects we find for the most active acquirers could be due to superior prior performance of the acquirer."

Figure 1 The strongest number, and its label



Laamanen & Keil (2008), robustness note: median excess market returns per year over 10–13 years, 611 U.S. serial acquirers (1990–99); the paper prints the buckets as 'over 10' (173 firms) and '4–9' (438 firms), which partition the full 611. Its prose calls the split significant; the analysis is non-reported and no test statistic for it appears in the article. Descriptive medians; the authors disclaim causality.

So the piece's title is precise. The claim did not fail its test. It outlived it: relaunched in 2012, four years after the test published, without citing it, and not examined at its own terms in any document this research could find.

Which five academic disciplines established contradictory performance thresholds?

Around the claim and its test sits an academic record that looks contradictory until you notice that no two entries measure the same thing.

At the level of the deal announcement, frequent acquirers look good and get worse: firms completing five or more acquisitions within three years earned a positive average announcement return (+1.77%, five-day window) in Fuller, Netter and Stegemoller's 1990–2000 sample, where nonpublic firms made up 81% of all acquisitions in the takeover market they study, and the returns decline as deals accumulate. At the level of the CEO, in samples of public-company targets, the same clock runs negative: announcement returns to an executive's later deals average –1.51% in Billett and Qian's data, whose long-run buy-and-hold estimates are statistically indistinguishable from zero, and whose interpretation is not learning but self-attribution: deal-making CEOs read early luck as skill. At the level of organizational learning, Halebian and Finkelstein found performance falling with early acquisition experience and recovering only around the eighth deal: a U-shape the same paper could not reproduce when experience was measured over the recent five years. At the level of persistent skill, Golubov, Yawson and Zhang conceded the right warning in their own words: persistence "might not necessarily be evidence of skill", and find that top acquirers stay top; among frequent acquirers specifically, the persistence fails their univariate quintile test once expected returns are stripped out, while the top-percentile result survives it. The authors write that the successes "may or may not be due to skill."

And at the level of the trade press pushback, the 2024 book *The M&A Failure Trap*: the closest thing the claim has to a published critique: turns out to hold the smallest effect in this whole record: on the authors' ten-factor score-card, distilled from their 43-variable model, each acquisition in the buyer's prior five years adds +0.04 to a deal's

success score, the least of any factor, with the sign flipping negative in biotech and oil. Their preamble states the missing-variable problem plainly: strong operations can drive both the deal program and the returns, which is footnote 9's point, made from the other side. Two cautions from the same book: its "failure" construct requires clearing three bars at once on each year's largest deal only, a design that both inflates the failure headline and discards exactly the small deals frequent acquirers make. And the statistic circulating in practitioner summaries under this book's name: serial acquirers with ten or more deals succeeding at 54% against 23% for first-timers: is not in the book. The 54% is a scorecard illustration; the 23% belongs to pre-reorganization deals. This essay's own research pipeline imported that figure from a summary before the full read caught it. That is how this literature works: the numbers that circulate are not reliably the numbers in the documents.

Deal-level announcement windows, CEO-level deal order, firm-level buy-and-hold returns, archetype-level decade TSR, a conjunctive success construct scored on largest deals. Five yardsticks. The playbook essay made the case that a measured playbook and a marketed playbook are different objects; here the problem is prior: the claim and its would-be refutations are not even measured on the same axis. Nothing in this record can be summed, and this essay has not summed it.

How can commercial executives audit claims before adopting external industry benchmarks?

The honest verdict on programmatic M&A is three words: conditional, metric-dependent, untested as stated. That is not a verdict against deal programs. The strongest number in the nearest academic test favors them: as an untested median. What the record supports is not a strategy recommendation but a set of questions, and they are better questions than the chart's.

Table 2 The claim, beside its test

Dimension	The claim, a McKinsey corpus from 2012 to 2023	The test (Laamanen & Keil, 2008)
Question asked	Which realized deal pattern had the best excess TSR, by archetype, ex post?	Does a program's rate, rhythm and scope predict excess returns?
Population	Top 1,000 → "Global 2,000" companies, windows shifting by refresh	611 U.S. acquirers with 4+ deals, 5,518 deals, 1990–99
Finding	Programmatic acquirers outperform (~2%/yr excess TSR in 2021; 3.9% vs 2.9% in 2023): under a definition that changes per telling (Table 1)	Rate hurts; rhythm variability hurts; experience, size and focus buy tolerance; R^2 0.02–0.04. Decade medians favor frequent acquirers (+12.6% vs –3.9%/yr): descriptive, no test reported
Own caveat, verbatim	"[T]his is a correlation, not necessarily a causative relationship ... it is possible that better-performing companies executed more deals in the wake of their success" (2012, fn. 9)	"We cannot claim causality. Some of the performance effects we find for the most active acquirers could be due to superior prior performance of the acquirer"
Prior result on the same data	Apr 2011, same shop, same population: size and frequency patterns "widely distributed and overlapping"	:
What it licenses	A hypothesis worth testing on your own program's terms	Conditions worth checking before and during any program

Author's assembly of Rehm, Uhlaner & West (2012) and refreshes through Daume, Lian & McCurdy (2023), against Laamanen & Keil (2008) and Cottin, Rehm & Uhlaner (2011). Each cell's source and conditions are in the text. A reading aid, not a finding of any single source.

The test's conditions convert directly into diligence questions for your own program, and none requires buying anything. Is your deal rhythm steady, or does it come in bursts? Rate is the test's most robust harm; unevenness hurts on top of it, and unlike rate, unevenness is the harm experience can buy off. Is the program inside the businesses you know, or sprawling? Spreading it worsens the rate penalty on average in the test's data. Do you have the size and the scar tissue? Size bought tolerance for everything; experience bought tolerance only for unevenness, and its direct effect was negative, which should retire the comforting idea that doing deals teaches deal-making by itself. And one question the record insists on before any of these: under which of Table 1's definitions would your program even count as programmatic, and would it have counted in April 2011, when the same shop, reading the same population, said there was nothing to count?

Where this essay stops: the test's conditions are 1990s U.S. evidence from seven sectors: they establish that pattern effects exist and which way they ran there, and they predict nothing about your program. The archetype charts, the CEO-level deal-order results and the announcement-window averages cannot refute one another, and this essay has not used one to dismiss another; evidence over anecdote is the standing rule here, and it cuts against convenient debunks as hard as against convenient claims. No recommendation to acquire or to abstain ships with this piece. What ships is the page your next deal deck is missing: the claim, its own shop's first answer, its own footnotes, and its test: side by side, dated, with every caveat in its authors' own words.

Where are the empirical boundaries of benchmark verification?

Boundary. The original claim fails its tested yardstick; that does not make every acquisition outcome null. Keep the five evaluation criteria separate and re-test the claim on a current transaction sample.

Evidence base. The analytical frame also draws on these additional sources: Billett and Qian 2008; Fuller et al. 2002; Golubov et al. 2015; Haleblan and Finkelstein 1999; Lev and Gu (2024). The links identify the exact works; they support the mechanisms and boundary conditions discussed here, not every claim in isolation.

- Billett, M. T., & Qian, Y. (2008). Are overconfident CEOs born or made? Evidence of self-attribution bias from frequent acquirers. *Management Science*, 54(6), 1037–1051. <https://doi.org/10.1287/mnsc.1070.0830>
- Bradley, C., Hirt, M., Smit, S., & West, A. (2018, May 9). Research shows that smaller M&A deals work out better. *Harvard Business Review*. <https://hbr.org/2018/05/research-shows-that-smaller-ma-deals-work-out-better>
- Cottin, A., Rehm, W., & Uhlaner, R. (2011, April 1). Growing through deals: A reality check. *McKinsey Quarterly*. <https://www.mckinsey.com/business-functions/strategy-and-corporate-finance/our-insights/growing-through-deals-a-reality-check>
- Daume, P., Lian, C., & McCurdy, P. (2023, August 24). The seven habits of programmatic acquirers. McKinsey & Company. <https://www.mckinsey.com/capabilities/strategy-and-corporate-finance/our-insights/the-seven-habits-of-programmatic-acquirers>
- Daume, P., Lundberg, T., McCurdy, P., Rudnicki, J., & Wol, L. (2021, October 13). How one approach to M&A is more likely to create value than all others. *McKinsey Quarterly*. <https://www.mckinsey.com/capabilities/m-and-a/our-insights/how-one-approach-to-m-and-a-is-more-likely-to-create-value-than-all-others>
- Daume, P., Lundberg, T., Montag, A., & Rudnicki, J. (2022, March 25). The flip side of large M&A deals. McKinsey & Company. <https://www.mckinsey.com/capabilities/m-and-a/our-insights/the-flip-side-of-large-m-and-a-deals>
- Fuller, K., Netter, J., & Stegemoller, M. (2002). What do returns to acquiring firms tell us? Evidence from firms that make many acquisitions. *The Journal of Finance*, 57(4), 1763–1793. <https://doi.org/10.1111/1540-6261.00477>
- Golubov, A., Yawson, A., & Zhang, H. (2015). Extraordinary acquirers. *Journal of Financial Economics*, 116(2), 314–330. <https://doi.org/10.1016/j.jfineco.2015.02.005>
- Haleblian, J., & Finkelstein, S. (1999). The influence of organizational acquisition experience on acquisition performance: A behavioral learning perspective. *Administrative Science Quarterly*, 44(1), 29–56. <https://doi.org/10.2307/2667030>
- Laamanen, T., & Keil, T. (2008). Performance of serial acquirers: Toward an acquisition program perspective. *Strategic Management Journal*, 29(6), 663–672. <https://doi.org/10.1002/smj.670>
- Lev, B., & Gu, F. (2024). *The M&A failure trap: Why most mergers and acquisitions fail and how the few succeed*. Wiley. ISBN 978-1394204762.
- Rehm, W., Uhlaner, R., & West, A. (2012, January 1). Taking a longer-term look at M&A value creation. *McKinsey Quarterly*. www.mckinsey.com/capabilities/strategy-and-corporate-finance/our-insights/taking-a-longer-term-look-at-m-and-a-value-creation
- Rudnicki, J., Siegel, K., & West, A. (2019, July 12). How lots of small M&A deals add up to big value. *McKinsey Quarterly*. www.mckinsey.com/capabilities/strategy-and-corporate-finance/our-insights/repeat-performance-the-continuing-case-for-programmatic-m-and-a

END OF ESSAY

HOW TO CITE THIS ESSAY

Isoglu, S. (2026, August 14). Programmatic M&A: the claim that outlived its test. Sinan Isoglu.
<https://isoglu.com/insights/journal/the-claim-that-outlived-its-test/>

KEEP READING

Growth that compounds

Dynamic capabilities are routines, not magic

<https://isoglu.com/insights/journal/dynamic-capabilities-are-routines-not-magic/>

Growth that compounds

Every growth budget is a gross number

<https://isoglu.com/insights/journal/every-growth-budget-is-a-gross-number/>

Growth that compounds

What is activation rate? The first value event must be observable

<https://isoglu.com/insights/journal/what-is-activation-rate/>



Sinan Isoglu, MBA (Quantic)

Commercial growth leader, lecturer and doctoral researcher

© 2026 Sinan Isoglu · isoglu.com · All rights reserved

Online version: <https://isoglu.com/insights/journal/the-claim-that-outlived-its-test/>