



WACHSTUM, DAS SICH POTENZIERT

Programmatische M&A: die Behauptung, die ihren Test überlebte

Jedes Deal-Deck sagt: Gewinner machen viele kleine Deals. Im eigenen Haus fand sich nichts – und der nächste Test antwortet mit deutlichen Bedingungen.

Sinan Isoglu, MBA (Quantic)

Führungskraft für Umsatzwachstum, Dozent und Doktorand

VERÖFFENTLICHT

14. August 2026

AKTUALISIERT

20. August 2026

LESEZEIT

16 Min.

WÖRTER

3.438

<https://isoglu.com/de/insights/journal/die-behauptung-die-ihren-test-ueberlebte/>

Management Summary	2
Warum lieferte die ursprüngliche Untersuchung eine eindeutige Nullmeldung?	3
Wie verschoben sich Definitionen, um die unbestätigte These zu stützen?	3
Was offenbaren historische Originalzitate über etablierte Marketing-Benchmarks?	5
Warum hielt die populäre Faustformel wissenschaftlichen Replikationstests nicht stand?	5
Welche fünf Fachdisziplinen etablierten widersprüchliche Leistungskriterien?	6
Wie sollten Führungskräfte externe Benchmarks vor deren Einführung prüfen?	7
Wo liegen die empirischen Grenzen der Benchmark-Überprüfung?	9
Literatur	10

MANAGEMENT SUMMARY

Programmatische M&A – die Behauptung, dass Firmen mit vielen kleinen Deals besser abschneiden – wird von McKinsey seit 2012 etwa alle zwei Jahre aufgefrischt und über HBR verstärkt. Die eigene Aktenlage: Im April 2011 fand die erste publizierte Lesart derselben Population im selben Haus, dass Deal-Größe und -Frequenz kaum einen Unterschied machen; im Januar 2012 wurde der Gewinner-Archetyp geboren, die Überlappung in eine Fußnote verschoben; und die Definition von „programmatisch“ ändert sich in jedem Dokument, das eine nennt – zweimal innerhalb eines einzigen Artikels von 2021. Der nächstgelegene akademische Test des Programm-Konstrukts, Laamanen & Keil 2008, antwortet mit Bedingungen: Tempo schadet, Rhythmus, Fokus und Größe kaufen Toleranz. Kein auffindbares öffentliches Dokument enthält beides zugleich. Dieser Essay legt sie auf eine Seite.

Schlagwörter: Programmatische M&A · Serielle Käufer · M&A-Strategie · Total Shareholder Return · Evidenzqualität

Irgendwo in Ihrem Haus entsteht gerade eine Folie mit einem McKinsey-Chart darauf. Das Chart sagt, dass Unternehmen mit „programmatischer M&A“ – viele kleine Akquisitionen, stetig, Jahr für Jahr – jeden anderen Deal-Ansatz übertreffen: die große Wette, den gelegentlichen Opportunitätskauf, den organischen Weg. Die Behauptung wird seit 2012 etwa alle zwei Jahre aufgefrischt, im Harvard Business Review unter den Bylines ihrer eigenen Autoren verstärkt und in einer Investorenliteratur der „Serial-Acquirer“-Portfolios weitergetragen. Nach Verbreitung gemessen ist sie eine der erfolgreichsten empirischen Behauptungen der Unternehmensstrategie.

Dieser Essay tut, was keine dieser Folien tut: Er legt die Behauptung neben ihre eigene Aktenlage und neben ihren Test. Beides ist merkwürdiger als das Chart. Die Aktenlage beginnt damit, dass das eigene Haus neun Monate vor der Geburt der Behauptung das gegenteilige Ergebnis publizierte. Der Test – die Studie, die gebaut wurde, um die Programm-Frage der Behauptung zu stellen – antwortet mit Bedingungen statt mit einem Urteil. Und in jedem Dokument, das die Recherche dieses Stücks finden konnte – ein Sweep über sechs Orte, gelaufen, bevor ein Wort entworfen war –, teilen Behauptung und Test nie eine Seite. Diese Abwesenheit ist kein Vorwurf; sie ist die Lücke, für deren Schließung dieses Stück existiert.

Warum lieferte die ursprüngliche Untersuchung eine eindeutige Nullmeldung?

Im April 2011 publizierten drei McKinsey-Autoren – Andres Cottin, Werner Rehm und Robert Uhlaner – die erste Lesart der Population, die alles Folgende tragen sollte: die 1.000 größten Unternehmen der Welt nach Marktkapitalisierung, 917 Firmen nach Ausschlüssen, mehr als 30.000 Deals. Ihr Befund, in den Worten des Artikels: Muster von Deal-Größe und -Frequenz „have made little difference in performance“ – haben für die Performance, gemessen am Excess Total Shareholder Return, kaum einen Unterschied gemacht. Die Verteilungen: „widely distributed and overlapping“ – breit gestreut und überlappend. Ihr Schluss: Größe und Zahl der Deals zählen weniger als die Disziplin, mit der sie identifiziert, bepreist, integriert und geführt werden.

Neun Monate später, im Januar 2012, re-segmentierten zwei der drei Autoren – Rehm und Uhlaner, nun mit Andy West – dieselbe Population in Archetypen: Large-Deal-Käufer, programmatische Käufer, taktische Käufer, selektive Käufer, organische Wachser. Aus dieser Segmentierung kam der Exhibit-Titel, auf den das nächste Jahrzehnt gebaut wurde: „Companies using a programmatic strategy are the most successful“ – Unternehmen mit programmatischer Strategie sind die erfolgreichsten. Die Überlappung war nicht verschwunden. Sie war in Fußnote 4 gezogen: „However, the confidence intervals for average returns are overlapping“ – die Konfidenzintervalle der Durchschnittsrenditen überlappen.

Zwei datierte Befunde aus einem Haus, neun Monate auseinander, dieselbe Population. Dieses Stück zieht keinen Schluss über das Warum. Es protokolliert die Abfolge – denn keine Auffrischung seither hat den ersten zitiert.

Wie verschoben sich Definitionen, um die unbestätigte These zu stützen?

Um als programmatischer Käufer besser abzuschneiden, muss ein Unternehmen zuerst als einer klassifiziert werden – und die Klassifikation hat nie stillgehalten. Jedes McKinsey-Dokument, das dieser Essay hält und das eine Schwelle nennt, nennt eine andere, und eines nennt zwei.

Tabelle 1 Programmatisch, siebenfach definiert

Dokument	„Programmatisch“ heißt	Datensatz
Apr 2011, McKinsey Quarterly	Kein Gewinnermuster – Größen- und Frequenzverteilungen „widely distributed and overlapping“	Top 1.000 nach Marktkapitalisierung; 917 Firmen, 30.000+ Deals
Jan 2012, McKinsey Quarterly	Viele kleine Deals, zusammen 19 % oder mehr der Marktkapitalisierung über die Dekade – der Cutoff ist der eigene Median der Stichprobe	Top 1.000 ohne Banken; 15.000+ Deals
Mai 2018, HBR	Mindestens ein Deal pro Jahr, kumuliert mehr als 30 % der Marktkapitalisierung über 10 Jahre, kein Einzeldeal über 30 %	2.393 größte Unternehmen, 2010–2014
Jul 2019, McKinsey Quarterly	Mehr als zwei kleine oder mittlere Deals pro Jahr, Median 15 % der Marktkapitalisierung erworben	„Global 1.000“, 2007–2017
Okt 2021, Fn. 3	Mehr als zwei kleine oder mittlere Deals pro Jahr, Gesamterwerb „meaningful (median of 19 percent)“	„Global 2.000“
Okt 2021, Fn. 4 – derselbe Artikel	„Ein Minimum von zwei kleinen oder mittleren Deals pro Jahr“, erworbene Marktkapitalisierung 20 bis 30 %	„Global 2.000“
Mär 2022 / Aug 2023	Keine Schwelle genannt – „multiple small or medium-size acquisitions“	„Global 2.000“

Eigene Zusammenstellung aus Text und Fußnoten der zitierten Dokumente: 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 und Fn. 4; Daume, Lundberg, Montag & Rudnicki (2022); Daume, Lian & McCurdy (2023). Definitionen ausschließlich aus Fließtext und Fußnoten zitiert oder kondensiert.

Die Schlagzeile wandert mit der Definition. Die Auffrischung von 2021 berichtet für programmatische Käufer rund 2 % mehr Excess Total Shareholder Return pro Jahr, mit 65 % Gewinneranteil – „two out of the three companies“, zwei von drei Unternehmen. Die Auffrischung von 2023 berichtet 3,9 % für die vergangene Dekade gegen 2,9 % in den 2010ern – und ergänzt, dass unter den programmatischen Käufern die mit den meisten Deals am häufigsten gewannen: 70 % schlugen ihre programmatischen Peers mit weniger Deals. Keine dieser Zahlen ist mit der vorigen vergleichbar, denn Population, Fenster und Archetyp-Grenzen haben sich zwischen den Fassungen verschoben – und die Reklassifikation bekommt der Leser nie gezeigt.

Eine wandernde Definition belegt nichts außer diesem: Was immer Ihrem Board gezeigt wird, es handelt sich nicht um einen Befund, der fünfmal repliziert wurde. Es ist eine These, fünfmal neu hergeleitet, von einem Forschungsprogramm, dessen eigene erste Lesart kein Muster fand.

Was offenbaren historische Originalzitate über etablierte Marketing-Benchmarks?

Der Artikel vom Januar 2012 trägt eine zweite Fußnote, die mehr wert ist als seine Exhibits. Fußnote 9, vollständig: „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.“ – Dies ist eine Korrelation, nicht notwendig eine kausale Beziehung. Obwohl wir zuversichtlich sind, dass die Deal-Strategie zur Outperformance beitrug, ist es möglich, dass besser performende Unternehmen im Zuge ihres Erfolgs mehr Deals machten.

Das ist das Problem der umgekehrten Kausalität, formuliert von den Autoren der Behauptung selbst, in ihrem Gründungsdokument: Vielleicht gewinnen stetige Käufer – vielleicht kaufen Gewinner stetig. Unternehmen, die gut laufen, haben die Währung, das Selbstvertrauen und die Verschuldungskapazität für jährliche Deals; wer Unternehmen nachträglich nach ihrem realisierten Deal-Muster klassifiziert, kann beides nicht trennen. Von den fünf späteren Dokumenten, die dieser Essay hält – 2018, 2019, 2021, 2022, 2023 –, wiederholt keines den Vorbehalt. Die Auffrischung von 2019 zitiert den Artikel von 2012 für seinen Befund; der Vorbehalt reist nicht mit. Fußnote 4 mit der Überlappung auch nicht. Die Behauptung verzinst sich; ihre Einschränkungen nicht.

Es ist dasselbe Muster, das der Integrations-Essay in der Synergie-Literatur fand: Die eigene Studie des Anbieters enthält den Vorbehalt, die Zitierkette streift ihn ab. Hier wurden Vorbehalt und Behauptung im selben Dokument geboren – was das Abstreifen sichtbarer macht. Und prüfbarer: Beide Fußnoten sind von jeder Folie, die das Chart zitiert, drei Klicks entfernt.

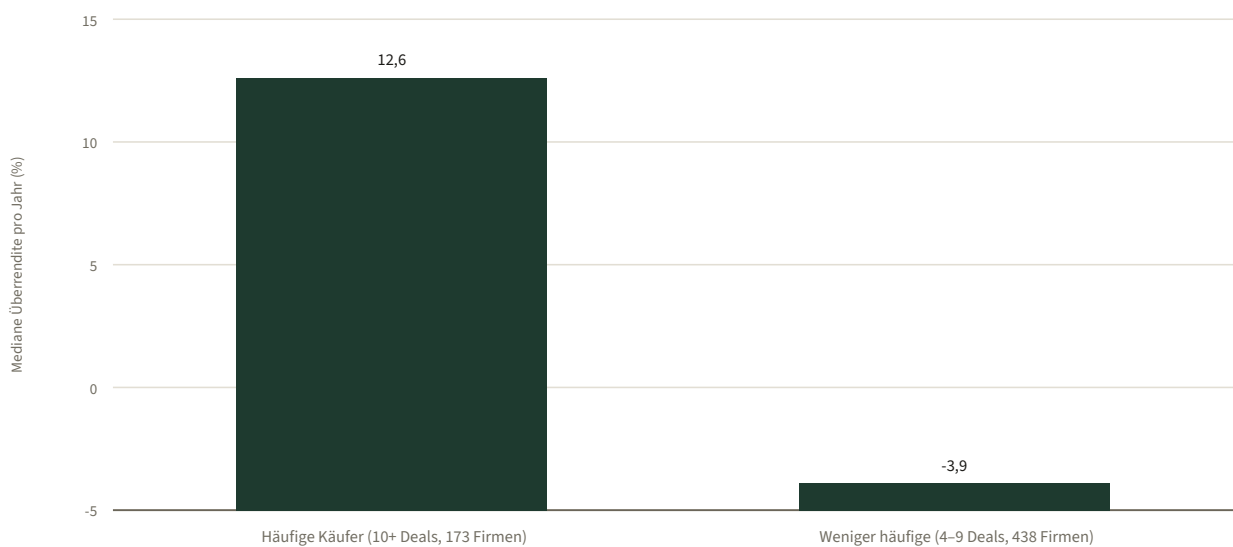
Warum hielt die populäre Faustformel wissenschaftlichen Replikationstests nicht stand?

Eine begutachtete Studie wurde gebaut, um das Programm-Konstrukt als Ganzes zu testen – nicht „lohnen sich Akquisitionen?“, dazu gibt es Hunderte, und nicht einzelne Teile des Musters, die ihre eigene Einleitung früherer Arbeit gutschreibt, sondern Tempo, Rhythmus und Streuung zusammen, mit den Moderatoren: „sagt das Muster eines Deal-Programms Renditen voraus?“. Sie erschien 2008 im Strategic Management Journal, vier Jahre vor dem Archetyp-Chart: Laamanen und Keil, „Performance of serial acquirers“. Ihre Motivation zitiert den Beratungsoptimismus über serielle Käufer beim Namen – Frick und Torres im McKinsey Quarterly 2002, Rovit und Lemire im HBR 2003 – als das zu testende Phänomen. Die Vorfahren der Behauptung stehen in den Eröffnungsabsätzen des Tests. Kein McKinsey-Dokument, das dieser Essay hält oder das seine Recherche finden konnte, zitiert den Test zurück.

Was der Test fand, in einem Panel von 611 US-Unternehmen mit vier oder mehr Akquisitionen zwischen 1990 und 1999 – 5.518 Deals, gemessen an Drei-Jahres-Überrenditen: Ein hohes Akquisitions-Tempo schadet ($-0,071$, $p < 0,01$). Hohe Variabilität des Tempos – Schübe und Pausen – schadet ebenfalls ($-0,104$ als Direkteffekt, auf Zehn-Prozent-Niveau; $-0,197$, $p < 0,01$, für Firmen ohne Akquisitionserfahrung im Modell, in dem Erfahrung eingeht). Erfahrung kauft Toleranz für ein unstetes Programm, nicht für ein schnelles – ihr Direkteffekt ist negativ. Größe kauft Toleranz für beides. Streuung über Geschäftsfelder verschärft den Tempo-Malus; ihre Interaktion mit der Variabilität ist nicht signifikant. Die Modelle erklären 2–4 % der Varianz – das Papier sagt es, statt es zu vergraben. Das sind Bedingungen, kein Urteil: Stetiger Rhythmus, Fokus, Größe und Narbengewebe sind die Umstände, unter denen ein Deal-Programm weniger wehtut. Nichts in dem Papier sagt, dass Programme gewinnen.

Ein weiteres Resultat verdient einen eigenen Absatz, denn es läuft in McKinseys Richtung, und ohne es wäre dieser Essay unehrlich. In einer Robustheitsnotiz berichten Laamanen und Keil, dass über den vollen Horizont von 10–13 Jahren die häufigsten Käufer ihrer Stichprobe – 173 Firmen mit zehn oder mehr Deals – mediane Überrenditen von +12,6 % pro Jahr erzielten, gegen –3,9 % für die 438 Firmen mit vier bis neun Deals. Die Prosa des Papiers nennt die Differenz signifikant; die Analyse ist als nicht berichtet abgelegt, und nirgends im Artikel erscheint eine Teststatistik dafür. Es ist ein deskriptiver Split – und die stärkste Programm-Zahl, die die These vom häufigen Käufer hat: stärkere Stütze als alles in den Archetyp-Charts, unter demselben Vorbehalt, in den Worten der Autoren aus ihrer eigenen Limitations-Sektion: „we cannot claim causality“ – wir können keine Kausalität beanspruchen; einige der Effekte könnten auf überlegene frühere Performance der Käufer zurückgehen.

Abbildung 1 Die stärkste Zahl, und ihr Etikett



Laamanen & Keil (2008), Robustheitsnotiz – mediane Überrenditen pro Jahr über 10–13 Jahre, 611 serielle US-Käufer (1990–99); das Papier druckt die Gruppen als 'over 10' (173 Firmen) und '4–9' (438 Firmen), die die vollen 611 aufteilen. Die Prosa nennt den Split signifikant; die Analyse ist nicht berichtet, eine Teststatistik dafür erscheint nirgends im Artikel. Deskriptive Mediane; die Autoren beanspruchen keine Kausalität.

Der Titel dieses Stücks ist also präzise. Die Behauptung ist nicht an ihrem Test gescheitert. Sie hat ihn überlebt – 2012 neu lanciert, vier Jahre nach der Publikation des Tests, ohne ihn zu zitieren, und in keinem auffindbaren Dokument seither zu ihren eigenen Bedingungen geprüft.

Welche fünf Fachdisziplinen etablierten widersprüchliche Leistungskriterien?

Um Behauptung und Test herum sitzt eine akademische Aktenlage, die widersprüchlich wirkt, bis man bemerkt: Keine zwei Einträge messen dasselbe.

Auf Ebene der Deal-Ankündigung sehen häufige Käufer gut aus und werden schlechter: Firmen mit fünf oder mehr Akquisitionen binnen drei Jahren erzielten in der Stichprobe von Fuller, Netter und Stegemoller (1990–2000) eine positive durchschnittliche Ankündigungsrendite (+1,77 %, Fünf-Tages-Fenster) – in einem Übernahmemarkt, in dem nichtbörsennotierte Firmen 81 % aller Akquisitionen ausmachten –, und die Renditen sinken, je mehr Deals sich ansammeln. Auf Ebene des CEO, in Stichproben börsennotierter Ziele, läuft dieselbe Uhr negativ: Ankündigungsrenditen der späteren Deals eines Managers liegen bei Billett und Qian im Schnitt bei –1,51 %; deren langfristige Buy-and-Hold-Schätzungen sind statistisch nicht von null zu unterscheiden, und ihre Deutung ist nicht

Lernen, sondern Selbstattribution – dealende CEOs lesen frühes Glück als Können. Auf Ebene des organisationalen Lernens fanden Halebian und Finkelstein Performance, die mit früher Akquisitionserfahrung fällt und sich erst um den achten Deal erholt – ein U, das dasselbe Papier mit Fünf-Jahres-Erfahrung nicht reproduzieren konnte. Auf Ebene persistenten Könnens räumen Golubov, Yawson und Zhang die richtige Warnung selbst ein – Persistenz sei nicht notwendig Beleg für Können – und finden, dass Top-Käufer oben bleiben; speziell unter häufigen Käufern scheitert die Persistenz an ihrem univariaten Quintil-Test, sobald erwartete Renditen herausgerechnet sind, während das Top-Perzentil-Resultat ihn übersteht. Die Erfolge, schreiben die Autoren, „may or may not be due to skill“ – mögen auf Können beruhen oder nicht.

Und auf Ebene des Fachbuch-Widerspruchs hält das Buch *The M&A Failure Trap* von 2024 – das Nächste an einer publizierten Kritik, das die Behauptung hat – den kleinsten Effekt der gesamten Aktenlage: Auf der Zehn-Faktoren-Scorecard der Autoren, destilliert aus ihrem 43-Variablen-Modell, fügt jede Akquisition des Käufers aus den fünf Vorjahren dem Erfolgswert eines Deals +0,04 hinzu, der geringste aller Faktoren, mit kippendem Vorzeichen in Biotech und Öl. Ihre Präambel benennt das Problem der fehlenden Variable klar – starke operative Substanz kann Deal-Programm und Renditen zugleich treiben –, was der Punkt von Fußnote 9 ist, von der anderen Seite formuliert. Zwei Vorsichten zum selben Buch: Sein „Failure“-Konstrukt verlangt drei Hürden zugleich, gewertet nur am größten Deal jedes Jahres – ein Design, das die Misserfolgs-Schlagzeile aufbläht und genau die kleinen Deals verwirft, die häufige Käufer machen. Und die Statistik, die unter dem Namen dieses Buches durch Praktiker-Zusammenfassungen läuft – serielle Käufer mit zehn und mehr Deals erfolgreich zu 54 %, Ersttäter zu 23 % – steht nicht im Buch. Die 54 % sind eine Scorecard-Illustration; die 23 % gehören zu Deals vor Restrukturierungen. Die Recherche-Pipeline dieses Essays hat die Zahl selbst aus einer Zusammenfassung importiert, bevor die Vollektüre sie fing. So funktioniert diese Literatur: Die Zahlen, die kursieren, sind nicht verlässlich die Zahlen in den Dokumenten.

Ankündigungsfenster auf Deal-Ebene, Deal-Reihenfolge auf CEO-Ebene, Buy-and-Hold-Renditen auf Firmenebene, Dekaden-TSR auf Archetyp-Ebene, ein konjunktives Erfolgskonstrukt am größten Deal. Fünf Maßstäbe. Der Playbook-Essay argumentierte, dass ein gemessenes und ein vermarktetes Playbook verschiedene Objekte sind; hier liegt das Problem davor – Behauptung und vermeintliche Widerlegungen sind nicht einmal auf derselben Achse gemessen. Nichts an dieser Aktenlage lässt sich addieren, und dieser Essay hat nichts addiert.

Wie sollten Führungskräfte externe Benchmarks vor deren Einführung prüfen?

Das ehrliche Urteil über programmatische M&A hat drei Worte: bedingt, maßstabsabhängig, ungetestet wie behauptet. Das ist kein Urteil gegen Deal-Programme. Die stärkste Zahl im nächstgelegenen akademischen Test spricht für sie – als ungetesteter Median. Was die Aktenlage trägt, ist keine Strategieempfehlung, sondern ein Satz Fragen – und es sind bessere Fragen als die des Charts.

Tabelle 2 Die Behauptung, neben ihrem Test

Dimension	Die Behauptung, ein McKinsey-Korpus von 2012 bis 2023	Der Test (Laamanen & Keil, 2008)
Gestellte Frage	Welches realisierte Deal-Muster hatte ex post den besten Excess TSR, nach Archetyp?	Sagen Tempo, Rhythmus und Streuung eines Programms Überrenditen voraus?
Population	Top 1.000 → „Global 2.000“, Fenster je Auffrischung verschoben	611 US-Käufer mit 4+ Deals, 5.518 Deals, 1990–99
Befund	Programmatische Käufer schneiden besser ab (rund 2 %/Jahr Excess TSR 2021; 3,9 % vs. 2,9 % 2023) – unter einer Definition, die sich je Fassung ändert (Tabelle 1)	Tempo schadet; Rhythmus-Variabilität schadet; Erfahrung, Größe und Fokus kaufen Toleranz; R^2 0,02–0,04. Dekaden-Mediane favorisieren häufige Käufer (+12,6 % vs. –3,9 %/Jahr) – deskriptiv, kein Test berichtet
Eigener Vorbehalt, wörtlich	„[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“
Früheres Resultat, gleiche Daten	Apr 2011, gleiches Haus, gleiche Population: Größen- und Frequenzmuster „widely distributed and overlapping“	–
Was es erlaubt	Eine Hypothese, die auf den Bedingungen des eigenen Programms zu testen wäre	Bedingungen, die vor und während jedes Programms zu prüfen sind

Eigene Zusammenstellung: Rehm, Uhlaner & West (2012) und Auffrischungen bis Daume, Lian & McCurdy (2023), gegen Laamanen & Keil (2008) und Cottin, Rehm & Uhlaner (2011). Quelle und Bedingungen jeder Zelle stehen im Text. Eine Lesehilfe, kein Befund einer einzelnen Quelle.

Die Bedingungen des Tests übersetzen sich direkt in Diligence-Fragen für das eigene Programm, und keine verlangt einen Kauf. Ist Ihr Deal-Rhythmus stetig, oder kommt er in Schüben? Das Tempo ist der robusteste Schaden des Tests; Unstetigkeit schadet obendrein – und anders als das Tempo ist sie der Schaden, den Erfahrung abkaufen kann. Liegt das Programm in Geschäften, die Sie kennen, oder streut es? Streuung verschärft den Tempo-Malus im Schnitt in den Daten des Tests. Haben Sie die Größe und das Narbengewebe? Größe kaufte Toleranz für alles; Erfahrung nur für Unstetigkeit – und ihr Direkteffekt war negativ, was die bequeme Idee beerdigen sollte, Deals machen lehre Deals machen von selbst. Und eine Frage, auf der die Aktenlage vor all dem besteht: Unter welcher Definition aus Tabelle 1 zählte Ihr Programm überhaupt als programmatisch – und hätte es im April 2011 gezählt, als dasselbe Haus in derselben Population nichts zu zählen fand?

Wo dieser Essay endet: Die Bedingungen des Tests sind US-Evidenz der 1990er aus sieben Sektoren – sie belegen, dass Mustereffekte existieren und wohin sie dort liefen, und sie sagen nichts über Ihr Programm voraus. Die Archetyp-Charts, die CEO-Ebene und die Ankündigungsfenster können einander nicht widerlegen, und dieser Essay hat keines benutzt, um ein anderes abzutun; Evidenz statt Anekdote ist hier die stehende Regel, und sie schneidet gegen bequeme Widerlegungen so hart wie gegen bequeme Behauptungen. Keine Empfehlung, zu kaufen oder zu lassen, verlässt dieses Stück. Was es verlässt, ist die Seite, die Ihrem nächsten Deal-Deck fehlt: die Behauptung, die erste Antwort ihres eigenen Hauses, ihre eigenen Fußnoten und ihr Test – nebeneinander, datiert, mit jedem Vorbehalt in den Worten seiner Autoren.

Wo liegen die empirischen Grenzen der Benchmark-Überprüfung?

Grenze. Die ursprüngliche Behauptung verfehlt das geprüfte Kriterium, nicht jedes Akquisitionsergebnis ist damit null. Trenne die fünf Kriterien und teste eine aktuelle Stichprobe.

Evidenzbasis. Der analytische Rahmen stützt sich außerdem auf diese zusätzlichen Quellen: Billett and Qian 2008; Fuller et al. 2002; Golubov et al. 2015; Halebian and Finkelstein 1999; Lev und Gu (2024). Die Links identifizieren die konkreten Werke; sie stützen die hier diskutierten Mechanismen und Geltungsgrenzen, nicht jede einzelne Aussage isoliert.

- 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, 9. Mai). 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, 1. April). 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, 24. August). 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, 13. Oktober). 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, 25. März). 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, 1. Januar). 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, 12. Juli). 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

ZITIERWEISE

Isoglu, S. (2026, 14. August). Programmatische M&A: die Behauptung, die ihren Test überlebte. Sinan Isoglu.
<https://isoglu.com/de/insights/journal/die-behauptung-die-ihren-test-ueberlebte/>

WEITERLESEN

Wachstum, das sich potenziert

Ein Kundenmodell kann nicht jedes Ergebnis prognostizieren

<https://isoglu.com/de/insights/journal/ein-kundenmodell-kann-nicht-jedes-ergebnis-prognostizieren/>

Wachstum, das sich potenziert

Die Megadeal-Welle trifft auf den Größeneffekt

<https://isoglu.com/de/insights/journal/die-megadeal-welle-trifft-auf-den-groesseneffekt/>

Wachstum, das sich potenziert

Customer Lifetime Value ist eine Prognose, keine Tatsache

<https://isoglu.com/de/insights/journal/customer-lifetime-value-ist-eine-prognose-keine-tatsache/>



Sinan Isoglu, MBA (Quantic)

Führungskraft für Umsatzwachstum, Dozent und Doktorand