



CMO Value Creation on Digital Platforms under Distorted Performance Feedback

Luis Maximilian Kade

Dissertation written under the supervision of

Professor Peter V. Rajsingh, PhD

Dissertation submitted in partial fulfillment of requirements for the MSc in
Management with Specialization in Strategic Marketing, at the
Universidade Católica Portuguesa, 28.05.2026.

Abstract

Digital platforms have become the main infrastructure through which Chief Marketing Officers (CMOs) allocate growth budgets and justify marketing expenditure. Yet, platform performance signals are produced inside systems whose attribution rules, reporting logic, and optimization incentives may diverge from true causal impact. This dissertation examines how CMOs can measure and defend value creation under distorted platform feedback.

The study develops and tests a mixed-methods explanation centered on Performance Learning Distortion (PLD), Perceived CMO Value Creation (PVC), and Growth Investment Capability (GIC). An exploratory sequential design combines 17 semi-structured expert interviews with a 2×2 experimental vignette study. The qualitative phase identifies Evidence Integration as the focal microfoundation. The quantitative phase tests whether stronger evidence reduces distorted learning and changes value and budget judgments.

The findings support the core mechanism. Evidence Triangulation reduced PLD and lowered perceived value claims for the focal investment. Serial mediation showed that evidence quality shaped budget decisions through lower PLD, lower PVC, and more restrained allocation. Lower PVC after Evidence Triangulation should therefore be interpreted as better calibration of inflated value claims, not as lower objective CMO value creation. The hypothesized moderating role of GIC was not supported. Capability appears to operate upstream by shaping evidence architecture, validation routines, and decision thresholds.

The dissertation contributes to marketing accountability theory by explaining how platform-governed feedback can distort executive learning. It identifies Evidence Integration as a microfoundation that protects decision quality under incomplete or biased platform evidence.

Keywords: Platform governance, performance learning distortion, CMO value creation, evidence integration, dynamic capabilities, digital advertising, incrementality measurement

Title: CMO Value Creation on Digital Platforms under Distorted Performance Feedback

Author: Luis Maximilian Kade

Resumo

As plataformas digitais tornaram-se a principal infraestrutura através da qual os CMOs alocam orçamentos de crescimento e justificam investimentos em marketing. No entanto, os sinais de desempenho vêm de sistemas cujas regras de atribuição, lógica de reporte e incentivos de otimização podem divergir do impacto causal real. Esta dissertação analisa como CMOs podem medir e defender o valor diante de feedback distorcido proveniente dessas plataformas.

O estudo desenvolve e testa uma explicação de métodos mistos centrada em Performance Learning Distortion (PLD), Perceived CMO Value Creation (PVC) e Growth Investment Capability (GIC). Um desenho sequencial combina 17 entrevistas semiestruturadas com um estudo experimental de vinhetas 2×2. A fase qualitativa identifica Evidence Integration como microfundamento central. A fase quantitativa testa se evidências mais robustas reduzem a aprendizagem distorcida e alteram julgamentos de valor e decisões orçamentais.

Os resultados sustentam o mecanismo. A triangulação das evidências reduziu o PLD e as percepções de valor inflacionadas. A mediação mostrou que a qualidade da evidência moldou decisões orçamentais por meio de menor PLD, menor PVC e uma alocação mais contida. A relação negativa entre Evidence Triangulation e PVC reflete uma calibração melhor, não uma menor criação de valor. A GIC não apresentou efeito moderador. A capacidade parece atuar a montante, moldando a arquitetura de evidência, as rotinas de validação e os limiares de decisão.

A dissertação contribui para a teoria de marketing ao identificar a Evidence Integration como um microfundamento que protege as decisões diante de evidências incompletas ou enviesadas.

Palavras-chave: Governança de Plataformas, Distorção da Aprendizagem de Desempenho (PLD), Criação de Valor do CMO (PVC), Integração de Evidência, Capacidades Dinâmicas, Publicidade Digital, Mensuração de Incrementalidade

Título: Criação de Valor do CMO em Plataformas Digitais sob Feedback de Desempenho Distorcido

Autor: Luis Maximilian Kade

Acknowledgments

Throughout this journey, I have been fortunate to receive support from many individuals, and I would like to express my gratitude to them. First and foremost, I am deeply thankful to my supervisor, Prof. Peter V. Rajsingh, for his guidance, support, and responsiveness, which were fundamental to the development of this dissertation.

I am also sincerely grateful to the 17 industry experts who generously shared their time and insights with me, enriching this dissertation with invaluable contributions. Likewise, I extend my thanks to everyone who responded to my survey and to those who facilitated warm introductions and referrals, enabling me to expand my network.

A special thanks goes to my family and partner, whose support has been a steady source of strength throughout the writing of this dissertation. I am equally grateful to my friends, who have enriched my life during this journey, challenged my thinking, and kept me motivated and inspired.

Lastly, I would like to thank everyone who has influenced my strategic thinking and scientific approach to work over the last few years.

Luis

Table of Contents

Abstract.....	I
Resumo	II
Acknowledgments.....	III
Table of Contents	IV
List of Figures	VI
List of Tables.....	VII
List of Abbreviations.....	IX
1. Introduction	1
2. Literature Review.....	4
2.1 Scope and review focus.....	4
2.2 Platform governance and information control.....	4
2.3 Platform dependence and platform opacity	5
2.4 Incrementality and attribution in digital advertising.....	6
2.5 Organizational learning under distorted performance feedback.....	7
2.6 Marketing accountability, legitimacy, and perceived marketing value.....	9
2.7 Capability and governance responses that protect decision quality	10
2.8 Synthesis and conceptual implications for this dissertation	11
2.9 Hypotheses Development.....	12
3. Methodology	15
3.1 Research Design	15
3.2 Data Collection	17
3.2.1 Secondary Data Collection	17
3.2.2 Primary Data Collection – Qualitative.....	17
3.2.3 Primary Data Collection – Quantitative.....	18
3.3 Data Analysis.....	20
3.4 Research Quality & Ethics.....	20
4. Results	22
4.1 Qualitative Results.....	22
4.2 Quantitative Results.....	26
4.2.1 Descriptive Statistics and Correlations	26

4.2.2 Manipulation Checks	27
4.2.3 Hypothesis Tests	28
4.2.4 Behavioral Mechanism Test: Serial Mediation to Budget Change.....	31
4.2.5 Additional Behavioral Analysis.....	32
4.2.6 Robustness Checks.....	32
4.2.7 Summary of Hypothesis Tests	33
4.3 Final Synthesis – Triangulation	33
5. Conclusion.....	35
5.1 Discussion of Main Findings	35
5.2 Theoretical Implications.....	37
5.3 Practical Implications.....	38
5.4 Limitations.....	39
5.5 Future Research	40
5.6 Final Conclusion	42
Reference List	A
Appendix	H
Appendix A. Semi-Structured Interview Guide.....	H
Appendix B. Qualitative Coding Structure	K
Appendix C. Quantitative Survey Instrument and Experimental Vignette.....	T
Appendix D. Scale Construction.....	Z
Appendix E. Coding Notes	AA
Appendix F. Data Cleaning and Variable Construction	AA
Appendix G. Supplementary Quantitative Analyses.....	FF

List of Figures

Figure 1 Conceptual model	14
Figure 2 Research design and methodological sequence	16

List of Tables

Table 1	Expert overview.....	22
Table 2	Qualitative findings and construct links	25
Table 3	Condition means by experimental cell	26
Table 4	Descriptive statistics and intercorrelations	27
Table 5	Manipulation checks.....	27
Table 6	Direct and indirect path tests for H1–H4.....	29
Table 7	Moderated mediation analysis with PVC as outcome	30
Table 8	Serial mediation from ET to Budget Change via PLD and PVC	31
Table 9	Summary of hypothesis tests	33
Table 10	Semi-structured interview guide.....	H
Table 11	Qualitative coding structure and construct derivation.....	L
Table 12	Selection logic for evidence integration as focal microfoundation	Q
Table 13	Coded expert interview matrix	R
Table 14	Survey structure and objectives	T
Table 15	Scale construction summary	Z
Table 16	Data-cleaning funnel	AA
Table 17	Experimental condition coding.....	BB
Table 18	Construct and scale construction	CC
Table 19	PLD construction.....	CC
Table 20	Budget change coding	DD
Table 21	Additional analytical variables	EE
Table 22	Statistical analysis workflow	EE
Table 23	Randomization balance checks.....	FF
Table 24	Reliability checks	GG
Table 25	Manipulation checks.....	GG

Table 26	Planned contrast results for PLD	HH
Table 27	Robustness and planned contrast results	HH
Table 28	Robustness check excluding influential cases 30 and 183	II
Table 29	Behavioral validation model.....	JJ

List of Abbreviations

AI	Artificial Intelligence
ANOVA	Analysis of Variance
API	Application Programming Interface
B2B	Business-to-Business
CCO	Chief Commercial Officer
CEO	Chief Executive Officer
CFO	Chief Financial Officer
CI	Confidence Interval
CMO	Chief Marketing Officer
CRM	Customer Relationship Management
CRO	Chief Revenue Officer
CTR	Click-Through Rate
CTO	Chief Technology Officer
D2C	Direct-to-Consumer
DSP	Demand-Side Platform
EI	Evidence Integration
ET	Evidence Triangulation
FMCG	Fast-Moving Consumer Goods
GIC	Growth Investment Capability
IAB	Interactive Advertising Bureau
KPI	Key Performance Indicator
LTV	Lifetime Value
OLS	Ordinary Least Squares
PLD	Performance Learning Distortion
PROCESS	PROCESS Macro
PVC	Perceived CMO Value Creation
QSR	Quick-Service Restaurant
ROAS	Return on Ad Spend
SDK	Software Development Kit
SEO	Search Engine Optimization
SD	Standard Deviation
SE	Standard Error

SPSS	Statistical Package for the Social Sciences
TV	Television
VIF	Variance Inflation Factor-

1. Introduction

“Half my advertising spend is wasted; the trouble is, I don’t know which half.” This quote, often attributed to John Wanamaker, captures a persistent problem in marketing management. Digital platforms have made advertising faster, more scalable, and more measurable. But they have not made the causal impact self-evident. The same systems that help firms reach customers also shape the evidence used to judge whether marketing spend creates value. In Europe, digital advertising now accounts for 67.2% of total ad spend, directing a large share of growth investment into platform-mediated channels (IAB Europe, 2025). Europe’s advertising market is expected to reach €188.5bn in 2025, raising the stakes for how growth budgets are governed and justified (McKinsey & Company, 2026). Chief Marketing Officers (CMOs) can reallocate substantial budgets quickly, yet the evidence used to justify these reallocations increasingly comes from platform-controlled dashboards. Platforms enable scalable growth, but they also shape the rules of access, measurement, and attribution that define what appears to work.

This problem matters because platform feedback can diverge from causal business impact. Many firms rely on a small set of dominant platforms, while transparency into algorithms, data access, and attribution remains limited (Cutolo & Kenney, 2021; Gordon et al., 2019).

These conditions constrain managerial discretion and make it difficult to verify whether observed performance reflects true incrementality, defined as the impact that would not have occurred absent the spend (Cutolo & Kenney, 2021; Gordon et al., 2019). When marketing evidence is partly produced outside the firm, organizational learning becomes central to value creation (Argote, 2012; Zollo & Winter, 2002).

The gap between reported performance and causal impact has consequences beyond measurement. It can distort how CMOs learn from performance feedback and how they assess the value of their decisions. When platform-reported metrics are treated as proof of causal impact, managers may update their beliefs too optimistically and reallocate budgets based on misleading signals. The strategic problem, therefore, is whether value can be credibly measured and defended under imperfect feedback.

The issue also concerns how firms allocate and defend growth resources under uncertainty. Under market turbulence, performance differences reflect not only industry structure but also firm-level capabilities that shape how resources are allocated and reconfigured over time (Barney et al., 2021; Barreto, 2010; Teece, 2007). In platform-mediated marketing, CMOs increasingly serve as strategic resource allocators (McKinsey & Company, 2026). They make

recurring growth-investment decisions in contexts where key customer interfaces and decision signals are partially controlled by external parties (Gawer, 2014; Jacobides et al., 2018; Wichmann et al., 2022).

This dissertation focuses on three constructs. Perceived CMO Value Creation (PVC) refers to the CMO's assessment that platform-mediated marketing spend delivers incremental business impact that justifies its cost. It also reflects whether this value claim remains defensible under scrutiny. PVC is treated as a perceptual outcome because independent causal verification is often constrained in platform environments.

Performance Learning Distortion (PLD) captures systematic distortion in executive learning. Conceptually, PLD refers to the gap between platform-reported signals and the CMO's best causal assessment of incrementality. In the experiment, PLD is measured as belief updating following exposure to platform-mediated evidence, rather than as a verified gap relative to true incrementality. This distinction matters because repeated reallocations can overweight within-platform metrics and compound misallocation across cycles (Gordon et al., 2019; Kohavi et al., 2009; Lewis & Rao, 2015).

Growth Investment Capability (GIC) refers to an organization's ability to govern, evaluate, and reconfigure growth investments. It entails explicit evidence standards, validation routines, decision rights, and disciplined reallocation practices beyond platform key performance indicators (KPIs) (Day, 2011; Homburg & Wielgos, 2022; Kalaignanam et al., 2021). The study tests whether a stronger GIC makes PLD less likely to translate into inflated value claims by strengthening evidence standards and validation capacity.

The gap lies in how these streams connect at the CMO level. Platform governance research explains how platforms create power asymmetries and informational constraints (Cutolo & Kenney, 2021; Gawer, 2014; Jacobides et al., 2018). Marketing accountability research shows that attribution and incrementality remain difficult to establish in digital advertising (Gordon et al., 2019; Lewis & Rao, 2015). Organizational learning research shows that noisy feedback can distort belief updating and resource allocation (Greve, 2003; Levitt & March, 1988; March & Olsen, 1975). Yet these streams have not been integrated into an account of how platform-governed feedback shapes CMO learning within recurring growth-allocation cycles.

This study addresses that gap by offering a dynamic-capabilities explanation grounded in microfoundations. It identifies Evidence Integration as the focal routine through which managers synthesize platform data, internal outcomes, experimental evidence, customer quality

signals, and cross-functional interpretation before making allocation decisions. Evidence Integration should reduce PLD by preventing platform-reported evidence from being treated as sufficient to establish causal impact. The study tests this explanation using an exploratory sequential mixed-methods design, combining 17 semi-structured interviews with senior marketing decision-makers and a 2×2 experimental vignette study ($n = 209$).

The research project makes three contributions. First, it introduces PLD as a mechanism linking platform-governed feedback to executive learning and budget allocation. Second, it specifies GIC as a capability architecture that shapes evidence standards, validation routines, and disciplined reallocation under uncertainty. Third, it identifies Evidence Integration as an empirically grounded microfoundation that protects decision quality in platform-mediated marketing settings.

These contributions are organized around the following research question:

Research Question: How can CMOs measure and defend value creation on digital platforms under distorted performance feedback?

2. Literature Review

2.1 Scope and review focus

Platform-mediated marketing creates a measurement problem. Platforms enable customer access, but they also shape the data, attribution rules, and reporting interfaces used to evaluate marketing performance. As a result, reported performance can diverge from true incremental impact, making it difficult for CMOs to judge whether platform-mediated spend creates value (Gordon et al., 2019; Lewis & Rao, 2015).

Three distinct research streams explain this problem. Platform governance explains how platforms shape access, dependence, and observability. Digital advertising measurement explains why attribution outputs and common KPIs can diverge from causal impact. Organizational learning and marketing accountability explain how ambiguous feedback is accepted, challenged, or translated into resource allocation.

These streams support three constructs. Performance Learning Distortion (PLD) captures distorted belief updating when platform-governed signals are overweighted (Greve, 2003; March & Olsen, 1975). Perceived CMO Value Creation (PVC) captures the CMO's assessment that platform-mediated spend creates incremental value worth its cost and remains defensible under scrutiny (Hanssens & Pauwels, 2016; Rust et al., 2004). Growth Investment Capability (GIC) captures the governance and measurement routines that protect decision quality under imperfect feedback (Day, 2011; Teece, 2007).

2.2 Platform governance and information control

Platform governance encompasses the technical and institutional arrangements that shape participation, access, and coordination within a platform ecosystem. For marketing decision-making, governance matters because it shapes both action and evidence. It defines what firms can do on the platform and what they can observe about the outcomes of those actions.

Platforms enable market participation, but they also constrain independent observation and validation. Access and data use are governed by boundary resources such as application programming interfaces (APIs), software development kits (SDKs), reporting interfaces, and experimentation tools (Eaton et al., 2015; Ghazawneh & Henfridsson, 2013; Helmond et al., 2021). APIs, for instance, translate platform policy into the technical conditions for access and permissible uses of information. They therefore shape both managerial action and the evidence

available for audit (Gawer, 2021; Helmond et al., 2021; Schmeiss et al., 2019; Staub et al., 2022).

These constraints extend to reporting standards, attribution interfaces, and experimentation tools. They are not purely technical. They also reflect platform incentives to retain control over measurement, optimization, and advertiser spending (Cutolo & Kenney, 2021; Gawer, 2021; Jacobides et al., 2018). As a result, firms often evaluate marketing performance using platform dashboards, attribution reports, optimization tools, and experimentation interfaces they do not fully control.

Design choices, privacy regulations, and tracking restrictions further constrain what firms can infer. They reduce user-level observability and make independent assessment of causal impact more difficult (Geradin & Katsifis, 2020; Goldfarb & Tucker, 2011). Measuring these constructs remains difficult for two reasons. Platform governance limits access to independent evidence, and causal inference in advertising is structurally challenging (Gordon et al., 2019).

This creates a link to PLD. When platforms control both action and evidence, managers may learn from performance signals they cannot fully audit. Over time, these signals can shape belief updating and resource allocation even when they diverge from true incremental impact (Cutolo & Kenney, 2021; Greve, 2003; March & Olsen, 1975).

2.3 Platform dependence and platform opacity

Platform dependence constrains substitution. Firms rely on a small set of dominant platforms because high-quality alternatives are scarce and switching costs are high. Platform opacity constrains interpretation. It limits observability, transparency, and access to performance information. This distinction is important: dependence limits where firms can act; opacity constrains what firms can know.

Platforms create lock-in by controlling key nodes and standards, including identity, tracking, APIs, and campaign tools. This raises switching costs and narrows the set of viable alternatives (Micova & Jacques, 2020; Nieborg & Poell, 2025; van der Vlist & Helmond, 2021). Dependence is reinforced by data advantages and intermediation structures, such as platform-controlled access to audiences, advertising inventory, and performance information. Platform-native targeting, optimization, and reporting systems are difficult to replicate externally, while market concentration shifts bargaining power and the distribution of surplus in advertising

markets (Decarolis & Rovigatti, 2021). Spending tends to move to other incumbents only after significant market shocks (Donati & Fong, 2025).

Opacity limits participants' ability to understand why outcomes occurred and how to align actions with platform delivery, ranking, and optimization criteria. Restricted measurement data and limited reporting standards reduce the granularity, comparability, and auditability of performance signals (Nieborg & Poell, 2025; Pang et al., 2022; van der Vlist & Helmond, 2021). This increases reliance on platform-native measurement interfaces (Benlian et al., 2015; Decarolis & Rovigatti, 2021; Rahman, 2021).

Together, dependence and opacity create the conditions in which PLD becomes likely. Dependence keeps firms within platform-mediated channels, while opacity limits their ability to validate the feedback those channels provide. This combination links platform governance to distorted learning under ambiguous feedback (Cutolo & Kenney, 2021; March & Olsen, 1975).

2.4 Incrementality and attribution in digital advertising

Performance and attribution indicators often diverge from causal effects, or incrementality. Exposure is rarely random in digital advertising. Platform delivery and auction systems adapt dynamically to user behavior and campaign conditions, while attribution conventions compress complex response dynamics into simplified metrics. As a result, what appears effective in a platform dashboard or an observational attribution model may not reflect true causal impact.

The literature identifies four mechanisms behind this divergence: selective exposure, delivery confounding, interference, and metric integrity.

Selective exposure occurs when targeting and optimization make exposed users systematically different from non-exposed users. Gordon et al. (2022) show that observational estimators applied to the same campaigns can produce biased estimates, including cases where estimates diverge substantially from experimental effects. Barajas et al. (2016) similarly find that many attributed conversions may have occurred without marketing exposure. These findings show that attribution errors are not edge cases. They are predictable when observed variables only partially capture a user's likelihood of converting.

Delivery confounding arises when platform systems allocate ads to optimized audiences during the experiment itself. Braun and Schwartz (2024) show that online experiments can be

confounded by divergent delivery, where measured differences reflect both ad effects and differences in exposure. This weakens what standard A/B tests can infer about causal response because platform delivery algorithms can undermine effective randomization. The problem is especially relevant when platforms optimize campaigns while firms are trying to test them.

Interference creates a further source of bias. In advertising marketplaces, one firm's experiment can affect the exposure, prices, and outcomes of other firms. Waisman et al. (2024) show that parallel experimentation and competitive interference can skew data when multiple advertisers adjust bids and budgets simultaneously. Even well-designed experiments can therefore produce misleading inferences if they ignore marketplace interactions.

Metric integrity and interpretation form the fourth mechanism. Edelman (2014) documents false metrics and shows how advertising effectiveness can be overstated when recorded activity is treated as incremental impact. Chan et al. (2011) examine incremental clicks and show that observed click activity does not automatically reveal counterfactual outcomes. Dong et al. (2023) further show that channels can produce different relationships between intermediate KPIs and downstream outcomes: artificial intelligence (AI)-recommended channels generated higher click-through rates but lower conversion rates than subscription channels. Common KPIs such as Click-Through Rate (CTR) can therefore misrepresent business impact when they are interpreted without downstream validation against sales, retention, or customer quality.

These mechanisms point towards the same conclusion. Attribution outputs and platform KPIs can guide campaign execution, but they are weak evidence of incrementality when used in isolation (Gordon et al., 2019; Lewis & Rao, 2015). This matters for recurring budget decisions. If CMOs treat platform-reported performance as causal evidence, uncertain attribution can become distorted learning. Incrementality uncertainty is therefore the measurement basis for PLD.

2.5 Organizational learning under distorted performance feedback

Organizational learning research shows that feedback does not automatically lead to accurate adaptation. Learning depends on how decision-makers interpret signals that are noisy, biased, delayed, or ambiguous. In such contexts, organizations face uncertainty about what observed outcomes mean and why they occurred (Levitt & March, 1988; March & Olsen, 1975). This logic builds on the behavioral theory of the firm, which treats organizational decisions as shaped

by bounded rationality, routines, search, and aspiration-based adjustment (Cyert & March, 1963; Greve, 2003).

Performance feedback can also be distorted by cognition and prior beliefs. Schumacher et al. (2020) show that CEO overconfidence is linked to more optimistic interpretations of performance feedback and weaker behavioral adjustments. Individual cognition determines whether feedback leads to corrective action. Biased interpretation can reduce sensitivity to negative feedback and weaken responses to performance shortfalls (Schumacher et al., 2020). Organizations also interpret new feedback through routines and beliefs shaped by past experience, even as environments change (Levitt & March, 1988).

Organizations adjust to feedback by comparing performance data with reference points (Greve, 2003). Albert (2025) reports that managers who underestimate feedback noise revise their beliefs more often. Frequent updating can produce early errors, but it may also enable faster self-correction in complex settings. This creates a trade-off between speed and accuracy. Frequent updating can make beliefs and allocation decisions more volatile in the short run, while reliable feedback can improve adaptation over time (Albert, 2025).

Decision routines also shape learning quality. Beil et al. (2025) find that decision-makers often struggle with dynamic allocation problems, especially when cost structures change over time. Higher performers improve by decomposing complex problems into simpler subproblems. Sharing best practices can also improve performance, indicating that routines and heuristics can support allocation quality over repeated cycles (Beil et al., 2025). Feedback is therefore filtered through both cognition and decision routines, not only through the informational properties of the signal.

For CMO budget cycles, this matters because platform dashboards are not neutral feedback objects. They provide frequent, visible signals, but those signals are often causally incomplete. If managers update beliefs from platform feedback without validation routines, small interpretive errors can compound across allocation cycles. PLD extends organizational learning logic to platform-mediated growth allocation, where feedback is visible, repeated, and difficult to causally verify. In this dissertation, PLD is therefore treated as directional belief updating under potentially inflated platform feedback, not as a verified deviation from an external incrementality benchmark.

2.6 Marketing accountability, legitimacy, and perceived marketing value

Marketing accountability depends on credible measures that link marketing activity to enterprise outcomes such as shareholder value, profit, and growth (Rogers, 2017; Rust et al., 2004). Performance measurement shapes how marketing value claims are assessed, defended, and challenged in senior decision-making contexts (O'Sullivan & Butler, 2010; Rust et al., 2004). Marketing value is therefore evaluated not only by observed outcomes but also by whether the claimed contribution can withstand budgetary and managerial scrutiny. Measurement reliability is central to closing marketing's credibility gap (Rust et al., 2004). O'Sullivan et al. (2009) show that measurement capabilities can serve as a strategic asset rather than an administrative practice. O'Sullivan and Butler (2010) further show that perceived marketing accountability shapes how marketing is evaluated within top management.

Marketing performance assessment must connect marketing activity to customer, strategic, and financial outcomes (Katsikeas et al., 2016; Morgan et al., 2021). In digital contexts, web analytics research warns against over-reliance on vanity metrics and emphasizes measures that better reflect customer value and business impact (Järvinen & Karjaluoto, 2015; Vollrath & Villegas, 2022). Marketing evaluations, therefore, need to link investments to enterprise value, profit, and growth rather than to platform-visible activity alone (Rogers, 2017).

The CMO role makes this accountability problem more salient. CMOs are expected to link customer-facing growth activities to firm-level performance, yet their strategic influence hinges on whether top management views marketing as financially credible. Research on CMO influence shows that CMOs gain strategic standing when they contribute to firm strategy and demonstrate value beyond communication execution (Homburg et al., 2015). Evidence on CMO tenure further suggests that weak perceived contribution and performance pressure can make the role vulnerable when marketing impact is difficult to defend (Germann et al., 2015).

Platform-mediated advertising creates a specific accountability tension. Platform-controlled reporting makes KPI narratives easy to generate and communicate, but the evidence needed for independent verification is partly external to the firm. Senior stakeholders may therefore treat dashboard metrics as useful while still questioning whether observed performance reflects incremental impact. This makes PVC a relevant outcome: in platform settings, CMO value creation depends not only on performance claims but also on whether those claims can be defended with credible evidence (Hanssens & Pauwels, 2016; Rust et al., 2004).

2.7 Capability and governance responses that protect decision quality

Resource-Based Theory holds that capabilities matter because they shape how firms allocate scarce resources under uncertainty (Barney et al., 2021). Over time, these allocation differences can compound into durable performance differences (Barney et al., 2021). Dynamic capabilities are especially relevant in changing environments because they explain how firms adapt their resource base as markets, technologies, and competitive conditions shift (Barreto, 2010; Teece et al., 1997; Teece, 2007). They should be specified as distinct, measurable constructs rather than treated as residual explanations for performance differences (Barreto, 2010). Marketing capabilities research applies this logic to market sensing, digital capability, and adaptive resource deployment, showing that firms differ in how they translate market information into performance-relevant action (Day, 2011; Homburg & Wielgos, 2022; Kalaignanam et al., 2021).

In platform-mediated marketing, decision quality depends not only on available data but also on the routines and governance arrangements through which evidence is generated, evaluated, challenged, and translated into action. Capabilities can therefore be specified as observable decision-process mechanisms. Three capability families are especially relevant for learning quality and defensible value claims: experimentation and measurement infrastructure, governance and decision rights, and disciplined reallocation routines. These dimensions translate dynamic-capability logic into platform-mediated marketing judgment. Evidence infrastructure supports sensing by helping firms detect whether platform signals reflect business impact. Governance and decision rights support seizing by defining when evidence is strong enough to scale, pause, or stop spend. Disciplined reallocation routines support reconfiguring by carrying learning across budget cycles (Teece, 2007).

Growth Investment Capability (GIC) is an organization's ability to govern, evaluate, and reconfigure growth investments under uncertainty. GIC institutionalizes evidence standards across recurring allocation cycles. It comprises three operational dimensions: evidence infrastructure, governance and decision rights, and disciplined reallocation cadence. Evidence infrastructure encompasses experiments and measurement approaches that generate decision-relevant evidence beyond platform-native KPIs. Governance and decision rights define who evaluates evidence, who challenges weak evidence, and when spending is scaled, paused, or stopped. Disciplined reallocation cadence refers to regular budget-review cycles in which teams document why spending is scaled, paused, or stopped, so that learning carries across decisions.

This framing distinguishes GIC from analytics. Analytics concern the production of information. GIC concerns the governance of how evidence is used, challenged, retained, and translated into allocation decisions. This aligns with calls to specify capabilities through constituent processes and observable operations (Barreto, 2010). Prior work on dynamic-capability measurement and experimentation similarly emphasizes specifying capabilities through observable routines and evidence-generating practices (Hilliard & Goldstein, 2019; Liu et al., 2020).

The dynamic capabilities literature also motivates examining microfoundations, meaning the concrete routines through which capabilities operate (Teece, 2007). Decision routines can protect learning quality by shaping how organizations sense signals, validate interpretations, and reconfigure allocations (Teece, 2007). In this dissertation, Evidence Integration is treated as such a routine. It specifies how managers combine platform data with broader evidence before making allocation decisions, thereby linking dynamic-capability theory to the problem of distorted platform feedback.

2.8 Synthesis and conceptual implications for this dissertation

PLD translates the measurement problem into a learning problem at the budgeting cadence. PLD occurs when platform-governed performance signals diverge from the best available causal assessment of incrementality and are overweighted in reallocation decisions. This treats the distortion as a recurring learning outcome that can compound across cycles, rather than a one-off measurement error. The construct builds on evidence that causal measurement in platform environments is difficult, that performance signals can lack validity, and that belief updating can be biased under ambiguity and noise (Albert, 2025; Braun & Schwartz, 2024; Gordon et al., 2022; March & Olsen, 1975; Schumacher et al., 2020; Waisman et al., 2024).

Perceived CMO Value Creation (PVC) captures the CMO's assessment that platform-mediated growth investment generates incremental impact worth its cost and remains defensible under cross-functional scrutiny. This construct follows from marketing accountability research, which shows that marketing value depends not only on outcomes but also on whether performance claims can be linked to enterprise value, profit, and growth (Morgan et al., 2021; O'Sullivan & Butler, 2010; Rust et al., 2004).

Growth Investment Capability (GIC) captures the governance and learning capability through which firms institutionalize evidence standards and validation capacity across recurring

allocation cycles. The dynamic capability literature supports a distinction between a focal microfoundation and a broader capability architecture. Evidence Integration is the focal microfoundation. It refers to the repeatable decision-making routine through which managers combine platform data with internal outcomes, experimental evidence, customer quality signals, and cross-functional interpretation before making allocation decisions. The dynamic capabilities literature underlines the need for such routines, while the qualitative phase identifies Evidence Integration as their empirical form. GIC is the broader architecture. It shapes decision rights, evidence thresholds, experimentation and measurement infrastructure, and disciplined reallocation routines (Barreto, 2010; Hilliard & Goldstein, 2019; Liu et al., 2020; Teece, 2007). This logic leads to the conceptual model in Figure 1, which structures the hypotheses developed below.

2.9 Hypotheses Development

The conceptual model separates perceived value claims from objective decision quality. Perceived CMO Value Creation (PVC) captures the perceived value and defensibility of the focal platform-mediated investment, based on the evidence available at the time of judgment. It does not measure objective CMO performance or the objective quality of the final allocation decision. When platform feedback overstates incremental impact, higher PVC can reflect an inflated value claim rather than superior value creation.

Prior advertising-measurement research shows that platform and attribution metrics can diverge from causal effects because observed conversions are not always incremental (Gordon et al., 2019; Gordon et al., 2022; Lewis & Rao, 2015). Evidence Integration addresses this problem by adding backend, experimental, and customer-quality evidence before budget judgment. In the quantitative experiment, this mechanism is operationalized through Evidence Triangulation. In the focal scenario, this additional evidence calls into question the incrementality of the platform-reported results. Evidence Triangulation should therefore reduce the perceived value and defensibility of the focal investment. This leads to the first hypothesis:

H1: Evidence Triangulation is negatively related to Perceived CMO Value Creation.

Performance Learning Distortion captures distorted belief updating following exposure to platform-mediated performance feedback. Organizational learning research shows that ambiguous feedback can distort managerial interpretation and subsequent adjustment (Greve, 2003; March & Olsen, 1975; Schumacher et al., 2020). In digital advertising, this risk is stronger

because platform-reported performance can overstate true incrementality (Gordon et al., 2019; Lewis & Rao, 2015). When respondents receive only platform-side evidence, they are more likely to update their incrementality beliefs upward. When respondents receive triangulated evidence, they can compare platform signals with backend outcomes, experimental evidence, and customer-quality indicators. Evidence Triangulation should therefore reduce distorted learning. The same logic implies a second prediction for distorted learning:

H2: Evidence Triangulation is negatively related to Performance Learning Distortion.

PLD should increase perceived value claims. Prior research on performance feedback shows that managers can interpret ambiguous feedback too optimistically, especially when visible signals are overweighted (Greve, 2003; Schumacher et al., 2020). Marketing accountability research further shows that marketing value depends on whether performance claims appear credible and defensible (Hanssens & Pauwels, 2016; Rust et al., 2004). If respondents update their beliefs about incrementality more optimistically, they should also judge the campaign as more valuable and defensible. This does not imply that distortion improves true decision quality. It means that distorted learning can inflate the perceived value and defensibility of the focal investment. Thus, hypothesis three states:

H3: Performance Learning Distortion is positively related to Perceived CMO Value Creation.

The preceding arguments imply a mediation mechanism. Research on advertising measurement shows that broader evidence can improve causal interpretation when platform-reported metrics are incomplete or biased (Gordon et al., 2019; Gordon et al., 2022; Kohavi et al., 2009). Evidence Triangulation should therefore reduce PLD by providing corrective evidence. Lower PLD should, in turn, reduce inflated perceptions of the focal investment's value and defensibility. The indirect effect is therefore corrective rather than value-destructive. Evidence Triangulation should improve judgment quality by reducing inflated value claims. Combining these arguments yields the mediation hypothesis:

H4: Evidence Triangulation has an indirect effect on Perceived CMO Value Creation through Performance Learning Distortion.

Growth Investment Capability should shape how distorted learning translates into perceived value claims. Dynamic-capabilities and marketing-capabilities research shows that firms differ in how they sense, validate, and reconfigure resources under uncertainty (Barreto, 2010; Day, 2011; Teece, 2007). Marketing accountability research also suggests that stronger measurement routines and financial justification improve the credibility of marketing decisions (Hanssens &

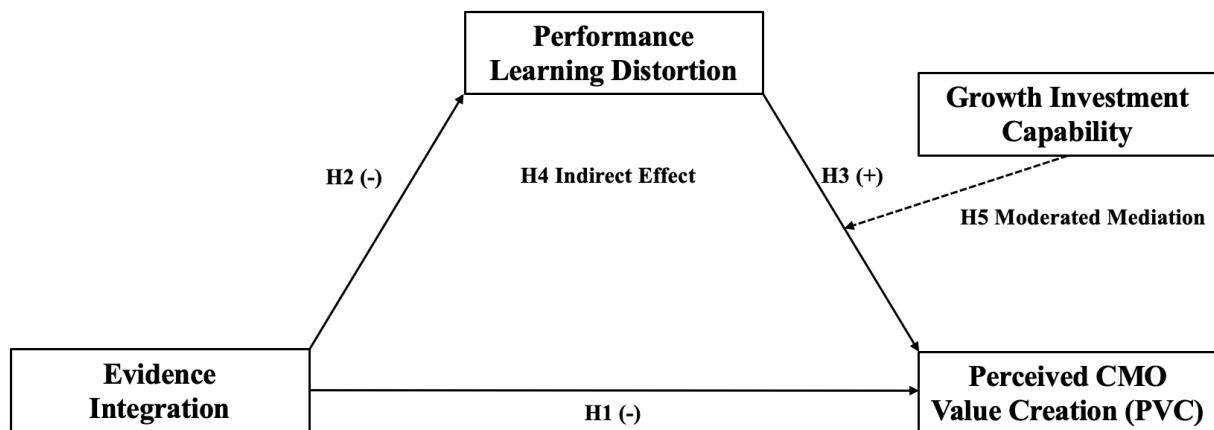
Pauwels, 2016; O’Sullivan et al., 2009). Firms with stronger GIC have clearer evidence standards, stronger validation routines, cross-functional review, and more disciplined reallocation processes. These conditions should make weak or inflated evidence less likely to be accepted as proof of value. The PLD on PVC relationship should therefore be weaker when GIC is high. As a result, the indirect effect of Evidence Triangulation on PVC through PLD should depend on GIC. The final hypothesis introduces GIC as a boundary condition:

H5: The indirect effect of Evidence Triangulation on Perceived CMO Value Creation through Performance Learning Distortion is conditional on Growth Investment Capability.

Together, these hypotheses test whether triangulated evidence corrects inflated perceptions of platform-based value. The model does not assume that lower PVC is worse in all cases. In the focal scenario, lower PVC after triangulated evidence indicates better calibration of an over-attributed campaign. The behavioral implication is tested using Budget Change, which captures whether corrected beliefs and value perceptions translate into more restrained allocation decisions.

Figure 1

Conceptual model



Note. Evidence Integration is the broader microfoundation and is operationalized as Evidence Triangulation in the vignette. PVC captures the perceived value and defensibility of the focal investment, not objective decision quality. H4 tests the indirect effect on PVC through PLD. H5 tests whether GIC conditions the PLD → PVC path. ET = Evidence Triangulation; PLD = Performance Learning Distortion; PVC = Perceived CMO Value Creation; GIC = Growth Investment Capability.

3. Methodology

The dissertation employs an exploratory sequential mixed-methods design because the focal routine first had to be identified in executive practice and then tested under controlled conditions. The qualitative strand draws on semi-structured expert interviews to identify Evidence Integration (EI) as the focal dynamic-capability microfoundation and to refine construct definitions and measurement items. The quantitative strand uses a scenario-based experimental vignette to test whether Evidence Triangulation (ET), the experimental operationalization of EI, reduces Performance Learning Distortion (PLD) and alters Perceived CMO Value Creation (PVC). Findings are integrated by assessing convergence, complementarity, and divergence across the qualitative and quantitative strands (Creswell & Plano Clark, 2018; Saunders et al., 2019).

3.1 Research Design

The study examined how senior marketing decision-makers interpret platform-provided performance signals during recurring budget reviews. This focus made an exploratory sequential design appropriate for two reasons. First, Evidence Integration (EI) needed to be grounded in executive practice before it could be operationalized and tested quantitatively. Second, platform environments often lack observable ground truth for incrementality, limiting causal inference from naturally occurring dashboard metrics. A controlled vignette, therefore, enabled a focused test of the proposed pathway while holding key conditions constant (Aguinis & Bradley, 2014; Creswell & Plano Clark, 2018; Saunders et al., 2019; Shadish et al., 2002).

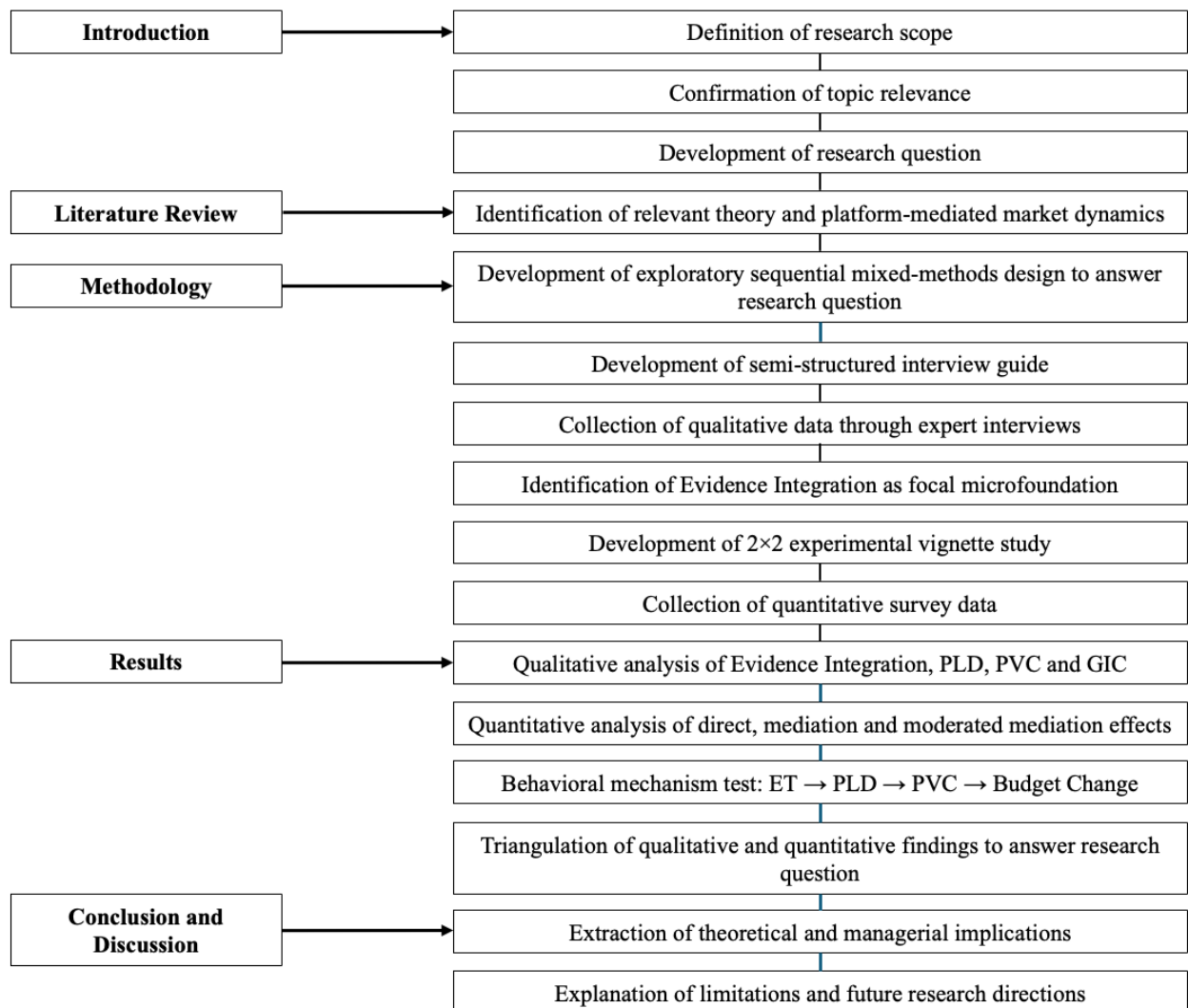
The research logic was abductive. The literature review provided initial constructs, the interviews refined the mechanism, and the experiment tested it under controlled conditions.

EI was selected using a transparent protocol. Candidate routines were extracted from coded interview data and evaluated against four criteria: salience in executive accounts, causal proximity to distorted learning, managerial actionability, and measurability in a standardized scenario. Evidence Integration showed the strongest fit and was selected as the focal microfoundation (Appendix B). Coding followed a Gioia-style approach, moving from first-order informant terms to second-order themes and aggregate dimensions (Gioia, 2021). Coding notes and construct-development decisions were retained to support auditability.

The quantitative strand tested whether Evidence Triangulation (ET), the experimental operationalization of EI, reduced Performance Learning Distortion (PLD) and altered Perceived CMO Value Creation (PVC). It also tested whether Growth Investment Capability (GIC) moderated the PLD on PVC relationship. GIC was treated as a theoretically motivated boundary condition rather than a statistical control. Hypotheses were evaluated using regression-based conditional process logic with bootstrapped inference (Hayes, 2018). Indirect effects were assessed using bootstrapped confidence intervals, consistent with current mediation-analysis practice (Hayes, 2018). Mixed-methods integration compared qualitative mechanism claims with quantitative results to assess convergence, complementarity, and divergence across strands (Denzin, 1978; Jick, 1979). Figure 2 summarizes the exploratory sequential design.

Figure 2

Research design and methodological sequence



3.2 Data Collection

3.2.1 Secondary Data Collection

The literature review synthesized research across four areas: platform governance and information asymmetries, causal measurement in digital advertising, learning from performance feedback, and dynamic capabilities and microfoundations. Searches were conducted in Scopus, Web of Science, and Google Scholar using keyword combinations aligned with these themes. The search prioritized recent peer-reviewed research. Older foundational sources were included when central to the construction of definitions or causal logic.

Inclusion focused on peer-reviewed articles and foundational books. Practitioner and regulatory sources were used selectively for contextual evidence when their methods were sufficiently transparent (Saunders et al., 2019; Tranfield et al., 2003). Studies were included if they addressed platform governance and information asymmetries, causal measurement or incrementality in digital advertising, or executive learning, accountability, and capability building under noisy feedback. Studies outside these domains were excluded unless conceptually foundational. The review informed the interview guide (Appendix A, Table 10), the platform-style signals in the vignette, the distortion pattern, and construct-consistent item wording.

3.2.2 Primary Data Collection – Qualitative

Seventeen CMOs or equivalent senior marketing decision-makers were interviewed. Participants were eligible if they had direct decision rights or strong recommendation authority over recurring marketing or growth allocations. They also needed experience with material budget decisions and exposure to platform-mediated growth channels in which platform evidence was used to justify budgets.

Sampling was purposive and aimed at information-rich cases across industries and business models (Patton, 2002). This variation helped identify recurring mechanisms across different platform-dependent decision contexts. Participants were recruited through the author's professional network, warm referrals, and targeted outreach. The final sample included senior decision-makers from consumer brands, subscription businesses, marketplaces, automotive, retail, hospitality, fintech, adtech, and advisory contexts.

Interviews were semi-structured to balance comparability with flexibility (Kvale & Brinkmann, 2009). The interview guide followed a funnel structure. It first covered the participant's role and allocation context. It then examined platform dependence and opacity, evidence use, suspected cases of distorted learning, and routines used to validate or challenge platform-reported performance. Later questions focused on governance practices, cross-functional decision processes, and capability conditions linked to Growth Investment Capability (GIC).

Interviews were conducted online, recorded with consent, transcribed, and pseudonymized. They focused on concrete decision episodes rather than general opinions. This approach helped identify how executives interpreted platform feedback, defended budget decisions, and integrated evidence before reallocating spend. An anonymized expert overview is reported in the Results section, and the interview guide is provided in Appendix A.

Data collection followed an explicit stopping logic. Core themes had largely stabilized by interview twelve. The remaining interviews were used to test pattern stability, add contextual contrast, and explore boundary conditions (Francis et al., 2010; Guest et al., 2006). This process supported the selection of Evidence Integration as the focal microfoundation. Evidence Integration combines platform data, internal outcomes, testing evidence, and cross-functional interpretation before allocation decisions.

3.2.3 Primary Data Collection – Quantitative

Quantitative data were collected through a scenario-based experimental vignette study to test the causal mechanism under controlled feedback conditions (Aguinis & Bradley, 2014; Shadish et al., 2002). The vignette mirrored a simplified CMO budget review. Participants first made an initial budget decision and estimated campaign incrementality. They then reviewed platform-style performance feedback and made a second budget decision. This structure allowed the study to observe how the assigned evidence changed beliefs and allocation choices.

Respondents were instructed to answer as the CMO responsible for the decision. The design tested CMO-like decision logic rather than claiming that all respondents were actual CMO decision-makers. The scenario elicited incrementality beliefs before and after exposure to the assigned evidence, enabling PLD to be measured as belief updating under controlled feedback conditions. This operationalization separates the conceptual label from the empirical measure. Conceptually, PLD refers to distorted learning under platform-mediated feedback.

Empirically, the experiment measures directional belief updating after exposure to potentially inflated platform evidence. It does not verify whether each individual update deviates from an objective incrementality benchmark.

The study used a between-subjects 2×2 design. It manipulated Evidence Triangulation (absent vs. present/strong) and GIC (low vs. high). The two factors were varied independently, allowing the effects of Evidence Triangulation, GIC, and their interaction to be estimated separately. Manipulation checks assessed whether triangulated-evidence conditions were perceived as involving stronger Evidence Integration and whether high-GIC conditions were perceived as involving stronger validation capability.

Evidence Triangulation operationalized Evidence Integration by requiring respondents to consider platform data alongside backend, experimental, and customer-quality evidence. GIC was manipulated as an organizational governance context, including evidence standards, validation capacity, decision rights, and reallocation discipline. Platform opacity and feedback distortion were held constant at a high level across conditions. Because the respondent pool was not restricted to CMOs, GIC was operationalized at the vignette level rather than as a self-report trait. This reduced referent ambiguity and key-informant validity concerns.

The survey was fielded between 10 April and 22 April 2026 and generated 316 initial responses. After removing incomplete and low-quality responses, the final analytical sample comprised 209 respondents. Data were exported from Qualtrics to SPSS and cleaned using a pre-specified sequence. Remaining cases were screened for attention-check failures, implausibly short completion times, and inconsistent response patterns.

The final sample was balanced by gender, with 103 male respondents (49.3%) and 106 female respondents (50.7%). Respondents ranged from 18 to 74 years old ($M = 38.13$, $SD = 13.67$). Most respondents had a marketing or growth background (52.2%), and 32.1% had more than 10 years of marketing experience. In addition, 65.5% reported being at least very familiar with digital marketing or performance metrics.

Measures combined behavioral belief updating with perceptual scales. PLD was calculated as post-feedback incrementality belief minus prior incrementality belief. PVC was measured with a multi-item Likert scale assessing perceived value-for-money and defensibility. GIC was manipulated at the vignette level and assessed through treatment-fidelity checks. Exact item wording, response anchors, vignette text, and construct mappings are documented in Appendices C and D.

3.3 Data Analysis

Quantitative analysis proceeded in four steps. First, descriptive statistics and condition-balance checks assessed sample structure and randomization quality. Second, scale diagnostics assessed internal consistency using Cronbach's α . Evidence Integration ($\alpha = .935$), Growth Investment Capability ($\alpha = .927$), and Perceived CMO Value Creation ($\alpha = .948$) all showed excellent reliability. Mean scores were therefore used as composite measures in the subsequent analyses. Third, manipulation checks tested whether participants perceived the Evidence Triangulation and GIC conditions as intended. A realism check assessed whether respondents found the vignette plausible. Fourth, hypotheses were tested using regression-based conditional process analysis.

Hypotheses were tested using PROCESS Model 4, Model 14, and Model 6, each with 5,000 bootstrap samples (Hayes, 2018). Model 4 tested the indirect effect of Evidence Triangulation on PVC through PLD. Model 14 tested whether this indirect effect was conditional on GIC. Model 6 tested the behavioral extension from Evidence Triangulation to Budget Change through PLD and PVC. Evidence Triangulation was modeled as the independent variable, behavioral PLD as the primary mediator, PVC as the main outcome, and perceived GIC as the moderator of the PLD–PVC path. The binary GIC condition was used for experimental assignment and manipulation checks, whereas the moderation test used the perceived GIC composite. PLD was entered as the behavioral belief-updating measure described in Section 3.2.3. Inference relied on bootstrapped confidence intervals for indirect and conditional indirect effects.

3.4 Research Quality & Ethics

Qualitative rigor was assessed using Lincoln and Guba's (1985) criteria of credibility, transferability, dependability, and confirmability. Credibility was supported by purposive heterogeneity, triangulation with the quantitative strand, and an explicit saturation rule (Francis et al., 2010; Guest et al., 2006). Dependability and confirmability were supported by an audit trail comprising the interview guide, coding rules, coding notes, and construct-development logic (Gioia, 2021). Transferability was supported by anonymized contextual detail and excerpt-grounded interpretation (Krippendorff, 2004).

Quantitative validity was supported by random assignment, standardized exposure, piloting, and pre-specified manipulation checks (Shadish et al., 2002). Construct validity was

strengthened through domain-anchored item development, refinement of interview language, and item-construct mapping (Gioia, 2021; Hinkin, 1998; O’Sullivan et al., 2009). Procedural remedies reduced common-method and demand effects by separating belief elicitation from judgment items, using neutral wording, and applying attention and consistency checks. Analysis scripts, coding materials, and anonymized outputs were retained to support auditability.

All participants provided informed consent. Interview data were pseudonymized, identifying details were removed, and survey data were handled in accordance with confidentiality principles for academic research (Saunders et al., 2019).

4. Results

4.1 Qualitative Results

Table 1 provides an anonymized overview of the expert sample and shows how each interview contributed to the study's decision-making context.

Table 1

Expert overview

Code	Strategic lens	Decision seniority	Business context	Primary decision domain	Distinctive relevance to study
E01	Enterprise CMO / growth transformation	Full budget owner	Global QSR / loyalty ecosystem	Enterprise allocation, digital acquisition, measurement infrastructure	Shows how data infrastructure, finance alignment, LTV logic, and machine learning shape allocation quality.
E02	CMO / CRO / traffic-platform leadership	Full budget owner	Subscription D2C / e-commerce	Channel allocation, pricing investment, commercial investment	Provides strong evidence on backend measurement, incrementality testing, predicted LTV, and platform-versus-backend logic.
E03	CEO / ex-CMO / transformation leader	Full budget owner	Buy-and-build / home-services platform	Growth investment, M&A, performance marketing	Highlights testing over false precision and shows how channel interdependencies can distort attribution.
E04	Managing director / FMCG brand leadership	Full budget owner	FMCG / omnichannel retail	Brand, trade, and digital allocation	Shows how digital performance depends on trade activation, distribution, and store execution.
E05	Automotive CMO	Full budget owner	Automotive	Media allocation and funnel interpretation	Highlights funnel-based decision logic, configurator metrics, downstream signals, and partial visibility between spend and sale.
E06	Consumer brand CMO	Full budget owner	Better-for-you snacking / FMCG	Brand building, awareness, and	Demonstrates investor pressure, awareness measurement, retail

Code	Strategic lens	Decision seniority	Business context	Primary decision domain	Distinctive relevance to study
				performance trade-offs	constraints, and internal brand-performance tensions.
E07	Luxury CMO / brand steward	Full budget owner	Luxury automotive	Brand investment and long-term equity management	Provides a contrast case where desirability, scarcity, and brand meaning outweigh short-term platform efficiency.
E08	Media-commerce / platform growth executive	Full budget owner	E-commerce / media platform	Traffic allocation, attention quality, monetization	Brings an attention-quality lens and distinguishes operational precision from strategic understanding.
E09	Marketplace / mobility growth executive	Full budget owner	Two-sided marketplace	Allocation across pricing, incentives, supply, demand, and marketing	Reframes allocation as system equilibrium and shows why platform KPIs can miss marketplace-level effects.
E10	Content / CRM / customer-journey specialist	Major influence, not sole owner	Digital finance / content-led customer journey	Content, CRM, SEO, design, and journey quality	Adds a non-CMO perspective on trust, clarity, user journey quality, and less visible forms of value creation.
E11	B2B fintech / scale-up CMO	Full budget owner	B2B fintech / SaaS-like growth	Demand generation, pipeline quality, revenue quality	Shows threshold-based allocation, payback logic, pipeline quality, and finance-aligned budget defense.
E12	Global marketing / customer transformation leader	Full budget owner	Energy / banking / platform-enabled customer businesses	Customer transformation, capability building, strategic allocation	Emphasizes that platform data is useful for action but insufficient for interpretation.
E13	Retail CMO / CTO	Full budget owner	Grocery retail / retail media	Retail growth, loyalty, customer data, technology integration	Adds a combined commercial-and-technology lens on loyalty data, first-party signals, and retail media value.
E14	CCO / CTO	Full budget owner	Hospitality / loyalty ecosystem	Commercial strategy, digital product, loyalty, technology	Shows how customer data, digital experience, loyalty, and

Code	Strategic lens	Decision seniority	Business context	Primary decision domain	Distinctive relevance to study
					commercial action form one decision loop.
E15	Retail CMO	Full budget owner	Discount retail	Omnichannel retail, content systems, customer value perception	Highlights execution infrastructure, store reality, content consistency, and operational design as drivers of decision quality.
E16	Platform / adtech CMO	Full budget owner	Adtech / AI-driven advertising	Evidence standards, interpretation routines, capital allocation logic	Provides the clearest articulation of local platform truth versus whole-business truth.
E17	Senior marketing strategy advisor	Cross-company influence	Advisory / cross-industry	Marketing and sales decisioning across firms	Provides pattern-recognition across firms and serves as a challenge case for generalizability.

Note. Revenue ranges are not reported to preserve participant confidentiality and avoid disclosing or inferring commercially sensitive information. The table instead reports business context and decision relevance at a level sufficient to assess sample heterogeneity.

The sample provided cross-context evidence on how executives interpreted platform feedback and defended allocation decisions. Coding followed a hybrid abductive approach. First-order informant terms were grouped into second-order themes and aggregate dimensions, following Gioia's (2021) data-to-theory logic. Category definition and iterative refinement were supported by content-analytic guidance on coding and categorization (Elo & Kyngäs, 2008; Krippendorff, 2004).

Four qualitative findings emerged. First, interviewees described Performance Learning Distortion as a structural feature of platform-mediated marketing, not merely a technical error. Platform data was useful for campaign steering but insufficient for assessing broader business impact. One expert captured this distinction clearly: "The platform is optimizing a local environment. You are trying to optimize a business" (E16). Distortion appeared as partial visibility, misleading proxy metrics, and attribution ambiguity.

Second, stronger decision-makers did not search for a single perfect metric. They combined heterogeneous signals before making allocation decisions. As one expert explained, stronger

measurement involves building “a hierarchy of data points” rather than relying on a single dashboard signal (E02). Evidence Integration emerged through backend validation, downstream outcome checks, triangulation, testing, and cross-functional interpretation. It therefore served as the focal microfoundation linking signal interpretation to more robust investment decisions.

Third, perceived value creation was broader than platform efficiency. Interviewees linked value to profit, customer quality, retention, lifetime value, repeat behavior, brand strength, pricing power, and enterprise contribution.

Fourth, Evidence Integration depended on a broader capability context. Its usefulness relied on data infrastructure, cross-functional coordination, review routines, and the translation of marketing evidence into enterprise language. One senior executive summarized this logic directly: “You cannot revolutionize marketing from inside marketing alone” (E01). This supports Growth Investment Capability as an evidence architecture.

Table 2 summarizes the qualitative findings and links each evidence pattern to the corresponding thesis construct.

Table 2

Qualitative findings and construct links

Qualitative finding	Evidence pattern	Link to construct
Platform feedback is useful but partial	Experts described platform dashboards as helpful for campaign steering but insufficient for judging true incrementality.	PLD
Visible KPIs can overstate business impact	Interviewees mentioned lead volume, attributed conversions, CTR, and ROAS as signals that may look strong without downstream validation.	PLD / PVC
Stronger routines integrate heterogeneous evidence	Experts described backend validation, geo-tests, customer-quality checks, LTV logic, and cross-functional review.	EI
Value is broader than platform efficiency	Value was linked to profit, retention, customer quality, repeat behavior, pricing power, and enterprise contribution.	PVC

Qualitative finding	Evidence pattern	Link to construct
Capability operates through evidence architecture	Stronger firms had clearer evidence standards, better data infrastructure, and finance/analytics involvement.	GIC

4.2 Quantitative Results

Random assignment produced balanced experimental cells. Condition sizes were 51, 54, 53, and 51 respondents. Evidence Triangulation was absent for 104 respondents and present for 105 respondents. Growth Investment Capability was low for 105 respondents and high for 104 respondents.

Table 3 reports condition means for PLD, PVC, and Budget Change. The descriptive pattern already reflected the expected mechanism. Evidence Triangulation conditions showed lower PLD, lower PVC, and more restrained budget decisions than platform-only conditions.

Table 3

Condition means by experimental cell

Condition	ET	GIC	PLD M	PVC M	Budget Change M
C1	0	0	10.04	5.16	0.57
C2	1	0	-10.96	3.29	-0.78
C3	0	1	2.49	4.99	0.13
C4	1	1	-12.33	3.20	-0.86

Note. ET = Evidence Triangulation; GIC = Growth Investment Capability. PLD was calculated as post-feedback incrementality belief minus prior incrementality belief. Positive PLD values indicate upward belief updating in response to performance evidence; negative values indicate downward correction. Positive Budget Change values indicate more expansionary budget decisions.

4.2.1 Descriptive Statistics and Correlations

Table 4 presents descriptive statistics and correlations for the core variables. PVC had a mean of 4.153, $SD = 1.425$. PLD had a mean of -2.761, $SD = 16.641$. Budget Change had a mean of -0.239, $SD = 1.193$. On average, respondents showed a slight tendency toward downward belief correction and more restrained budget allocation.

PLD was positively correlated with PVC, $r = .664, p < .001$, and with Budget Change, $r = .524, p < .001$. PVC was also positively correlated with Budget Change, $r = .644, p < .001$. More optimistic belief updating was therefore associated with higher perceived CMO value creation and more expansionary budget decisions.

Table 4

Descriptive statistics and intercorrelations

Variable	M	SD	1	2	3
1. PLD	-2.761	16.641	—		
2. PVC	4.153	1.425	.664***	—	
3. Budget Change	-0.239	1.193	.524***	.644***	—

Note. N = 209. Positive PLD values indicate upward belief updating following the review of performance evidence. Positive Budget Change values indicate more expansionary budget decisions. *** $p < .001$.

4.2.2 Manipulation Checks

Manipulation checks supported treatment validity, as summarized in Table 5. Evidence Triangulation increased perceived Evidence Integration, $F(3, 205) = 58.83, p < .001, \eta^2 = .463$. Mean EI was 3.14 in Condition 1, 5.20 in Condition 2, 3.17 in Condition 3, and 5.58 in Condition 4.

The GIC manipulation also worked as intended. A two-way ANOVA showed a significant GIC effect on perceived capability, $F(1, 205) = 122.67, p < .001, \eta^2 = .374$. The ET main effect, $F(1, 205) = 0.93, p = .336$, and the $ET \times GIC$ interaction, $F(1, 205) = 0.60, p = .441$, were not significant. The vignette was perceived as realistic, $M = 3.99, SD = 0.75$, with 81.3% rating it as realistic or very realistic.

Table 5

Manipulation checks

Dependent variable	Effect	Statistic
Evidence Integration	Condition	$F(3, 205) = 58.83, p < .001, \eta^2 = .463$
GIC perception	GIC	$F(1, 205) = 122.67, p < .001, \eta^2 = .374$
GIC perception	ET	$F(1, 205) = 0.93, p = .336$

Dependent variable	Effect	Statistic
GIC perception	ET × GIC	$F(1, 205) = 0.60, p = .441$

4.2.3 Hypothesis Tests

Hypothesis 1 predicted that Evidence Triangulation would reduce Perceived CMO Value Creation for the focal platform-mediated investment. A simple OLS regression showed a significant negative effect of Evidence Triangulation on PVC, $b = -1.829$, $SE = 0.151$, $t = -12.092$, $p < .001$, 95% CI $[-2.128, -1.531]$. The model explained 41.4% of the variance in PVC, $R^2 = .414$, $F(1, 207) = 146.215$, $p < .001$. Respondents exposed to triangulated evidence therefore perceived the focal investment as less valuable and less defensible than those exposed to platform-only evidence. This negative effect should not be interpreted as evidence that better evidence reduces true CMO value creation. It indicates that triangulated evidence reduced an inflated perceived value claim for the focal investment. Hypothesis 1 is supported.

Hypothesis 2 predicted that Evidence Triangulation would reduce Performance Learning Distortion. A two-way ANOVA showed a significant main effect of Evidence Triangulation on PLD, $F(1, 205) = 87.01$, $p < .001$, $\eta^2 = .298$. Without triangulated evidence, respondents updated their beliefs upward or only slightly downward. Mean PLD was 10.04 in Condition 1 and 2.49 in Condition 3. With triangulated evidence, respondents corrected their beliefs downward. Mean PLD was -10.96 in Condition 2 and -12.33 in Condition 4. GIC also showed a smaller main effect on PLD, $F(1, 205) = 5.39$, $p = .021$, $\eta^2 = .026$. The ET × GIC interaction was not significant, $F(1, 205) = 2.59$, $p = .109$, $\eta^2 = .012$. Evidence Triangulation therefore reduced PLD across both capability contexts. The same pattern appeared in PROCESS Model 4. Evidence Triangulation reduced PLD, $b = -17.821$, $SE = 1.947$, $t = -9.152$, $p < .001$, 95% CI $[-21.660, -13.982]$. The model explained 28.8% of the variance in PLD, $R^2 = .288$. Hypothesis 2 is supported.

Hypothesis 3 predicted that higher PLD would increase PVC. PROCESS Model 4 showed a positive effect of PLD on PVC, $b = 0.038$, $SE = 0.005$, $t = 8.127$, $p < .001$, 95% CI $[0.029, 0.048]$. More optimistic belief updating was therefore associated with stronger perceived value and defensibility of the focal platform-mediated investment. Hypothesis 3 is supported.

Hypothesis 4 predicted that Evidence Triangulation would have an indirect effect on PVC through PLD. PROCESS Model 4 with 5,000 bootstrap samples supported this mediation pathway. Evidence Triangulation reduced PLD, $b = -17.821$, $SE = 1.947$, $t = -9.152$, $p < .001$, 95% CI $[-21.660, -13.982]$. PLD positively predicted PVC, $b = 0.038$, $SE = 0.005$, $t = 8.127$, $p < .001$, 95% CI $[0.029, 0.048]$. The indirect effect was significant, effect = -0.682 , BootSE = 0.115 , 95% BootCI $[-0.927, -0.478]$. Because the bootstrap confidence interval did not include zero, Hypothesis 4 is supported. Table 6 summarizes the direct and indirect path tests for H1–H4.

Table 6

Direct and indirect path tests for H1–H4

Hypothesis	Path / Effect	Estimate	SE / BootSE	t	p	95% CI
H1	ET → PVC	-1.829	0.151	-12.092	< .001	[-2.128, -1.531]
H2	ET → PLD	-17.821	1.947	-9.152	< .001	[-21.660, -13.982]
H3	PLD → PVC	0.038	0.005	8.127	< .001	[0.029, 0.048]
H4	ET → PLD → PVC	-0.682	0.115	-	-	[-0.927, -0.478]

Note. H1–H3 report regression coefficients. H4 reports the bootstrapped indirect effect from PROCESS Model 4 with 5,000 bootstrap samples. For H4, inference is based on the bootstrap confidence interval. ET = Evidence Triangulation; PLD = Performance Learning Distortion; PVC = Perceived CMO Value Creation. PVC captures the perceived value and defensibility of the focal platform-mediated investment.

Hypothesis 5 predicted that the indirect effect of ET on PVC through PLD would depend on GIC. PROCESS Model 14 was estimated with PVC as the dependent variable, ET as the independent variable, PLD as the mediator, and perceived GIC as the moderator of the PLD → PVC path. The overall model was significant and explained 56.0% of the variance in PVC, $R^2 = .560$.

PLD positively predicted PVC, $b = 0.0387$, $SE = 0.0047$, $t = 8.1915$, $p < .001$, 95% CI $[0.0294, 0.0480]$. The direct effect of ET on PVC remained negative and significant, $b = -1.1401$, $SE = 0.1600$, $t = -7.1256$, $p < .001$, 95% CI $[-1.4556, -0.8247]$. This pattern indicates partial mediation and is consistent with an additional direct pathway from triangulated evidence to perceived value claims beyond belief updating alone. Substantively, triangulated evidence may reduce PVC not only by correcting incrementality beliefs, but also by making the investment less defensible under broader customer-quality and backend-performance evidence.

GIC had no significant direct effect, $b = 0.0609$, $SE = 0.0457$, $t = 1.3342$, $p = .184$, 95% CI [-0.0291, 0.1509]. The $PLD \times GIC$ interaction was also not significant, $b = -0.0006$, $SE = 0.0028$, $t = -0.2164$, $p = .829$, 95% CI [-0.0061, 0.0049]. The conditional indirect effect of ET on PVC through PLD was significant at low, mean, and high levels of GIC. However, the index of moderated mediation was not significant, index = 0.0107, BootSE = 0.0450, 95% CI [-0.0718, 0.1028].

The indirect effect was therefore present, but it did not vary significantly across levels of GIC. Hypothesis 5 is not supported. Table 7 reports the moderated mediation results for PVC as the outcome variable.

Table 7

Moderated mediation analysis with PVC as outcome

Panel A. Outcome variable: PVC

Predictor	B	SE	t	p	95% CI
ET	-1.1401	0.1600	-7.1256	< .001	[-1.4556, -0.8247]
PLD	0.0387	0.0047	8.1915	< .001	[0.0294, 0.0480]
GIC	0.0609	0.0457	1.3342	.184	[-0.0291, 0.1509]
PLD $\times$ GIC	-0.0006	0.0028	-0.2164	.829	[-0.0061, 0.0049]

Panel B. Conditional indirect effect of ET on PVC via PLD

Level of GIC	Effect	BootSE	BootLLCI	BootULCI
Low (-1.6579)	-0.7076	0.1388	-0.9871	-0.4385
Mean (0.3421)	-0.6862	0.1189	-0.9344	-0.4642
High (1.5921)	-0.6728	0.1384	-0.9555	-0.4125

Panel C. Index of moderated mediation

Index	BootSE	BootLLCI	BootULCI
0.0107	0.0450	-0.0718	0.1028

Note. Conditional indirect effects were probed at low, mean, and high observed levels of perceived GIC as estimated by PROCESS. Bootstrap confidence intervals are based on 5,000 samples. ET = Evidence Triangulation; PLD = Performance Learning Distortion; PVC = Perceived CMO Value Creation; GIC = Growth Investment Capability.

4.2.4 Behavioral Mechanism Test: Serial Mediation to Budget Change

Because the research question centers on budget judgments, a serial mediation model tested whether Evidence Triangulation affected Budget Change through PLD and PVC in sequence. PROCESS Model 6 showed a significant serial mediation effect.

Evidence Triangulation reduced PLD, $b = -17.821$, $p < .001$. PLD positively predicted PVC, $b = 0.038$, $p < .001$. Both PLD, $b = 0.011$, $p = .032$, and PVC, $b = 0.400$, $p < .001$, positively predicted Budget Change. The total indirect effect of ET on Budget Change was significant, effect = -0.928 , 95% bootstrap CI $[-1.230, -0.675]$. The serial indirect effect ET → PLD → PVC → Budget Change was also significant, effect = -0.273 , 95% bootstrap CI $[-0.418, -0.160]$. Once the mediators were included, the direct effect of ET on Budget Change was no longer significant, $b = -0.237$, $p = .155$.

This result supports the central behavioral mechanism. Evidence Triangulation affected budget decisions primarily through corrected belief updating and better-calibrated value claims. Table 8 reports the serial mediation results.

Table 8

Serial mediation from ET to Budget Change via PLD and PVC

Effect	Estimate	BootSE / SE	p	95% CI
Total effect of ET on Budget	-1.165	0.144	< .001	[-1.450, -0.881]
Direct effect of ET on Budget	-0.237	0.166	.155	[-0.565, 0.090]
Total indirect effect	-0.928	0.141	—	[-1.230, -0.675]
ET → PLD → Budget	-0.197	0.095	—	[-0.392, -0.020]
ET → PVC → Budget	-0.459	0.110	—	[-0.702, -0.268]
ET → PLD → PVC → Budget	-0.273	0.065	—	[-0.418, -0.160]

Note. Confidence intervals for indirect effects are bootstrap-based and may differ from symmetric intervals derived from standard errors. ET = Evidence Triangulation; PLD = Performance Learning Distortion; PVC = Perceived CMO Value Creation.

4.2.5 Additional Behavioral Analysis

A simple behavioral validation regressed Budget Change on PVC. PVC positively predicted Budget Change, $b = 0.539$, $SE = 0.045$, $t = 12.106$, $p < .001$, 95% CI [0.451, 0.627]. The model explained 41.5% of the variance in Budget Change, $R^2 = .415$. Higher perceived CMO value creation was therefore associated with larger budget increases.

4.2.6 Robustness Checks

Robustness checks supported the stability of the findings. Balance checks showed no significant differences across conditions in age, $F(3, 205) = 0.632$, $p = .595$, marketing experience, $F(3, 205) = 0.764$, $p = .515$, or gender, $\chi^2(3) = 0.939$, $p = .816$.

The serial mediation model was re-estimated after excluding two influential cases identified by PROCESS diagnostics. The substantive pattern remained unchanged. Evidence Triangulation still reduced PLD, $b = -18.805$, $p < .001$, 95% CI [-22.390, -15.220]. PLD still increased PVC, $b = 0.0408$, $p < .001$, 95% CI [0.0309, 0.0507]. The serial indirect effect also remained significant, effect = -0.2901, 95% BootCI [-0.4432, -0.1659].

Collinearity diagnostics were satisfactory. The highest VIF was 2.254, and all tolerance values exceeded .44.

Planned contrasts confirmed that Evidence Triangulation reduced PLD under both low GIC, $t(103) = 9.06$, $p < .001$, $d = 1.77$, and high GIC, $t(102) = 4.83$, $p < .001$, $d = 0.95$. GIC reduced PLD only when triangulated evidence was absent, $t(88.94) = 2.71$, $p = .008$, $d = 0.53$. With triangulated evidence, the GIC difference was not significant, $t(103) = 0.52$, $p = .603$, $d = 0.10$. Full results are reported in Appendix G.

4.2.7 Summary of Hypothesis Tests

Table 9 summarizes the hypothesis tests. Four hypotheses were supported. The moderated mediation hypothesis was not supported.

Table 9

Summary of hypothesis tests

Hypothesis	Statement	Result
H1	Evidence Triangulation is negatively related to PVC	Supported
H2	Evidence Triangulation is negatively related to PLD	Supported
H3	PLD is positively related to PVC	Supported
H4	ET has an indirect effect on PVC through PLD	Supported
H5	The indirect effect of ET on PVC through PLD is conditional on GIC	Not supported

Note. ET = Evidence Triangulation; PLD = Performance Learning Distortion; PVC = Perceived CMO Value Creation; GIC = Growth Investment Capability. PVC captures the perceived value and defensibility of the focal platform-mediated investment, not objective decision quality or verified causal impact.

4.3 Final Synthesis – Triangulation

The qualitative and quantitative findings converged on a core point. Platform-governed feedback is not a neutral representation of business impact. It is partial, contestable, and checked against broader evidence in stronger decision environments.

The qualitative phase identified Evidence Integration as the main corrective routine. The experiment supported this mechanism: Evidence Triangulation increased Evidence Integration, reduced PLD, and lowered inflated value and budget judgments.

Both strands also converged on the role of PLD. In the interviews, distorted learning appeared as a structural feature of platform-mediated decision-making rather than a rare measurement error. The experiment supported that interpretation. When respondents relied on platform-style feedback without triangulated evidence, their beliefs remained more optimistic, and their budget decisions became more expansionary. When additional evidence was available, beliefs corrected downward, and budget decisions became more conservative.

The serial mediation analysis sharpened this result. Evidence Triangulation reduced PLD. Lower PLD reduced PVC. Lower PVC then translated into more restrained budget decisions. The significant serial indirect effect linked the informational treatment to the final allocation outcome through a coherent decision mechanism.

The main divergence concerned Growth Investment Capability. The qualitative phase suggested that a stronger capability should reduce the consequences of distorted feedback. Interviewees repeatedly emphasized review routines, stronger evidence standards, cross-functional challenge, and better data flow across the firm. The quantitative phase did not support the hypothesized moderation of the PLD on PVC path. This result allows two interpretations. The first is theoretical: GIC may operate upstream by shaping evidence architecture, validation routines, and decision thresholds before distorted learning arises. The second is methodological: the vignette manipulation may have been too narrow to capture a multidimensional organizational capability. The present design cannot fully distinguish between these explanations. It therefore supports an upstream interpretation of GIC, but only as a plausible refinement rather than a conclusive finding. The non-significant moderation result refines the argument: GIC appears more relevant to the conditions under which PLD emerges than to the downstream translation of PLD into PVC.

This divergence did not weaken the core mechanism. It left open whether GIC is a weaker boundary condition, operates through another pathway, or was under-activated by the vignette.

Overall, triangulation across both strands supported a clear conclusion. CMOs operate in a setting in which platform signals can systematically distort learning. Broader evidence integration reduced that distortion. Lower distortion led to better-calibrated perceived value claims and more restrained budget allocation. Both strands converged on this mechanism. The moderating role of GIC remained unresolved.

5. Conclusion

5.1 Discussion of Main Findings

The central finding is straightforward: platform-governed feedback supports managerial execution but is insufficient for strategic judgment when considered in isolation. Across the literature review, the interviews, and the experiment, the same pattern emerged. Platform-reported performance can shape learning in systematically biased ways when treated as a proxy for true incrementality without broader validation (Gordon et al., 2019; Greve, 2003; Lewis & Rao, 2015).

The literature review established the structural tension between action and observability. Digital platforms increase speed, reach, and optimization capacity, but they also limit independent verification through opacity, platform-specific attribution logic, and incomplete visibility into downstream outcomes.

The qualitative findings showed how senior decision-makers respond to this tension in practice. Interviewees repeatedly described platform dashboards as operationally useful, but strategically partial. The recurring problem was not too little data. It was learning too much from the wrong data. Stronger decision processes therefore relied less on a single decisive KPI and more on a broader evidentiary base before scaling spend.

The quantitative findings tested this mechanism under controlled conditions. Evidence Triangulation strongly increased Evidence Integration and substantially reduced Performance Learning Distortion. Lower PLD was associated with lower Perceived CMO Value Creation and more conservative budget decisions. The serial mediation results sharpened this interpretation. Evidence Triangulation did not merely change how much information respondents saw. It changed how they interpreted causality, judged value, and reallocated budget.

The integrated implication is that CMOs need routines that connect platform metrics to broader business outcomes before budget decisions.

The findings also refined the meaning of value creation in platform settings. In the interviews, value was rarely defined solely by platform efficiency. Executives instead referred to profit, customer quality, retention, lifetime value, repeat behavior, and enterprise contribution. The quantitative results are consistent with that broader view. As distorted learning declined, PVC also declined because the perceived value claim became better calibrated. This is substantively

important. It suggests that inflated confidence in platform-reported performance can exaggerate perceived value creation and support budget decisions that are more optimistic than the evidence warrants. The negative ET on PVC effect should not be interpreted as showing that Evidence Integration reduces CMO value creation. It showed that triangulated evidence corrected an inflated value claim for the focal investment.

The main divergence concerned Growth Investment Capability. The qualitative phase suggested that a stronger capability should reduce the consequences of distorted feedback. Interviewees repeatedly emphasized review routines, stronger evidence standards, cross-functional challenge, and better data flow across the firm. The quantitative phase did not support the hypothesized moderation of the PLD on PVC path. The lack of support did not imply that GIC is irrelevant. It suggests, but does not conclusively establish, that GIC may operate upstream by shaping evidence architecture, validation routines, and decision thresholds. Once evidence was presented, GIC did not significantly alter how PLD translates into PVC. That interpretation is consistent with the interview material, although the design cannot fully rule out that the vignette captured GIC too narrowly. The non-significant moderation result, therefore, refined rather than weakened the argument: GIC appears to shape the evidence conditions under which PLD arises, rather than moderating the downstream relationship between PLD and PVC.

Two interview insights were especially valuable for interpreting these results. First, one senior executive described strong marketing decision-making as an enterprise measurement problem rather than a narrow media problem. Better decisions become possible when marketing, finance, and technology work from a shared evidence base and translate marketing outcomes into business language. Second, another executive stressed that the real objective is not simply to spend more efficiently but to avoid spending that should not be spent. Both insights align closely with the experimental results. Evidence Triangulation reduced PLD, and lower PLD translated into more restrained budget allocation.

Overall, the findings support a clear conclusion. CMOs operate in settings where platform-governed signals can systematically distort learning. Integrating broader evidence reduces that distortion. Lower distortion leads to better-calibrated perceived value claims and more restrained budget allocation. The central managerial challenge is therefore not simply to collect more performance data. It is to create better conditions for causal interpretation.

5.2 Theoretical Implications

This thesis makes three main theoretical contributions and identifies two extensions for future theory development. First, it introduces Performance Learning Distortion (PLD) as a mechanism linking platform governance to executive learning and resource allocation. Prior work has shown that attribution is noisy, that platform reporting can diverge from causal effects, and that learning under ambiguity is vulnerable to biased interpretation. The contribution here is to connect those streams at the level of recurring budget decisions. PLD explains how platform-governed feedback can systematically misguide judgments about true incrementality and how repeated reliance on such feedback can bias both value judgments and allocation choices.

Second, the thesis identifies Evidence Integration as a relevant microfoundation in platform-mediated marketing contexts. The interviews did not point to a single perfect measurement model. Instead, they revealed a routine of combining heterogeneous evidence before judgment. The experiment then showed that this logic matters behaviorally. Once broader evidence is introduced, distorted learning declines. This makes the dynamic capabilities perspective more concrete. The relevant capability is not abstract adaptability alone. It is the disciplined integration of multiple evidence sources under imperfect observability.

Third, the thesis sharpens Perceived CMO Value Creation as a perceptual accountability outcome. In platform environments, causal truth is often only partially observable at the time of decision. Under these conditions, value creation is not only an objective outcome. It is also a judgment about whether the claimed value appears credible and defensible under scrutiny. PVC, therefore, extends marketing accountability theory by locating value judgments within a contested evidentiary environment rather than assuming that causal impact is directly knowable.

Beyond these contributions, the findings also point to a future theoretical extension: how machine learning changes the structure and accessibility of marketing judgment. The interview material suggests that predictive systems can reduce waste when tied to strong first-party signals and clearly defined decision problems. At the same time, machine learning may increase the capability gap between firms. Large firms may be better positioned to combine proprietary data, testing infrastructure, analytics talent, and Evidence Integration routines. Smaller and medium-sized firms may gain access to powerful tools but may struggle more to validate opaque outputs and separate prediction from causal evidence. Future research should therefore

examine whether marketing measurement becomes more accessible across firm sizes or further advantages firms with stronger data and evidence infrastructures.

A second extension concerns the relationship between prediction and causality. As machine learning becomes more central to digital decision-making, firms may become better at predicting near-term outcomes without necessarily becoming better at understanding incremental business impact. This thesis suggests that future theorizing should distinguish more clearly between predictive power and evidentiary quality. High-quality decision-making requires both, but they are not the same thing.

5.3 Practical Implications

The findings have direct implications for CMOs, CFOs, data leaders, and firms that allocate substantial budgets through digital platforms.

First, platform dashboards should not be treated as sufficient evidence of value creation. They are useful for campaign steering, bid management, and within-channel optimization, but they are weaker tools for judging whether measured performance reflects true incremental business impact. Firms should therefore separate operational optimization from strategic judgment. Campaign execution can remain platform-informed. Budget defense and cross-channel reallocation should rely on broader evidence.

Second, firms need explicit evidence standards before scaling spend. Stronger organizations rarely rely on a single signal. They define what counts as enough evidence to increase, hold, or reduce budget. A strong dashboard is not enough. The pattern should also align with backend outcomes, customer-quality signals, or downstream evidence that matters financially.

Third, Evidence Integration should become routine rather than an exception. Operationally, this means comparing platform KPIs with backend outcomes, customer-quality signals, experimental evidence, and cross-functional interpretation before scaling spend. Multi-touch and data-driven attribution can support pattern recognition when used cautiously. When feasible, conversion lift studies, geo-tests, holdouts, media mix models, and post-purchase surveys can provide stronger checks on incrementality.

Fourth, first-party measurement infrastructure is now central to marketing judgment. Better decisions depend on cleaner customer-level data, stronger event collection, CRM integration,

cohort tracking, and stronger links between marketing, finance, and technology. Without these conditions, firms remain structurally exposed to PLD.

Fifth, budgeting should follow marginal value logic. Budget should be treated as the result of what still makes economic sense, not as a fixed entitlement to spend. This requires credible evidence, tighter finance dialogue, and a clearer translation of marketing outcomes into enterprise value.

Sixth, CMOs must translate marketing evidence into business language. A defensible budget argument is rarely built on clicks, CTR, or ROAS alone. It is built on whether those signals connect to profit, retention, customer quality, repeat behavior, or other outcomes recognized by CEOs, CFOs, boards, and investors.

Finally, machine learning should be treated as a decision infrastructure, not a substitute for causal judgment. Predictive systems can reduce waste when tied to clear decision problems and strong first-party data. They do not eliminate the need for causal interpretation. Firms that automate weak assumptions may scale distortion faster. Firms that embed machine learning within a strong evidence architecture may improve allocation quality.

5.4 Limitations

This thesis has several limitations that constrain generalizability. First, the quantitative phase relied on a scenario-based experiment with role-taking respondents. This design strengthens internal validity and tests CMO-level decision logic, but it cannot fully replicate C-suite judgment or real budget cycles. Actual CMO decisions unfold over time, involve multiple actors, and are shaped by politics, incentives, prior commitments, and category-specific priors. The observed effect sizes should therefore not be treated as direct estimates of field effects. The large ET effect on PLD ($\eta^2 = .298$) may overstate the magnitude of belief correction that would occur in organizational settings. The vignette may also have created demand characteristics, as participants in the Evidence Triangulation condition may have inferred that broader evidence was expected to support a more cautious decision.

Second, the dependent process was observed in a single decision episode rather than over time. PLD may become more consequential when small interpretive errors accumulate across repeated reallocation cycles. The present design cannot capture the full dynamics of compounding mislearning, delayed outcomes, or institutional lock-in.

Third, PLD was measured without an objective external benchmark for true incrementality. Because PLD was constructed as a post-minus-prior difference score, it may be affected by regression to the mean and lower measurement reliability. The measure captures directional belief updating under experimentally manipulated evidence, not verified causal error. This distinction matters because distortion implies deviation from a more accurate benchmark. In this study, that benchmark is theoretically motivated by the scenario evidence, but not independently observed. Future research should compare managerial belief updating against field-based incrementality benchmarks, such as geo-tests, holdouts, or conversion-lift studies.

Fourth, GIC was operationalized as a vignette-level context rather than as a fully developed organizational construct. This was appropriate for experimental control but narrowed the construct. Qualitative evidence suggests that GIC includes review cadence, evidence standards, decision rights, infrastructure depth, and stronger interfaces among marketing, finance, and technology. A short vignette cannot reproduce these routines or their cumulative effects across real budget cycles. The non-significant moderation result should therefore not be read as conclusive evidence that capability is irrelevant. It may indicate that GIC operates upstream, but it may also reflect the limits of representing a multidimensional capability in a controlled scenario.

Fifth, the study focused on platform-mediated growth settings and used a decision task centered on digital allocation logic. This improves conceptual precision but limits direct generalization. The strength of the mechanism may vary across sectors, business models, and measurement environments. Subscription businesses may observe retention and lifetime value more directly, while fast-moving consumer goods (FMCG) and automotive firms often rely on retail, dealer, or offline sales signals that are harder to connect to platform exposure.

Sixth, the qualitative sample was strong in seniority and diversity, but it remains an interpretive sample rather than a statistically representative population of CMOs. The interviews were valuable for mechanism building and construct refinement, but they cannot establish population prevalence.

Seventh, the study may not have captured all feedback mechanisms that CMOs use in practice. Retail sell-out data, brand tracking, sales-team feedback, partner-level lift analyses, market-level demand indicators, and other commercial signals may all matter in specific contexts.

5.5 Future Research

Several directions follow directly from these limitations and findings. First, future research should examine PLD longitudinally. Field studies that track recurring budget reviews over time would help determine whether distorted learning compounds, stabilizes, or self-corrects. This would be especially valuable where short-term platform efficiency conflicts with longer-term customer quality, retention, or brand effects.

Second, future work should measure GIC more comprehensively. The present findings suggest that capability may matter, but not in the simple moderating form initially hypothesized. Future studies should therefore operationalize GIC through observable routines and structures, such as review cadence, experimentation maturity, evidence thresholds, finance-marketing interfaces, and first-party data depth. Multi-informant designs would be especially useful here.

Third, future research should compare specific corrective tools rather than treating Evidence Triangulation as a single broad category. This includes multi-touch attribution, data-driven attribution, media mix models, geo-tests, conversion lift studies, holdout designs, post-purchase surveys, and cohort-based lifetime value analysis. The critical question is not whether triangulation matters. It is the mix of tools that best reduces PLD across different platforms, channels, and firm conditions.

Fourth, future work should examine how machine learning changes the structure and accessibility of marketing judgment. Predictive systems may reduce waste when they are tied to strong first-party data, clear decision problems, and disciplined Evidence Integration routines. However, these benefits are unlikely to be evenly distributed. Larger firms may be better positioned to combine proprietary data, testing infrastructure, analytics talent, and cross-functional evidence governance. Smaller firms may gain access to powerful tools but may struggle more to validate opaque outputs and distinguish prediction from causal evidence. Future research should therefore examine whether machine learning democratizes marketing measurement or reinforces capability gaps between firms.

Fifth, future work should examine how changes in platform governance reshape the quality of evidence. The Trade Desk is useful here as an illustration of a broader issue rather than a settled answer. The architecture of media decisioning is becoming more fluid. Inventory access, demand-sided platform (DSP) competition, retail media expansion, evolving platform interfaces, dashboard design, default attribution windows, and platform-reported optimization metrics may all alter the informational basis on which CMOs evaluate value creation. Future

studies should therefore treat the evidentiary environment of digital marketing as changing rather than fixed.

Sixth, future research should test the mechanism across business models. Subscription businesses, marketplaces, retail networks, FMCG firms, hospitality groups, and automotive players face different signal environments and degrees of downstream observability. The current thesis suggests that the mechanism is portable, but its strength and the most useful correction devices likely vary by context.

Seventh, future work should push more directly into the relationship between perceived value creation and organizational legitimacy. The present findings show that distorted learning inflates perceived value and budget expansion. A potential next step is to examine how CFO scrutiny, board expectations, investor pressure, and the financial language available to the CMO shape PVC.

5.6 Final Conclusion

Digital platforms have made marketing faster, more measurable, and more scalable. They have not made value creation self-evident. This thesis shows that platform-governed feedback can systematically distort executive learning when treated as sufficient evidence of causal impact. Broader evidence integration reduces that distortion. Lower distortion leads to better-calibrated value claims and more restrained budget allocation. The practical implication is simple to understand but hard to solve. CMOs do not merely need more data. They need better evidence.

Reference List

- Aguinis, H., & Bradley, K. J. (2014). Best practice recommendations for designing and implementing experimental vignette methodology studies. *Organizational Research Methods*, 17(4), 351–371. <https://doi.org/10.1177/1094428114547952>
- Albert, D. (2025). Rapid learning and adaptive search in complex environments: How underestimating noise in performance feedback can leverage and resolve errors of commission. *Strategy Science*, 10(4), 316–337. <https://doi.org/10.1287/stsc.2024.0246>
- Argote, L. (2012). *Organizational learning: Creating, retaining and transferring Knowledge*. Springer Science & Business Media.
- Barajas, J., Akella, R., Holtan, M., & Flores, A. (2016). Experimental designs and estimation for online display advertising attribution in marketplaces. *Marketing Science*, 35(3), 465–483. <https://doi.org/10.1287/mksc.2016.0982>
- Barney, J. B., Ketchen, D. J., Jr., & Wright, M. (2021). Resource-based theory and the value creation framework. *Journal of Management*, 47(7), 1936–1955. <https://doi.org/10.1177/01492063211021655>
- Barreto, I. (2010). Dynamic capabilities: A review of past research and an agenda for the future. *Journal of Management*, 36(1), 256–280. <https://doi.org/10.1177/0149206309350776>
- Beil, D., Duenyas, I., Leider, S., Li, J., & Qi, A. (2025). Human decision making in dynamic resource allocation. *Management Science*. Advance online publication. <https://doi.org/10.1287/mnsc.2021.01097>
- Benlian, A., Hilkert, D., & Hess, T. (2015). How open is this platform? The meaning and measurement of platform openness from the complementors' perspective. *Journal of Information Technology*, 30(3), 209–228. <https://doi.org/10.1057/jit.2015.6>
- Braun, M., & Schwartz, E. M. (2024). Where A/B Testing Goes Wrong: How Divergent Delivery Affects What Online Experiments Cannot (and Can) Tell You About How Customers Respond to Advertising. *Journal of Marketing*, 89(2), 71–95. <https://doi.org/10.1177/00222429241275886>
- Chan, D. X., Yuan, Y., Koehler, J., & Kumar, D. (2011). Incremental clicks – The impact of search advertising. *Journal of Advertising Research*, 51(4), 643–647. <https://doi.org/10.2501/JAR-51-4-643-647>

- Creswell, J. W., & Plano Clark, V. L. (2018). *Designing and conducting mixed methods research* (3rd ed.). SAGE.
- Cutolo, D., & Kenney, M. (2021). Platform-dependent entrepreneurs: Power asymmetries, risks, and strategies in the platform economy. *Academy of Management Perspectives*, 35(4), 584–605. <https://doi.org/10.5465/amp.2019.0103>
- Cyert, R. M., & March, J. G. (1963). *A behavioral theory of the firm*. Prentice-Hall.
- Day, G. S. (2011). Closing the marketing capabilities gap. *Journal of Marketing*, 75(4), 183–195. <https://doi.org/10.1509/jmkg.75.4.183>
- Decarolis, F., & Rovigatti, G. (2021). From mad men to math men: Concentration and buyer power in online advertising. *American Economic Review*, 111(10), 3299–3327. <https://doi.org/10.1257/aer.20190811>
- Denzin, N. K. (1978). *The research act: A theoretical introduction to sociological methods* (2nd ed.). McGraw-Hill.
- Donati, D., & Fong, H. (2025). The cost of banning TikTok: Implications for the digital advertising market. *Proceedings of the National Academy of Sciences*, 122(38), e2512043122. <https://doi.org/10.1073/pnas.2512043122>
- Dong, B., Zhuang, M., Fang, E., & Huang, M. (2023). Tales of Two Channels: Digital Advertising Performance Between AI Recommendation and User Subscription Channels. *Journal of Marketing*, 88(2), 141–162. <https://doi.org/10.1177/00222429231190021>
- Eaton, B., Elaluf-Calderwood, S., Sørensen, C., & Yoo, Y. (2015). Distributed tuning of boundary resources: The case of Apple’s iOS service system. *MIS Quarterly*, 39(1), 217–243. <https://doi.org/10.25300/MISQ/2015/39.1.10>
- Edelman, B. (2014). Pitfalls and fraud in online advertising metrics. *Journal of Advertising Research*, 54(2), 127–132. <https://doi.org/10.2501/JAR-54-2-127-132>
- Elo, S., & Kyngäs, H. (2008). The qualitative content analysis process. *Journal of Advanced Nursing*, 62(1), 107–115. <https://doi.org/10.1111/j.1365-2648.2007.04569.x>
- Francis, J. J., Johnston, M., Robertson, C., Glidewell, L., Entwistle, V., Eccles, M. P., & Grimshaw, J. M. (2010). What is an adequate sample size? Operationalising data saturation for theory-based interview studies. *Psychology & Health*, 25(10), 1229–1245. <https://doi.org/10.1080/08870440903194015>

- Gawer, A. (2014). Bridging differing perspectives on technological platforms: Toward an integrative framework. *Research Policy*, 43(7), 1239–1249.
<https://doi.org/10.1016/j.respol.2014.03.006>
- Gawer, A. (2021). Digital platforms' boundaries: The interplay of firm scope, platform sides, and digital interfaces. *Long Range Planning*, 54(5), 102045.
<https://doi.org/10.1016/j.lrp.2020.102045>
- Geradin, D., & Katsifis, D. (2020). “Trust me, I’m fair”: Analysing Google’s latest practices in ad tech from the perspective of EU competition law. *European Competition Journal*, 16(1), 11–54. <https://doi.org/10.1080/17441056.2019.1706413>
- Germann, F., Ebbes, P., & Grewal, R. (2015). The chief marketing officer matters! *Journal of Marketing*, 79(3), 1–22. <https://doi.org/10.1509/jm.14.0244>
- Ghazawneh, A., & Henfridsson, O. (2013). Balancing platform control and external contribution in third-party development: The boundary resources model. *Information Systems Journal*, 23(2), 173–192. <https://doi.org/10.1111/j.1365-2575.2012.00406.x>
- Gioia, D. A. (2021). A systematic methodology for doing qualitative research. *The Journal of Applied Behavioral Science*, 57(1), 20–29. <https://doi.org/10.1177/0021886320982715>
- Goldfarb, A., & Tucker, C. E. (2011). Privacy regulation and online advertising. *Management Science*, 57(1), 57–71. <https://doi.org/10.1287/mnsc.1100.1246>
- Gordon, B. R., Moakler, R., & Zettermeyer, F. (2022). Close enough? A large-scale exploration of non-experimental approaches to advertising measurement. *Marketing Science*, 42(4), 768–793. <https://doi.org/10.1287/mksc.2022.1413>
- Gordon, B. R., Zettermeyer, F., Bhargava, N., & Chapsky, D. (2019). A comparison of approaches to advertising measurement: Evidence from big field experiments at Facebook. *Marketing Science*, 38(2), 193–225. <https://doi.org/10.1287/mksc.2018.1135>
- Greve, H. R. (2003). *Organizational learning from performance feedback: A behavioral perspective on innovation and change*. Cambridge University Press.
- Guest, G., Bunce, A., & Johnson, L. (2006). How many interviews are enough? An experiment with data saturation and variability. *Field Methods*, 18(1), 59–82.
<https://doi.org/10.1177/1525822X05279903>

- Hanssens, D. M., & Pauwels, K. H. (2016). Demonstrating the value of marketing. *Journal of Marketing*, 80(6), 173–190. <https://doi.org/10.1509/jm.15.0417>
- Hayes, A. F. (2018). *Introduction to mediation, moderation, and conditional process analysis: A regression-based approach* (2nd ed.). Guilford Press.
- Helmond, A., van der Vlist, F. N., Burkhardt, M., & Seitz, T. (2021). The governance of Facebook platform. *AoIR Selected Papers of Internet Research*. <https://doi.org/10.5210/spir.v2021i0.12181>
- Hilliard, R., & Goldstein, D. (2019). Identifying and measuring dynamic capability using search routines. *Strategic Organization*, 17(2), 210–240. <https://doi.org/10.1177/1476127018755001>
- Hinkin, T. R. (1998). A brief tutorial on the development of measures for use in survey questionnaires. *Organizational Research Methods*, 1(1), 104–121. <https://doi.org/10.1177/109442819800100106>
- Homburg, C., & Wielgos, D. M. (2022). The value relevance of digital marketing capabilities to firm performance. *Journal of the Academy of Marketing Science*, 50(4), 666–688. <https://doi.org/10.1007/s11747-022-00858-7>
- Homburg, C., Vomberg, A., Enke, M., & Grimm, P. H. (2015). The loss of the marketing department's influence: Is it really happening? And who cares? *Journal of the Academy of Marketing Science*, 43(1), 1–13. <https://doi.org/10.1007/s11747-014-0416-3>
- IAB Europe. (2025). *AdEx Benchmark 2024 report*. IAB Europe.
- Jacobides, M. G., Cennamo, C., & Gawer, A. (2018). Towards a theory of ecosystems. *Strategic Management Journal*, 39(8), 2255–2276. <https://doi.org/10.1002/smj.2904>
- Järvinen, J., & Karjaluoto, H. (2015). The use of web analytics for digital marketing performance measurement. *Industrial Marketing Management*, 50, 117–127. <https://doi.org/10.1016/j.indmarman.2015.04.009>
- Jick, T. D. (1979). Mixing qualitative and quantitative methods: Triangulation in action. *Administrative Science Quarterly*, 24(4), 602–611. <https://doi.org/10.2307/2392366>
- Kalaignanam, K., Tuli, K. R., Kushwaha, T., Lee, L., & Gal, D. (2021). Marketing agility: The concept, antecedents, and a research agenda. *Journal of Marketing*, 85(1), 35–58. <https://doi.org/10.1177/0022242920952760>

- Katsikeas, C. S., Morgan, N. A., Leonidou, L. C., & Hult, G. T. M. (2016). Assessing performance outcomes in marketing. *Journal of Marketing*, 80(2), 1–20.
<https://doi.org/10.1509/jm.15.0287>
- Kohavi, R., Longbotham, R., Sommerfield, D., & Henne, R. M. (2009). Controlled experiments on the web: Survey and practical guide. *Data Mining and Knowledge Discovery*, 18(1), 140–181. <https://doi.org/10.1007/s10618-008-0114-1>
- Krippendorff, K. (2004). *Content analysis: An introduction to its methodology* (2nd ed.). SAGE.
- Kvale, S., & Brinkmann, S. (2009). *InterViews: Learning the craft of qualitative research interviewing* (2nd ed.). SAGE.
- Levitt, B., & March, J. G. (1988). Organizational learning. *Annual Review of Sociology*, 14, 319–340. <https://doi.org/10.1146/annurev.so.14.080188.001535>
- Lewis, R. A., & Rao, J. M. (2015). The unfavorable economics of measuring the returns to advertising. *The Quarterly Journal of Economics*, 130(4), 1941–1973.
<https://doi.org/10.1093/qje/qjv023>
- Lincoln, Y. S., & Guba, E. G. (1985). *Naturalistic inquiry*. SAGE.
- Liu, C. H. B., Chamberlain, B. P., & McCoy, E. J. (2020). What is the value of experimentation and measurement? *Data Science and Engineering*, 5(2), 152–167.
<https://doi.org/10.1007/s41019-020-00121-5>
- March, J. G., & Olsen, J. P. (1975). The uncertainty of the past: Organizational learning under ambiguity. *European Journal of Political Research*, 3(2), 147–171.
<https://doi.org/10.1111/j.1475-6765.1975.tb00521.x>
- McKinsey & Company. (2026). *Past forward: The modern rethinking of marketing's core (State of Marketing Europe 2026)* [Report]. McKinsey & Company.
- Micova, S. B., & Jacques, S. (2020). Platform power in the video advertising ecosystem. *Internet Policy Review*, 9(4). <https://doi.org/10.14763/2020.4.1506>
- Morgan, N. A., Jayachandran, S., Hulland, J., Kumar, B., Katsikeas, C., & Somosi, A. (2021). Marketing performance assessment and accountability: Process and outcomes. *International Journal of Research in Marketing*, 39(2), 462–481.
<https://doi.org/10.1016/j.ijresmar.2021.10.008>

- Nieborg, D. B., & Poell, T. (2025). Analyzing institutional platform power: Evolving relations of dependence in the mobile digital advertising ecosystem. *New Media & Society*, 27(4), 1909–1927. <https://doi.org/10.1177/1461444825131440>
- O’Sullivan, D., Abela, A. V., & Hutchinson, M. (2009). Marketing performance measurement and firm performance: Evidence from the European high-technology sector. *European Journal of Marketing*, 43(5–6), 843–862. <https://doi.org/10.1108/03090560910947070>
- O’Sullivan, D., & Butler, P. (2010). Marketing accountability and marketing’s stature: An examination of senior executives’ perceptions. *Australasian Marketing Journal*, 18(2), 113–119. <https://doi.org/10.1016/j.ausmj.2010.05.002>
- Pang, J., Lin, W., Fu, H., Kleeman, J., Bitar, E., & Wierman, A. (2022). Transparency and control in platforms for networked markets. *Operations Research*, 70(3), 1665–1690. <https://doi.org/10.1287/opre.2021.2244>
- Patton, M. Q. (2002). *Qualitative research & evaluation methods* (3rd ed.). SAGE.
- Rahman, H. A. (2021). The invisible cage: Workers’ reactivity to opaque algorithmic evaluations. *Administrative Science Quarterly*, 66(4), 945–988. <https://doi.org/10.1177/00018392211010118>
- Rogers, B. H. (2017). The marketing performance credibility gap: The growing importance of developing financially credible measures of the contribution of marketing to enterprise value, profits and growth. *Applied Marketing Analytics*, 3(3), 194–204. <https://doi.org/10.69554/xqvk4428>
- Rust, R. T., Ambler, T., Carpenter, G. S., Kumar, V., & Srivastava, R. K. (2004). Measuring marketing productivity: Current knowledge and future directions. *Journal of Marketing*, 68(4), 76–89. <https://doi.org/10.1509/jmkg.68.4.76.42721>
- Saunders, M., Lewis, P., & Thornhill, A. (2019). *Research methods for business students* (8th ed.). Pearson.
- Schmeiss, J., Hoelzle, K., & Tech, R. P. G. (2019). Designing governance mechanisms in platform ecosystems: Addressing the paradox of openness through blockchain technology. *California Management Review*, 62(1), 121–143. <https://doi.org/10.1177/0008125619883618>
- Schumacher, C., Keck, S., & Tang, W. (2020). Biased interpretation of performance feedback: The role of CEO overconfidence. *Strategic Management Journal*, 41(6), 1139–1165. <https://doi.org/10.1002/smj.3138>

- Shadish, W. R., Cook, T. D., & Campbell, D. T. (2002). *Experimental and quasi-experimental designs for generalized causal inference*. Houghton Mifflin.
- Staub, N., Haki, K., Aier, S., & Winter, R. (2022). Governance mechanisms in digital platform ecosystems: Addressing the generativity–control tension. *Communications of the Association for Information Systems*, 51, 906–939. <https://doi.org/10.17705/1CAIS.05137>
- Teece, D. J. (2007). Explicating dynamic capabilities: The nature and microfoundations of (sustainable) enterprise performance. *Strategic Management Journal*, 28(13), 1319–1350. <https://doi.org/10.1002/smj.640>
- Teece, D. J., Pisano, G., & Shuen, A. (1997). Dynamic capabilities and strategic management. *Strategic Management Journal*, 18(7), 509–533. [https://doi.org/10.1002/\(SICI\)1097-0266\(199708\)18:7%3C509::AID-SMJ882%3E3.0.CO;2-Z](https://doi.org/10.1002/(SICI)1097-0266(199708)18:7%3C509::AID-SMJ882%3E3.0.CO;2-Z)
- Tranfield, D., Denyer, D., & Smart, P. (2003). Towards a methodology for developing evidence-informed management knowledge by means of systematic review. *British Journal of Management*, 14(3), 207–222. <https://doi.org/10.1111/1467-8551.00375>
- van der Vlist, F. N., & Helmond, A. (2021). How partners mediate platform power: Mapping business and data partnerships in the social media ecosystem. *Big Data & Society*, 8(1). <https://doi.org/10.1177/20539517211025061>
- Vollrath, M. D., & Villegas, S. G. (2022). Avoiding digital marketing analytics myopia: Revisiting the customer decision journey as a strategic marketing framework. *Journal of Marketing Analytics*, 10(2), 106–113. <https://doi.org/10.1057/s41270-020-00098-0>
- Waisman, C., Sahni, N. S., Nair, H. S., & Lin, X. (2024). Parallel Experimentation and Competitive Interference on Online Advertising Platforms. *Marketing Science*, 44(2), 437–456. <https://doi.org/10.1287/mksc.2022.0194>
- Wichmann, J. R. K., Wiegand, N., & Reinartz, W. J. (2022). The platformization of brands. *Journal of Marketing*, 86(1), 109–131. <https://doi.org/10.1177/00222429211054073>
- Zollo, M., & Winter, S. G. (2002). Deliberate learning and the evolution of dynamic capabilities. *Organization Science*, 13(3), 339–351. <https://doi.org/10.1287/orsc.13.3.339.2780>

Appendix

Appendix A. Semi-Structured Interview Guide

The interview guide followed a semi-structured funnel logic, which is illustrated in Table 10. It began with role and decision-context questions, then moved into platform feedback, distorted learning, concrete decision episodes, evidence integration routines, capability conditions, and future outlook. Questions were adapted to each participant's industry and role while preserving the same core structure.

Table 10

Semi-structured interview guide

Section	Core Question / Probe	Objective
A: Opening & Research Context		
	Thank you for taking the time. This study examines how senior marketing leaders allocate and defend budgets when performance feedback is imperfect, especially in platform-mediated environments. No sensitive data is required. With permission, the interview is recorded for accuracy and anonymized.	Establishes informed consent and frames the interview around executive decision-making under imperfect evidence. Semi-structured expert interviews are appropriate when the goal is to understand decision processes and allow role-specific probing (Kvale & Brinkmann, 2009).
B: Role & Decision Context		
	Q1: Could you briefly describe your current role and responsibilities related to marketing or growth decision-making?	Establishes the respondent's relevance and decision authority. This supports purposive sampling of information-rich cases (Patton, 2002).
	Q2: Which budgeting or allocation decisions are you directly responsible for, and which do you influence?	Identifies the unit of analysis: executive decision processes in recurring allocation cycles. This links to resource allocation and dynamic-capabilities logic (Teece, 2007; Barreto, 2010).
	Q3: How is the marketing or growth budget usually determined? Is it fixed, adaptive, or performance-based?	Clarifies the allocation setting and whether learning from feedback can affect reallocation decisions. This connects to organizational learning from performance feedback (Greve, 2003; Argote, 2012).
C: Allocation Logic & Value Criteria		
	Q4: How do you decide where to allocate marketing or growth spend across channels, campaigns, or business priorities?	Elicits allocation criteria and the role of managerial judgment under uncertainty. This links to resource-based and dynamic-capabilities

	views of resource deployment (Barney et al., 2021; Teece, 2007).
Q5: What role does incrementality play when deciding whether additional spend creates value?	Connects reported performance to causal impact. This grounds the interview in the advertising-measurement problem (Lewis & Rao, 2015; Gordon et al., 2019; Gordon et al., 2022).
Q6: When you evaluate whether marketing is working, what counts as real value?	Elicits how respondents define marketing value and defensibility. This links to marketing accountability and perceived marketing value (Rust et al., 2004; O'Sullivan & Butler, 2010; Morgan et al., 2021).
D: Platform Data & Performance Feedback	
Q7: How are platform-reported metrics from channels such as Google, Meta, Amazon, TikTok, or similar platforms used in decision-making?	Assesses how platform-governed signals enter budget decisions. This links to platform governance and information control (Gawer, 2014; Cutolo & Kenney, 2021; Wichmann et al., 2022).
Q8: Where do you see the biggest gaps between platform-reported performance and actual business impact?	Identifies sources of PLD. This connects platform opacity with causal-measurement limits in digital advertising (Lewis & Rao, 2015; Gordon et al., 2019; Braun & Schwartz, 2024).
Q9: Which metrics do teams tend to overtrust, and which downstream signals do they underuse?	Probes biased interpretation and overweighting of visible signals. This links to organizational learning under noisy or ambiguous feedback (March & Olsen, 1975; Levitt & March, 1988; Greve, 2003).
E: False-Positive Case	
Q10: Can you walk me through a case where available performance data looked strong, but real business impact turned out to be weaker?	Captures concrete evidence of distorted learning and false-positive performance feedback. This links to attribution bias, non-incremental conversions, and misleading KPIs (Barajas et al., 2016; Gordon et al., 2022; Edelman, 2014).
Q11: What exactly looked strong at first? When did you start doubting the result?	Identifies the misleading signal and the trigger for belief revision. This connects to feedback interpretation and belief updating under ambiguity (March & Olsen, 1975; Schumacher et al., 2020; Albert, 2025).
Q12: Which non-platform signals contradicted the initial interpretation?	Identifies validation sources such as backend outcomes, sales, margin, retention, customer quality, or experiments. This links to Evidence

	Triangulation and causal validation (Gordon et al., 2019; Kohavi et al., 2009).
Q13: What decision did you take, and what would you do differently today?	Links distorted learning to allocation behavior and corrective action. This grounds the PLD → PVC → budget decision mechanism.
F: False-Negative Case	
Q14: Have you seen the opposite case, where platform data looked weak, but actual business impact was stronger than reported?	Captures under-recognized value creation, especially brand, trust, loyalty, customer quality, and long-term value. This links to marketing accountability and longer-term performance outcomes (Rust et al., 2004; Katsikeas et al., 2016; Morgan et al., 2021).
Q15: What made you believe the impact was real despite weaker platform signals?	Identifies evidence used to defend investments not fully visible in platform dashboards. This links to budget defense and perceived value creation under incomplete observability.
G: Evidence Integration & Decision Routines	
Q16: What practices help you avoid making wrong investment decisions under imperfect performance feedback?	Identifies routines that reduce distorted learning. This supports the derivation of Evidence Integration as a dynamic-capability microfoundation (Teece, 2007; Gioia, 2021).
Q17: How do you combine platform data with internal data, experiments, customer data, financial outcomes, or managerial judgment?	Directly probes Evidence Integration. This reflects the disciplined combination of heterogeneous evidence before allocation decisions.
Q18: What rules, thresholds, or review routines determine whether you scale, hold, or stop spending?	Captures disciplined reallocation routines and decision thresholds. This links to dynamic capabilities, marketing agility, and decision-process discipline (Teece, 2007; Kalaigianam et al., 2021; Homburg & Wielgos, 2022).
H: Capability & Governance Conditions	
Q19: Which internal capabilities matter most for better marketing allocation decisions?	Identifies capability conditions linked to GIC. This connects to marketing capabilities and capability-building under uncertainty (Day, 2011; Barreto, 2010; Homburg & Wielgos, 2022).
Q20: How do marketing, finance, data, product, sales, or technology teams interact when performance evidence is ambiguous?	Captures cross-functional evidence interpretation and governance. This links to marketing accountability and financial credibility

	(O’Sullivan & Butler, 2010; Rogers, 2017; Morgan et al., 2021).
Q21: When budget decisions are challenged internally, what counts as credible evidence?	Connects evidence quality to budget defense and PVC. This supports the thesis’s focus on perceived, defensible value creation.
I: Future Outlook	
Q22: Do you think this problem is becoming easier or harder for CMOs, and why?	Explores how platform opacity, privacy, automation, and data fragmentation affect future decision quality. This links to platform dependence and changing measurement environments (Cutolo & Kenney, 2021; Wichmann et al., 2022).
Q23: What role will AI or machine learning play in marketing measurement and allocation decisions?	Explores emerging capability conditions and the distinction between prediction and causal understanding. This supports the thesis’s future research logic.
J: Closing	
Q24: If you had to rely on one practice to avoid bad decisions in this environment, what would it be?	Forces synthesis around the most important routine or capability.
Q25: Is there anything I did not ask that you regard as important for understanding this problem?	Allows inductive themes and boundary conditions to emerge.
Q26: Are there additional experts you would recommend speaking to?	Supports referral-based expert sampling and additional information-rich cases.

Note. The table reports the semi-structured interview guide used in the qualitative phase. Questions were adapted to each participant’s role and industry while preserving the same core structure. PLD = Performance Learning Distortion; EI = Evidence Integration; PVC = Perceived CMO Value Creation; GIC = Growth Investment Capability.

Appendix B. Qualitative Coding Structure

Appendix B summarizes the qualitative coding structure used to derive the focal microfoundation and connect interview evidence to the thesis constructs. The coding followed a structured thematic logic. First-order informant terms were grouped into second-order themes and then abstracted into aggregate dimensions. This structure supported the selection of Evidence Integration as the focal dynamic-capability microfoundation. Evidence Integration captures the disciplined combination of platform-reported metrics, internal performance data, experimental evidence, and cross-functional interpretation before allocation

decisions. Table 11 reports the coding structure used to derive the focal microfoundation and link interview evidence to the thesis constructs.

Table 11

Qualitative coding structure and construct derivation

Representative first-order evidence	Second-order theme	Aggregate dimension	Link to thesis construct
“Platform data is useful for campaign steering, but not sufficient for true ROI.” / “Platform data helps you drive the car, but not decide whether you should be on that road.” / “The platform is optimizing a local environment; the firm is optimizing a business.”	Platform feedback as local and partial evidence	Platform-governed feedback creates distorted learning	Performance Learning Distortion (PLD)
“Each platform shows its own part of the story.” / “Every platform has its own attribution logic.” / “Platforms report what they can see, using their own logic.”	Platform-specific attribution logic	Platform-governed feedback creates distorted learning	PLD
“Platform data is not wrong, but incomplete.” / “The problem starts when leaders mistake local truth for whole-business truth.” / “The platform’s view is not the full system.”	Local truth mistaken for causal truth	Platform-governed feedback creates distorted learning	PLD
“Digital KPIs create a false sense of safety because they are easier to measure.” / “What is visible online can be overestimated.” / “Teams can confuse activity with progress.”	Measurability bias in managerial interpretation	Platform-governed feedback creates distorted learning	PLD
“Lead quantity was high, but lead quality was weak.” / “Meta lead ads generated volume, but conversion to customer contracts was lower.” / “Strong surface metrics did not translate into purchase quality.”	Visible performance can overstate true business impact	Platform-governed feedback creates distorted learning	PLD
“The dashboard looked healthy, but sales or margin did not follow.” / “Revenue did not increase in line with the number of leads.” / “Strong engagement did not translate into meaningful commercial effect.”	False-positive learning from platform signals	Platform-governed feedback creates distorted learning	PLD
“Some channels look efficient because they harvest demand created elsewhere.” / “Branded search and retargeting can be over-credited.” / “A final visible touchpoint can receive too much credit.”	Demand capture mistaken for demand creation	Platform-governed feedback creates distorted learning	PLD
“Marketing may take credit for demand caused by weather, COVID, seasonality, or market conditions.” /	External forces misread as marketing effect	Platform-governed feedback creates distorted learning	PLD

Representative first-order evidence	Second-order theme	Aggregate dimension	Link to thesis construct
“The campaign looked brilliant, but external demand was rising anyway.”			
“Digital investment alone does not drive sales if trade activation or shelf presence is weak.” / “A campaign can perform well online but fail if store execution is missing.” / “Retail reality changes the interpretation of digital metrics.”	Channel metrics detached from execution context	Platform-governed feedback creates distorted learning	PLD
“In marketplaces, a conversion is not only about marketing exposure.” / “Price, availability, wait time, supply, and incentives also drive outcomes.” / “Marketing is only one lever in the system.”	System effects beyond marketing attribution	Platform-governed feedback creates distorted learning	PLD
“Promotions can create strong short-term metrics while attracting low-quality or price-sensitive users.” / “You may pull demand forward rather than create durable value.” / “Good-looking conversion rates can hide weaker long-term economics.”	Short-term response mistaken for sustainable value	Platform-governed feedback creates distorted learning	PLD / PVC
“Clicks, opens, or traffic can look strong while the customer journey does not improve.” / “A headline can get attention without creating trust.” / “Traffic is not valuable if it does not support the next action.”	Proxy metrics decoupled from downstream value	Platform-governed feedback creates distorted learning	PLD / PVC
“Some valuable effects do not show immediately in platform data.” / “Brand, loyalty, trust, and customer experience reveal value over time.” / “Weak short-term KPIs can hide long-term value.”	False-negative learning from narrow measurement windows	Platform-governed feedback creates distorted learning	PLD / PVC
“Brand initiatives can look inefficient in short-term ROI terms but improve the whole funnel.” / “TV can look weak alone but make Google more efficient.” / “Heritage, storytelling, and loyalty strengthen future interactions.”	Under-recognized indirect and long-term effects	Platform-governed feedback creates distorted learning	PLD / PVC
“The true measurement should happen in the backend.” / “You need to compare platform data with internal outcomes.” / “Downstream signals are more trustworthy than one isolated KPI.”	Backend validation and downstream outcome checking	Evidence Integration protects decision quality	Evidence Integration (EI)
“Platform data is essential for operating campaigns, but measurement should happen outside the platform.” / “Use platform data for execution, backend data for ROI.”	Separation of operational and strategic evidence	Evidence Integration protects decision quality	EI

Representative first-order evidence	Second-order theme	Aggregate dimension	Link to thesis construct
“We build a hierarchy of data points.” / “Some evidence is more reliable but less frequent; other evidence is less reliable but more available.” / “The task is to combine imperfect but useful signals.”	Evidence hierarchy	Evidence Integration protects decision quality	EI
“Geo-tests, holdouts, cohort analysis, and LTV help challenge platform reporting.” / “Testing reduces uncertainty but does not eliminate it.” / “Experiments create better evidence when dashboards are ambiguous.”	Experimental validation and causal checking	Evidence Integration protects decision quality	EI
“Good decisions require looking at the whole funnel.” / “Clicks, leads, configurations, sales, retention, and margin need to tell a coherent story.” / “No single metric is enough.”	Funnel-level and cross-signal interpretation	Evidence Integration protects decision quality	EI
“If data is not available, companies should test in order to generate it.” / “Good tests need a success metric, decision owner, and review date.” / “Without clear design, a test is just activity.”	Testing discipline	Evidence Integration protects decision quality	EI / GIC
“Platform signals need to be compared with loyalty data, transactions, repeat behavior, basket evolution, and customer value.” / “First-party data reduces ambiguity.”	First-party validation of platform signals	Evidence Integration protects decision quality	EI / GIC
“Power BI combines different data sources into one decision dashboard.” / “Dashboards help only when multiple sources are merged and interpreted together.”	Integrated decision dashboards	Evidence Integration protects decision quality	EI / GIC
“Strong decisions combine data, testing, context, and strategy without overreacting to noise.” / “The role of the CMO is not to eliminate uncertainty, but to navigate it.”	Disciplined interpretation under uncertainty	Evidence Integration protects decision quality	EI / GIC
“Good leaders ask: what else could explain this?” / “They do not allow easy narratives to win too quickly.” / “They separate correlation from causation as much as possible.”	Challenge routines and alternative explanations	Evidence Integration protects decision quality	EI / GIC
“We look at quantity and quality, not only lead volume.” / “Configurations matter because they correlate most with sales.” / “Customer quality, margin dilution, and repeat behavior matter more than activity volume.”	Signal quality over signal quantity	Evidence Integration protects decision quality	EI / PVC
“Marketing, finance, data, product, and technology need to work from a shared evidence base.” / “The CMO must	Cross-functional interpretation of evidence	GIC operates as evidence architecture	Growth Investment Capability (GIC)

Representative first-order evidence	Second-order theme	Aggregate dimension	Link to thesis construct
speak enterprise value, not only marketing metrics.” / “Finance and technology alignment are essential.” “You cannot revolutionize marketing from inside marketing alone.” / “Marketing effectiveness is an enterprise issue, not a functional one.” / “Brand, analytics, loyalty, digital, finance, and commercial teams must align.”	Marketing as enterprise decision process	GIC operates as evidence architecture	GIC
“Strong first-party data and event streams are critical.” / “Data must move across the enterprise.” / “The starting point is building the systems so data can move.”	Evidence infrastructure and signal quality	GIC operates as evidence architecture	GIC
“Better decisions require CRM integration, event collection, customer-level data, and backend signal quality.” / “Many companies have data, but not shared confidence in what it means.”	Infrastructure for trusted evidence	GIC operates as evidence architecture	GIC
“Scale only when performance is stable and downstream quality holds.” / “Weekly data should not drive structural decisions.” / “Strategic decisions need a different cadence from campaign optimization.”	Disciplined reallocation cadence	GIC operates as evidence architecture	GIC
“A lot of organizations are good at launching initiatives and bad at killing them.” / “The willingness to stop is part of budget quality.”	Termination discipline	GIC operates as evidence architecture	GIC
“Prediction can reduce waste when tied to the right problem.” / “Machine learning is useless unless you put the right problem into it.”	Prediction tied to decision problem	GIC operates as evidence architecture	GIC / EI
“AI can improve execution, but interpretation becomes harder.” / “The mechanics get easier while judgment gets harder.” / “Machine confidence is not business truth.”	Prediction–causality distinction	GIC operates as evidence architecture	GIC / EI
“If the wrong question is asked, AI amplifies confusion.” / “Advanced analytics are not a substitute for strategic clarity.”	Automation requires stronger evidence governance	GIC operates as evidence architecture	GIC
“The strongest capability is connecting evidence to decision quality.” / “Not just collecting data, but turning partial evidence into better capital allocation decisions.”	Repeatable decision capability	GIC operates as evidence architecture	GIC
“The deeper issue is how organizations learn under partial information.” / “Marketing is where this becomes visible quickly.”	Organizational learning under partial observability	GIC operates as evidence architecture	GIC / PLD

Representative first-order evidence	Second-order theme	Aggregate dimension	Link to thesis construct
“More data, more tools, and more automation also create more noise.” / “The challenge shifts from access to data toward making sense of it.”	Complexity and signal interpretation	GIC operates as evidence architecture	GIC / EI
“Explicit thresholds matter: acceptable CAC, payback, pipeline quality, and conversion slippage.” / “If downstream conversion deteriorates, stop and investigate.”	Evidence thresholds for scaling and stopping	GIC operates as evidence architecture	GIC / EI
“Budget should be the output of what makes economic sense.” / “The better question is how to avoid spending one euro that should not be spent.” / “Budget is a consequence, not a starting point.”	Budget as economic output, not fixed entitlement	GIC operates as evidence architecture	GIC / PVC
“Long-term does not mean vague.” / “If an investment will not reveal itself in one quarter, define when it will be reviewed and what would cause it to stop.”	Structured review of long-term investments	GIC operates as evidence architecture	GIC / PVC
“Real value is profit, customer quality, LTV, retention, pricing power, or enterprise contribution.” / “Marketing value must be defensible to finance, CEOs, boards, franchisees, or investors.”	Value as defensible enterprise contribution	Value creation is defended as enterprise contribution	Perceived CMO Value Creation (PVC)
“If we were not driving customer value, margin return, cash flow, and reducing unnecessary cost, we were not effective.” / “LTV is not only a marketing formula; it affects the value of the business.”	CMO value as enterprise economics	Value creation is defended as enterprise contribution	PVC
“Customer lifetime value should be based on profit, not only revenue.” / “The core logic is incremental CAC versus predicted LTV.”	Profit-based customer value logic	Value creation is defended as enterprise contribution	PVC
“Customer quality matters more than cheap acquisition.” / “The question is not how cheaply we acquired users, but what quality of relationship we created.”	Quality of growth over volume of activity	Value creation is defended as enterprise contribution	PVC
“Real value is when customer value and business value move in the same direction.” / “Short-term business gains that weaken the customer relationship are fragile.”	Dual customer-business value logic	Value creation is defended as enterprise contribution	PVC
“In luxury, visibility and desirability are not the same.” / “Too much reach can dilute perceived value.” / “Data must serve the brand, not the other way around.”	Brand meaning as value constraint	Value creation is defended as enterprise contribution	PVC

Representative first-order evidence	Second-order theme	Aggregate dimension	Link to thesis construct
“Strong platform efficiency can weaken long-term brand strength.” / “Performance marketing can capture demand while starving demand creation.”	Short-term efficiency versus long-term value	Value creation is defended as enterprise contribution	PVC
“Content and design quality build trust, but this may not appear immediately in platform metrics.” / “Clear communication can improve the journey even without spectacular short-term results.”	Trust and journey quality as less visible value	Value creation is defended as enterprise contribution	PVC
“Retail media turns attention and first-party data into assets.” / “The value of customer attention and intelligence becomes part of the business model.”	Customer attention as strategic asset	Value creation is defended as enterprise contribution	PVC / GIC

Across interviews, Evidence Integration emerged as the most recurring and actionable routine. Interviewees did not describe one universally superior metric or model. Instead, they repeatedly emphasized the need to combine platform signals with backend data, downstream commercial outcomes, experiments, customer-quality indicators, and cross-functional judgment. Evidence Integration was therefore selected as the focal microfoundation because it was salient across interviews, closely linked to distorted learning, actionable as a managerial routine, and suitable for experimental operationalization through an evidence-triangulation manipulation. Table 12 summarizes the selection logic for Evidence Integration as the focal microfoundation.

Table 12

Selection logic for evidence integration as focal microfoundation

Selection criterion	Evidence Integration fit
Salience in executive accounts	Recurring across interviews as backend validation, triangulation, funnel interpretation, testing, and cross-functional evidence review.
Causal proximity to PLD	Directly targets the mechanism through which platform signals become over-weighted or corrected before allocation decisions.
Managerial actionability	Can be translated into concrete practices, including evidence standards, validation routines, testing, and review cadence.
Experimental operationalizability	Can be represented in a standardized vignette by comparing platform-only feedback with triangulated evidence.

Selection criterion	Evidence Integration fit
Link to dynamic capabilities	Functions as a micro-level routine that improves sensing, interpretation, and reallocation under uncertainty.

Table 13 documents how individual interviews contributed to the qualitative analysis. The matrix does not reproduce full transcripts. Instead, it summarizes the main construct-relevant evidence from each interview and links it to Performance Learning Distortion (PLD), Evidence Integration (EI), Perceived CMO Value Creation (PVC), and Growth Investment Capability (GIC). This structure preserves an audit trail while keeping the appendix focused and readable.

Table 13

Coded expert interview matrix

Expert	Context	Main evidence contribution	Construct link
E01	Enterprise CMO / global QSR and loyalty ecosystem	Shown that marketing value depends on enterprise evidence architecture. Platform signals could be over-attributed when external demand, timing, offer effects, or social-platform disruption were ignored. The interview also emphasized loyalty data, transaction data, finance alignment, and machine-learning-supported allocation.	PLD / EI / PVC / GIC
E02	CMO / CRO in subscription D2C and e-commerce	Provided the clearest evidence for backend validation and evidence hierarchy. Platform data was useful for campaign steering, but true ROI required internal measurement, incrementality testing, cohort analysis, predicted LTV, and backend-platform comparison.	PLD / EI / PVC / GIC
E03	CEO / ex-CMO in buy-and-build and home-services platform contexts	Highlighted channel interdependence and false precision. Platform dashboards could isolate effects that were not isolated in the business system. Practical testing, geo splits, time-interval pilots, and backend checks were treated as more useful than excessive attribution modeling.	PLD / EI / PVC / GIC
E04	Managing director / FMCG and omnichannel retail	Shown that digital performance can be misread when detached from trade activation, distribution, shelf presence, and store execution. Digital KPIs required validation against retail sell-out signals, trade feedback, and customer behavior.	PLD / EI / PVC / GIC
E05	Automotive CMO	Highlighted partial visibility between media spend, funnel activity, dealer signals, and final sales. Configurator completions and downstream sales indicators were more meaningful than generic traffic or lead volume.	PLD / EI / PVC / GIC

Expert Context		Main evidence contribution	Construct link
E06	Consumer brand CMO / FMCG	Showed that platform metrics can miss awareness, trial, brand building, distribution effects, and retail conversion. This supported the false-negative side of PLD and showed why broader evidence is needed to evaluate long-term value creation.	PLD / EI / PVC / GIC
E07	Luxury automotive CMO	Provided a contrast case where short-term platform efficiency could conflict with desirability, scarcity, pricing power, and brand equity. Value creation depended on protecting long-term brand meaning rather than maximizing visible reach or short-term activation.	PLD / EI / PVC / GIC
E08	Media-commerce / platform growth executive	Showed that traffic volume and engagement can be mistaken for valuable attention. Decision quality depended on attention quality, monetization evidence, portfolio-level traffic comparisons, and interpretation of downstream user behavior.	PLD / EI / PVC / GIC
E09	Marketplace / mobility growth executive	Reframed allocation as a system-equilibrium problem. Marketing outcomes depended on pricing, supply, demand, incentives, availability, and operational conditions, not advertising exposure alone. This extended PLD beyond platform dashboards into cross-lever interdependence.	PLD / EI / PVC / GIC
E10	Content / CRM / customer-journey specialist	Added evidence on trust, clarity, comprehension, and journey quality as less visible value drivers. Content and CRM metrics could show engagement without proving that customers understood the offer or moved confidently through the funnel.	PLD / EI / PVC / GIC
E11	B2B fintech / scale-up CMO	Showed that lead volume and campaign efficiency can overstate pipeline value. Stronger decisions required sales feedback, funnel conversion, CAC/payback logic, pipeline qualification, and finance-aligned review.	PLD / EI / PVC / GIC
E12	Global marketing / customer transformation leader	Emphasized that platform data is useful for action but insufficient for interpretation. Evidence quality depended on first-party data, customer behavior, financial outcomes, customer transformation logic, and cross-functional review.	PLD / EI / PVC / GIC
E13	Retail CMO / CTO	Strengthened the GIC construct by linking marketing judgment to technology ownership, loyalty systems, first-party transaction data, retail-media evidence, and commercial decision routines. The interview showed why evidence architecture requires marketing and technology integration.	EI / PVC / GIC
E14	CCO / CTO in hospitality and loyalty ecosystem	Showed that digital performance metrics can miss guest experience, loyalty effects, repeat behavior, and customer lifetime value. Evidence had to connect digital systems, booking behavior, loyalty data, and commercial outcomes.	PLD / EI / PVC / GIC

Expert Context		Main evidence contribution	Construct link
E15	Retail CMO / discount retail	Highlighted execution infrastructure, store experience, content consistency, and customer value perception as necessary checks on campaign performance. Platform evidence had to be interpreted in the context of retail execution and commercial outcomes.	PLD / EI / PVC / GIC
E16	Platform / adtech CMO	Provided the clearest articulation of local platform truth versus whole-business truth. Stronger evidence standards, interpretation routines, and governance of automation were needed to challenge platform claims and automated outputs.	PLD / EI / PVC / GIC
E17	Senior marketing strategy advisor	Provided cross-company pattern recognition. Stronger firms triangulated evidence, challenged assumptions, aligned marketing interpretation with sales and finance realities, and institutionalized learning across decision cycles.	PLD / EI / PVC / GIC

Note. The matrix summarizes construct-relevant evidence from the expert interviews and does not reproduce complete transcripts. “Evidence on PLD” captures distorted learning, attribution ambiguity, proxy metrics, or divergence between visible performance and business impact. “Evidence on Evidence Integration” captures routines that combine platform data with backend, experimental, financial, customer, or cross-functional evidence. “Evidence on PVC” captures how interviewees defined defensible CMO value creation. “Evidence on GIC” captures organizational conditions that improve evidence quality and disciplined reallocation.

Appendix C. Quantitative Survey Instrument and Experimental Vignette

Appendix C documents the quantitative survey instrument used in the experimental vignette study. The survey followed a sequential structure: consent, screening, role-taking instruction, campaign context, prior belief assessment, platform performance feedback, experimental condition, belief update, perceived CMO value creation, budget reallocation, manipulation checks, attention check, realism check, and demographics. Table 14 documents the survey structure, stimulus sequence, constructs, and response formats.

Table 14

Survey structure and objectives

Section	Question / Stimulus	Purpose / Construct	Question type / Response format
Consent	You are invited to participate in a research study for my Master Thesis on how marketing decisions are made in digital environments. The study will take approximately 8–10 minutes. Your responses are anonymous and will be used for academic research purposes only. Participation is voluntary, and you may	Obtain informed consent and confirm voluntary participation.	Multiple choice: I agree to participate; I do NOT agree to participate.

Section	Question / Stimulus	Purpose / Construct	Question type / Response format
	stop at any time. By proceeding, you confirm that you understand the above and agree to participate.		
Screening	How familiar are you with digital marketing or performance metrics, e.g., Cost Per Acquisition, conversion rates?	Screen for basic familiarity with digital performance marketing concepts.	Multiple choice: Not familiar at all; Slightly familiar; Moderately familiar; Very familiar; Extremely familiar.
Role-taking instruction	Please answer the following scenario as if you were the Chief Marketing Officer responsible for growth investment decisions. Assume that you are accountable for allocating marketing budget and defending those decisions based on available performance data. More information follows on the next page.	Induce role-taking and align respondent judgment with the CMO decision context.	Text stimulus.
Company context	You are the Chief Marketing Officer of VELOX, a mid-sized direct-to-consumer brand selling premium bicycle accessories online, including helmets, lights, and repair kits. The company generates approximately €8 million in annual revenue and relies heavily on digital channels for customer acquisition. In recent months, growth targets have increased, while profitability expectations remain strict. As a result, marketing investments are closely monitored. Your decision on whether to adjust the marketing budget for the next 4-week cycle will directly affect short-term growth and will be reviewed by senior management.	Establish decision context, business setting, platform dependence, and budget accountability.	Text stimulus + Company Logo.
Prior belief information	Before reviewing the detailed campaign results, you want to form an initial expectation. In digital advertising, not all platform-attributed purchases reported by Meta are truly caused by the campaign. Some would have happened anyway, for example through organic demand, existing brand awareness, or other channels. Without seeing the specific results yet, please estimate what you would generally expect in a situation like this.	Prepare prior belief elicitation and introduce incrementality logic.	Text stimulus.
Prior incrementality belief	Out of 100 purchases reported by a digital advertising platform, how many do you think are typically truly	Measure prior belief about incrementality	Numeric entry: 0–100.

Section	Question / Stimulus	Purpose / Construct	Question type / Response format
	caused by the campaign, i.e., would not have happened without it?	before detailed feedback.	
Prior confidence	How confident are you in your estimate?	Capture confidence in prior belief.	Multiple choice: Not confident at all; Slightly confident; Moderately confident; Very confident; Extremely confident.
Initial budget decision	Based on the general business context and your estimate of how many platform-attributed purchases are typically truly incremental, what would you initially do with the marketing budget for the next 4-week cycle relative to the current €120,000 spend?	Capture initial budget intention before detailed platform feedback and experimental manipulation.	Multiple choice: Increase spend by over 20%; Increase spend by up to 20%; Maintain spend; Decrease spend by up to 20%; Decrease spend by over 20%.
Platform performance instruction	Now, please carefully review the detailed performance report from the advertising platform. This is very important for the experiment.	Focus attention before platform feedback.	Text stimulus.
Platform performance overview	Meta Ads Campaign Performance Overview (Last 4 weeks). Campaign objective: Acquire new customers, measured via purchases. Channel: Meta Ads, including Facebook, Instagram, Messenger, and Audience Network. Spend: €120,000. Optimization goal: Purchases.	Establish platform campaign context and common information across all conditions.	Text stimulus.
Platform-reported results	Impressions: 4.2 million. Click-through rate: 2.8%. Cost per click: €0.94. Platform-attributed purchases: 3,050. Cost per attributed purchase: €39.34. Reported ROAS: 3.8x. Performance vs. prior period: +22% attributed conversions.	Provide strong platform-reported campaign performance held constant across conditions.	Text stimulus.
Context notes	Attribution is based on Meta's reporting system. Full cross-channel and offline effects are only partially observable. Recent tracking changes have reduced visibility into true incremental impact.	Hold platform opacity and feedback uncertainty constant across conditions.	Text stimulus.
Condition A: Platform-only evidence / Low GIC	Important: Before making your final decision, no additional validation or analysis is available. Decisions in this context are typically made based on platform-reported performance metrics, without	Experimental condition: Evidence Triangulation absent; GIC low.	Text stimulus.

Section	Question / Stimulus	Purpose / Construct	Question type / Response format
	structured cross-functional review or additional evidence.		
Condition B: Triangulated evidence / Low GIC	Important: Before making your final decision, you review additional internal and analytical evidence beyond the platform-reported results. This includes backend conversion data, a geo-based test indicating a lower incremental impact than suggested by the platform, and initial indications that newly acquired customers show weaker repeat purchase behavior than expected. However, decisions in this context are typically made without formal validation standards or structured cross-functional review.	Experimental condition: Evidence Triangulation present; GIC low.	Text stimulus.
Condition C: Platform-only evidence / High GIC	Important: Before making your final decision, no additional data beyond the platform-reported results is available. However, in this organization, investment decisions are typically subject to structured validation processes. Marketing performance is regularly reviewed together with finance and analytics teams, and decisions are expected to be supported by strong evidence and clear justification. As a result, weak or uncertain evidence is likely to be challenged before budget adjustments are approved.	Experimental condition: Evidence Triangulation absent; GIC high.	Text stimulus.
Condition D: Triangulated evidence / High GIC	Important: Before making your final decision, you review additional internal and analytical evidence beyond the platform-reported results. This includes backend conversion data, a geo-based test indicating a lower incremental impact than suggested by the platform, and initial indications that newly acquired customers show weaker repeat purchase behavior than expected. In this organization, investment decisions are also subject to structured validation processes. Marketing performance is regularly reviewed together with finance and analytics teams, and decisions are expected to be supported by strong evidence and clear justification. Weak or uncertain evidence is likely to be challenged before budget adjustments are approved.	Experimental condition: Evidence Triangulation present; GIC high.	Text stimulus.

Section	Question / Stimulus	Purpose / Construct	Question type / Response format
Post-feedback incrementality belief	After reviewing all available information, out of 100 purchases reported by Meta, how many do you believe were truly caused by the campaign, i.e., would not have happened without it?	Measure updated incrementality belief after condition exposure; used to construct PLD.	Numeric entry: 0–100.
Post-feedback confidence	How confident are you now in this estimate?	Capture confidence in updated belief.	Multiple choice: Not confident at all; Slightly confident; Moderately confident; Very confident; Extremely confident.
PVC instruction	Please indicate the extent to which you agree with the following statements regarding the decision.	Introduce perceived CMO value creation items.	Text stimulus.
PVC1	This investment creates value that justifies its cost.	Measure perceived value-for-money and defensibility.	7-point Likert scale: Strongly disagree to Strongly agree.
PVC2	The results provide a strong justification for the budget decision.	Measure perceived budget defensibility.	7-point Likert scale: Strongly disagree to Strongly agree.
PVC3	I would be able to defend this decision internally.	Measure perceived internal defensibility under scrutiny.	7-point Likert scale: Strongly disagree to Strongly agree.
PVC4	The marketing spend generates meaningful business impact.	Measure perceived business impact.	7-point Likert scale: Strongly disagree to Strongly agree.
Realism check	How realistic did you find the scenario?	Assess perceived realism of the vignette.	Multiple choice: Not at all realistic; Slightly realistic; Moderately realistic; Realistic; Very realistic.
Budget reallocation decision	For the next 4-week cycle, relative to the current €120,000 spend, what would you do with the budget?	Measure final budget allocation decision.	Multiple choice: Increase spend by over 20%; Increase spend by up to 20%; Maintain spend; Decrease spend by up to 20%; Decrease spend by over 20%.

Section	Question / Stimulus	Purpose / Construct	Question type / Response format
Attention check	This is a control question. Please select “Agree.”	Identify inattentive responses.	Multiple choice: Agree; Neutral; Disagree.
Qualitative decision reason	What was the main factor influencing your decision?	Capture open-ended reasoning for optional qualitative validation.	Open text.
Evidence Integration instruction	Please indicate the extent to which you agree with the following statements regarding the decision.	Introduce Evidence Integration manipulation-check items.	Text stimulus.
EI1	My decision process incorporated multiple sources of evidence.	Manipulation check for Evidence Integration / Evidence Triangulation.	7-point Likert scale: Strongly disagree to Strongly agree.
EI2	My decision was based on more than platform-reported metrics.	Manipulation check for use of non-platform evidence.	7-point Likert scale: Strongly disagree to Strongly agree.
EI3	Non-platform data was taken into account in the decision-making process.	Manipulation check for triangulated evidence use.	7-point Likert scale: Strongly disagree to Strongly agree.
GIC instruction	Please indicate the extent to which you agree with the following statements regarding the decision.	Introduce GIC manipulation-check items.	Text stimulus.
GIC1	Decisions in this context are subject to structured validation processes.	Manipulation check for perceived validation routines.	7-point Likert scale: Strongly disagree to Strongly agree.
GIC2	Weak or uncertain evidence would likely be challenged.	Manipulation check for challenge routines and evidence scrutiny.	7-point Likert scale: Strongly disagree to Strongly agree.
GIC3	Decisions are expected to be supported by strong evidence and clear justification.	Manipulation check for evidence standards.	7-point Likert scale: Strongly disagree to Strongly agree.
GIC4	Marketing performance is reviewed in a structured and disciplined way.	Manipulation check for disciplined review routines.	7-point Likert scale: Strongly disagree to Strongly agree.

Section	Question / Stimulus	Purpose / Construct	Question type / Response format
Gender	What is your gender?	Demographic control and sample description.	Multiple choice: Male; Female; Non-binary / third gender; Prefer not to say.
Age	What is your age?	Demographic control and sample description.	Numeric entry.
Professional background	Which of the following best describes your current professional background?	Sample description and relevance check.	Multiple choice: Marketing / Growth; Strategy / Consulting; Finance / Analytics; General Management; Student; Other.
Marketing experience	How many years of experience do you have in marketing, growth, or related decision-making roles?	Sample description and balance check.	Multiple choice: 0; 0–2 years; 3–5 years; 6–10 years; Above 10 years.
End of Survey	Thank you for taking the time to answer this survey. Your response has been saved.	End of Survey	

Note. Respondents were randomly assigned to one of four experimental conditions. Evidence Triangulation was coded 0 = absent and 1 = present. GIC condition was coded 0 = low and 1 = high. PVC items were averaged into a composite score after reliability checks. PVC captures the perceived value and defensibility of the focal investment, not verified causal impact. Evidence Integration and GIC items were used as manipulation checks. Budget reallocation was coded so that higher values indicate more expansionary budget decisions. PLD was constructed from the post-feedback incrementality belief relative to the prior belief.

Appendix D. Scale Construction

Multi-item constructs were created by averaging their respective items after reliability checks. Evidence Integration, GIC, and PVC showed high internal consistency and were therefore treated as composite mean scores in the analysis. Table 15 summarizes the construction of the main scales and behavioral variables.

Table 15

Scale construction summary

Construct	Items included	Scale construction	Reliability
Evidence Integration manipulation check	EI1–EI3	Mean score	$\alpha = .935$

Construct	Items included	Scale construction	Reliability
Growth Investment Capability manipulation check	GIC1–GIC4	Mean score	$\alpha = .927$
Perceived CMO Value Creation	PVC1–PVC4	Mean score	$\alpha = .948$
PLD	Prior and post-feedback incrementality belief	Behavioral belief-updating score	Not applicable
Budget Change	Budget allocation item	Ordinal-coded allocation decision	Not applicable

Appendix E. Coding Notes

Higher values on Evidence Integration indicate stronger perceived use of multiple evidence sources. Higher values on GIC indicate a stronger perceived organizational capability context. Higher values on PVC indicate stronger perceived value and defensibility of the focal platform-mediated investment. Higher values on Budget Change indicate more expansionary budget decisions. Positive PLD values indicate upward belief updating from prior to post-feedback belief; negative PLD values indicate downward correction.

Appendix F. Data Cleaning and Variable Construction

Survey data were collected in Qualtrics and exported to SPSS. Cleaning removed incomplete responses first, followed by quality screening for implausibly short completion times and inconsistent response patterns. The final analytical sample comprised 209 valid responses.

Data Cleaning Procedure

The initial Qualtrics export contained 316 responses. Data cleaning followed a documented sequence in SPSS. First, incomplete responses were removed. The remaining cases were screened for attention-check failure, implausibly short completion times, and inconsistent response patterns. No respondent failed the attention check. After data cleaning, the final analytical sample comprised 209 valid responses. Table 16 reports the data-cleaning funnel from the initial export to the final analytical sample.

Table 16

Data-cleaning funnel

Step	Action	Remaining N
Initial export	Raw Qualtrics responses exported to SPSS	316
Incomplete responses removed	Cases without sufficient survey completion were excluded	227
Quality screening applied	Remaining cases screened for attention-check failure, implausibly short completion time, and inconsistent response patterns	209
Final analytical sample	Cases used for hypothesis testing	209

Note. The cleaning sequence removed incomplete responses first, then screened for implausibly short completion times and inconsistent response patterns. No respondent failed the attention check.

Experimental Condition Variables

The experiment used a between-subjects 2×2 design with two manipulated factors: Evidence Triangulation and Growth Investment Capability (GIC). Evidence Triangulation operationalized Evidence Integration in the vignette. GIC was manipulated as an organizational governance context. Table 17 summarizes the coding of the experimental condition variables.

Table 17

Experimental condition coding

Variable	Coding	Interpretation
Evidence Triangulation (ET)	0 = absent; 1 = present	Indicates whether respondents received platform-only feedback or broader triangulated evidence.
Growth Investment Capability (GIC condition)	0 = low; 1 = high	Indicates whether respondents were assigned to a weaker or stronger organizational capability context.
Experimental condition	1–4	Combined 2×2 condition variable used for condition-level comparisons.

The final sample was balanced across conditions: Condition 1 included 51 respondents, Condition 2 included 54 respondents, Condition 3 included 53 respondents, and Condition 4 included 51 respondents.

Construct and Scale Construction

Multi-item constructs were created as mean scores after reliability checks. Higher scores indicate stronger perceived Evidence Integration, stronger perceived GIC, and stronger Perceived CMO Value Creation. PLD and Budget Change were behavioral variables derived from belief updating and allocation choices, respectively. Table 18 summarizes the construction and interpretation of the main analytical variables.

Table 18

Construct and scale construction

Construct	Construction	Interpretation	Reliability
Evidence Integration manipulation check	Mean of Evidence Integration items	Higher values indicate stronger perceived use of multiple evidence sources.	$\alpha = .935$
GIC manipulation check	Mean of GIC perception items	Higher values indicate stronger perceived organizational capability context.	$\alpha = .927$
Perceived CMO Value Creation (PVC)	Mean of PVC items	Higher values indicate stronger perceived incremental value and defensibility of the focal platform-mediated investment.	$\alpha = .948$
PLD	Post-feedback incrementality belief minus prior incrementality belief	Higher values indicate more upward belief updating after the performance evidence; lower values indicate more downward correction.	n/a
Budget Change	Ordinal-coded budget allocation item	Higher values indicate more expansionary budget decisions.	n/a

Performance Learning Distortion Variable

Performance Learning Distortion (PLD) was constructed as a behavioral measure based on respondents' incrementality beliefs in the vignette. Respondents estimated the campaign's incremental impact before and after reviewing the performance evidence. PLD captures the direction and magnitude of belief updating after exposure to the assigned performance evidence. Table 19 reports how PLD was constructed from prior and post-feedback incrementality beliefs.

Table 19

PLD construction

Variable	Construction logic	Interpretation
Prior incrementality belief	Respondent's estimate before detailed feedback	Initial belief about campaign incrementality
Post-feedback incrementality belief	Respondent's estimate after assigned evidence condition	Updated belief after reviewing performance evidence
PLD	Post-feedback belief minus prior incrementality belief	Positive values indicate upward belief updating; negative values indicate downward correction

Budget Change Variable

Budget Change captured respondents' allocation decision after reviewing the assigned vignette condition. The item was coded so that higher values indicated more expansionary budget decisions. Table 20 reports the coding logic for the Budget Change variable.

Table 20

Budget change coding

Variable	Construction logic	Coding / response format	Interpretation
Prior budget decision	Initial budget decision before detailed platform feedback and experimental condition exposure	1 = Reduce budget by more than 20%; 2 = Reduce budget by up to 20%; 3 = Keep budget unchanged; 4 = Increase budget by up to 20%; 5 = Increase budget by more than 20%	Baseline allocation intention before detailed performance evidence
Post-feedback budget decision	Final budget decision after reviewing the assigned evidence condition	1 = Reduce budget by more than 20%; 2 = Reduce budget by up to 20%; 3 = Keep budget unchanged; 4 = Increase budget by up to 20%; 5 = Increase budget by more than 20%	Updated allocation decision after exposure to performance evidence
Budget Change	Post-feedback budget decision minus prior budget decision	Difference score ranging from negative to positive values	Positive values indicate a more expansionary adjustment after the performance evidence; negative values indicate downward budget correction; zero indicates no change from the prior decision

Budget Change captured the direction and magnitude of respondents' allocation adjustment after reviewing the assigned evidence condition. It was calculated as the post-feedback budget

decision minus the prior budget decision. Higher values indicate more expansionary adjustment, while lower values indicate more conservative budget correction.

Additional variables were retained for descriptive statistics, balance checks, and robustness analyses. This included age, gender, marketing experience, and familiarity with digital performance metrics or growth decision contexts. Table 21 summarizes the additional variables retained for descriptive, balance, and robustness analyses.

Table 21

Additional analytical variables

Variable group	Purpose in analysis
Demographic variables	Used to describe the sample and assess condition balance.
Marketing experience	Used to assess whether experience differed across experimental conditions.
Digital marketing familiarity	Used to ensure respondents could interpret the vignette context.
Realism and comprehension checks	Used to assess whether respondents understood and accepted the decision scenario.

Statistical Analysis Workflow

All quantitative analyses were conducted in SPSS using the cleaned analytical dataset. The analysis proceeded in four steps. First, descriptive statistics and condition-balance checks were conducted. Second, reliability checks were computed for multi-item constructs. Third, manipulation checks tested whether the Evidence Triangulation and GIC manipulations worked as intended. Fourth, hypotheses were tested using ANOVA, OLS regression, and PROCESS models with 5,000 bootstrap samples. Table 22 summarizes the statistical analysis workflow.

Table 22

Statistical analysis workflow

Step	Analysis	Purpose
1	Descriptive statistics and condition balance checks	Assess sample structure and random assignment quality.
2	Reliability analysis	Evaluate internal consistency of EI, GIC, and PVC scales.
3	Manipulation checks	Test whether ET and GIC were perceived as intended.
4	OLS regressions	Test direct effects of ET on PVC for H1.

Step	Analysis	Purpose
5	PROCESS Model 4	Test ET → PLD, PLD → PVC, and the indirect effect for H2–H4.
6	PROCESS Model 14	Test whether the indirect effect of ET on PVC through PLD is conditional on GIC.
7	PROCESS Model 6	Test the behavioral extension from ET to Budget Change through PLD and PVC.
8	Robustness checks	Assess balance, influential cases, and planned contrasts.

Appendix G. Supplementary Quantitative Analyses

Appendix G reports supplementary checks on randomization, reliability, manipulation validity, planned contrasts, influential cases, and behavioral validation.

Randomization Balance Checks

Balance checks assessed whether the four experimental conditions differed on relevant pre-treatment characteristics. No statistically significant differences emerged for age, marketing experience, or gender. This supports the comparability of respondents across experimental conditions. Table 23 reports the randomization balance checks.

Table 23

Randomization balance checks

Variable	Test	Result	Interpretation
Age	ANOVA	$F(3, 205) = 0.632, p = .595$	No significant condition difference
Marketing experience	ANOVA	$F(3, 205) = 0.764, p = .515$	No significant condition difference
Gender	Chi-square test	$\chi^2(3) = 0.939, p = .816$	No significant condition difference

Internal consistency was assessed for all multi-item constructs before hypothesis testing.

Evidence Integration, Growth Investment Capability, and Perceived CMO Value Creation all showed excellent reliability. This supported the use of mean scores in the subsequent analyses. Table 24 reports the reliability checks for the multi-item constructs.

Table 24*Reliability checks*

Construct	Reliability	Interpretation
Evidence Integration	$\alpha = .935$	Excellent internal consistency
Growth Investment Capability	$\alpha = .927$	Excellent internal consistency
Perceived CMO Value Creation	$\alpha = .948$	Excellent internal consistency

Manipulation Checks

The Evidence Triangulation manipulation was assessed through perceived Evidence Integration. The manipulation produced a strong effect, indicating that respondents in triangulated-evidence conditions perceived stronger evidence integration than respondents in platform-only conditions.

The GIC manipulation was assessed through perceived Growth Investment Capability. The manipulation produced a strong main effect of GIC. The Evidence Triangulation main effect and the ET $\times$ GIC interaction were not significant, indicating that the GIC manipulation was empirically distinct from the Evidence Triangulation manipulation. Table 25 summarizes the manipulation-check results.

Table 25*Manipulation checks*

Dependent variable	Effect tested	Statistic	Interpretation
Evidence Integration	Experimental condition	$F(3, 205) = 58.83, p < .001, \eta^2 = .463$	Evidence Triangulation manipulation worked
GIC perception	GIC condition	$F(1, 205) = 122.67, p < .001, \eta^2 = .374$	GIC manipulation worked
GIC perception	Evidence Triangulation	$F(1, 205) = 0.93, p = .336$	No unintended ET effect on GIC perception
GIC perception	ET $\times$ GIC	$F(1, 205) = 0.60, p = .441$	No significant interaction

Planned Contrasts for PLD

Planned contrasts examined the 2 $\times$ 2 pattern for Performance Learning Distortion. Evidence Triangulation significantly reduced PLD under both low-GIC and high-GIC conditions. By

contrast, the GIC effect was smaller and depended on the evidence context. Without triangulated evidence, high GIC reduced PLD relative to low GIC. With triangulated evidence, no significant GIC difference remained. Table 26 reports the planned contrast results for PLD.

Table 26

Planned contrast results for PLD

Contrast	Test statistic	Result	Interpretation
ET effect under low GIC	$t(103) = 9.06, p < .001, d = 1.77$	Significant	Evidence Triangulation strongly reduced PLD under low GIC
ET effect under high GIC	$t(102) = 4.83, p < .001, d = 0.95$	Significant	Evidence Triangulation reduced PLD under high GIC
GIC effect without ET	$t(88.94) = 2.71, p = .008, d = 0.53$	Significant	High GIC reduced PLD when triangulated evidence was absent
GIC effect with ET	$t(103) = 0.52, p = .603, d = 0.10$	Not significant	GIC did not add correction when triangulated evidence was present

These contrasts support the interpretation that Evidence Triangulation was the dominant driver of belief correction. GIC showed a weaker corrective role only when triangulated evidence was absent.

Robustness Check: Influential Observations

PROCESS diagnostics identified two influential observations in the original serial mediation model. The serial mediation model was therefore re-estimated after excluding cases 30 and 183. The sample size decreased from 209 to 207. The substantive pattern remained unchanged.

Evidence Triangulation still reduced PLD. PLD still increased PVC. Both PLD and PVC still predicted Budget Change positively. The serial indirect effect remained significant. The direct effect of Evidence Triangulation on Budget Change remained non-significant after including the mediators. Table 27 summarizes the robustness and planned contrast results while Table 28 reports the robustness check excluding influential cases 30 and 183.

Table 27

Robustness and planned contrast results

Test	Statistic	Result
Age balance across conditions	$F(3, 205) = 0.632, p = .595$	No difference
Marketing experience balance across conditions	$F(3, 205) = 0.764, p = .515$	No difference
Gender balance across conditions	$\chi^2(3) = 0.939, p = .816$	No difference
ET effect on PLD under low GIC (C1 vs. C2)	$t(103) = 9.06, p < .001, d = 1.77$	Significant
ET effect on PLD under high GIC (C3 vs. C4)	$t(102) = 4.83, p < .001, d = 0.95$	Significant
GIC effect without ET (C1 vs. C3)	$t(88.94) = 2.71, p = .008, d = 0.53$	Significant
GIC effect with ET (C2 vs. C4)	$t(103) = 0.52, p = .603, d = 0.10$	Not significant
Collinearity Diagnostics	Highest VIF = 2.254; minimum tolerance = .444	No problematic multicollinearity

Table 28

Robustness check excluding influential cases 30 and 183

Path / Effect	Estimate	p-value / CI	Interpretation
ET → PLD	b = -18.805	$p < .001$; 95% CI [-22.390, -15.220]	ET still reduced PLD
PLD → PVC	b = 0.0408	$p < .001$; 95% CI [0.0309, 0.0507]	PLD still increased PVC
PLD → Budget Change	b = 0.0141	$p = .011$; 95% CI [0.0032, 0.0250]	PLD still predicted budget change
PVC → Budget Change	b = 0.3780	$p < .001$; 95% CI [0.2459, 0.5100]	PVC still predicted budget change
Direct effect ET → Budget Change	b = -0.2012	$p = .238$	Direct effect remained non-significant
Total indirect effect	effect = -0.9669	95% BootCI [-1.2706, -0.6882]	Indirect effect remained significant
Serial indirect effect ET → PLD → PVC → Budget Change	effect = -0.2901	95% BootCI [-0.4432, -0.1659]	Serial mediation remained significant

This robustness check indicates that the core behavioral mechanism did not depend on the two influential observations.

A supplementary behavioral validation model regressed Budget Change on PVC. PVC showed a strong positive relationship with Budget Change. Respondents who perceived stronger CMO value creation were more likely to increase the next-period budget. Table 29 reports the supplementary behavioral validation model.

Table 29

Behavioral validation model

Dependent variable	Predictor	B	SE	t	p	95% CI	R ²
Budget Change	PVC	0.539	0.045	12.106	< .001	[0.451, 0.627]	.415

This result supports the interpretation that PVC is behaviorally meaningful for allocation decisions.

Summary of Supplementary Evidence

The supplementary analyses support four conclusions. First, random assignment produced comparable groups across key background variables. Second, all multi-item scales showed excellent internal consistency. Third, both manipulations worked as intended and were empirically distinct. Fourth, the serial mediation mechanism remained robust after excluding influential observations. These analyses strengthen confidence that the reported findings reflect the experimental mechanism rather than group imbalance, unreliable measures, or isolated influential cases.