

Can AI Do Strategy?

Felipe A. Csaszar

Gwendolyn Lee

Peter Zemsky

Todd Zenger

March 1, 2026

Forthcoming in Strategy Science.

Can AI Do Strategy?

Abstract

Can artificial intelligence (AI) do strategy? This question is both urgent and foundational: urgent because AI is already reshaping strategic practice and foundational because answering it forces us to articulate what strategy actually is. In this introductory essay to the *Strategy Science* Special Issue on AI and Strategy, we propose a dual-ladder framework: a causal ladder that maps the cognitive hierarchy of strategic tasks and a delegation ladder that specifies when organizations will grant AI autonomy over those tasks. A core insight emerges: AI will not enter strategy where required reasoning is deepest but where its performance is most measurable. We organize the Special Issue contributions around what AI can do today, could do as capabilities develop, and should do given the imperatives of accountability and human judgment. We close with a challenge and an invitation: if strategy scholars do not define good strategizing precisely enough to be encoded, tested, and refined, other disciplines will, embedding thinner conceptions of strategy into the tools managers use. Teaching machines to strategize and support strategizing is ultimately a method for rediscovering what strategy is.

1 Introduction

The field of strategic management has long navigated waves of transformation, but the emergence of advanced artificial intelligence (AI) appears to mark a profound shift in how we will understand and practice strategy moving forward. Consider the remarkable trajectory of AI capabilities: just a few years ago, generating coherent narratives from unstructured data or simulating multifaceted decision scenarios seemed firmly beyond the reach of machines. Today, these feats are not only possible but increasingly sophisticated, powered by models that process vast information landscapes with speed and precision unattainable by human minds alone. This evolution compels us as a field to confront a foundational question: can AI do strategy? Or, perhaps, can humans and AI, by skillfully partitioning relevant tasks and decision rights, productively collaborate to do strategy?

At its essence, strategy involves sensing opportunities, crafting pathways to valuable future states, and adapting amid uncertainty—all tasks that demand foresight, synthesis, and bold choices. The rise of AI prompts two intertwined inquiries that lie at the heart of this special issue. First, how might AI transform the very processes of strategic reasoning from initial formulation to ongoing adaptation? Second, in what ways could AI redefine the building blocks of achieving valued strategic outcomes, altering how firms establish and sustain competitive advantages in dynamic markets?

The stakes of this inquiry extend far beyond theoretical curiosity; they may touch the core survival of strategic management as a field. As AI permeates strategic decision making in an organization, we argue that our field must lead the charge in exploring its implications lest we relinquish this decision-making terrain to adjacent domains such as computer science. Engaging deeply with AI offers a unique opportunity for self-discovery: by embedding strategic principles into AI systems—essentially teaching machines to “do” strategy or strategy tasks—we stand to gain profound insights into the nature of strategy itself. This reciprocal process could illuminate hidden patterns in how effective strategies emerge, refine our understanding of cognitive bottlenecks in human reasoning, and reveal novel ways to integrate judgment and

reasoning with computation. Through these efforts, strategy scholars are ideally positioned to shape AI’s trajectory, ensuring that it serves as a force for ethical innovation and sustained value creation in business and society.

To navigate this terrain, it is helpful to draw an analogy from another domain in which machine intelligence and autonomy have progressed in measured stages: autonomous driving. Frameworks such as the Society of Automotive Engineers (SAE) levels delineate a spectrum from basic driver assistance to full vehicular independence. Importantly, those levels are fundamentally about delegation and autonomy (who “drives”) even though capability improvements enable that autonomy. In this essay, we leverage this insight but depart from the SAE framing in a critical respect: we separate what the SAE levels tend to conflate. Specifically, we distinguish (i) a hierarchy of strategy capabilities (the cognitive tasks involved in strategizing) from (ii) a hierarchy of delegation (the discretion organizations grant to AI). This separation allows us to be more precise about what it would mean for AI to “do strategy” and why adoption will depend not only on technical capability but also on verification and governance.

Of course, framing the “can AI do strategy” question as binary overlooks the essential facts that, though AI has significant limitations, such as causal reasoning (“Mean Articulation Machines” by McBride 2026, *Strategy Science* (this issue), Felin and Holweg 2024, Song et al. 2026, Yang et al. 2026), along many dimensions, it is rapidly evolving (Dell’Acqua et al. 2023). Thus, a more productive approach contemplates answers along three interrelated dimensions: the present (“can”), the prospective (“could”), and the prescriptive (“should”). On the “can” front, early indicators suggest that AI is already making substantial inroads; for instance, contemporary models can generate and assess strategic proposals (Doshi et al. 2025) with quality rivaling that of seasoned professionals in generating plausible strategic narratives and options, especially in entrepreneurial contexts (Csaszar et al. 2024). The “could” dimension probes deeper and depends on how well AI progresses in its capacity to augment or automate strategic decision making as detailed in our framework. Finally, the “should” lens introduces

vital considerations of governance: even as AI’s potential expands, human elements such as ethical discernment, accountability for outcomes, and alignment with societal values remain indispensable. This human–machine interaction underscores that strategy’s future lies not in automation alone but in thoughtful partitioning, integration, and governance, in which AI augments the effectiveness of strategy without supplanting human responsibility.

A central goal of this special issue is to foster a shared language and a set of concepts that enable strategy researchers to build cumulatively on one another’s work, accelerating progress in this nascent area. Large language models (LLMs), in particular, represent a highly malleable technology, one that requires deliberate exploration to unlock its full value in strategic contexts. Whereas they may harbor the potential to facilitate or accelerate high-quality strategic decisions, we lack clear methods for extracting that value effectively. Open questions abound: What strategic tasks can LLMs effectively perform? How can we measure success and progress? How can AI prompting and human–AI workflows be engineered to elicit robust strategic insights? What combinations of human oversight and AI collaboration yield the best outcomes? By sharing early experiments and empirical findings on these system design choices, the contributions in this special issue help the field advance, transforming speculative possibilities into tangible steps toward practical tools and approaches for strategy formulation and evaluation.

This essay charts a course through these ideas to advance our collective understanding. Section 2 introduces a dual-ladder framework for analyzing AI in strategy: a causal ladder that classifies the cognitive sophistication of strategic tasks and a delegation ladder that describes the conditions under which organizations grant AI autonomy over those tasks. Section 3 reviews the contributions to this special issue, organizing them by their focus on current capabilities (can), future potential (could), and normative governance (should), identifying key research streams for the field. Finally, Section 4 considers how the field can propel itself forward, advocating a pivot toward more engineering-inspired methods that allow strategy scholars to have greater influence and voice to shape the intelligent machines

that will both aid and perform the essential tasks of strategizing. Together, these sections invite scholars, practitioners, and educators to join in building a robust agenda for strategy in an AI-infused era.

2 A Hierarchy of Strategy Tasks and Strategic Delegation

Strategy is not AI’s first domain of application. Strategy’s relative lateness likely reflects the complexity of strategy as a task, as strategy making involves both complex problem framing and complex problem solving. In domains such as autonomous driving, in which engineers and scientists have for decades infused intelligence into vehicles, the industry finds value in categorizing and benchmarking progress using a framework that maps levels of driving functionality and charts AI-enabled autonomy at each level, beginning with blind spot and lane departure warnings at the lowest level, moving to lane centering and cruise control at the midlevel, and culminating with fully autonomous driving in all driving conditions.

With this framework, the field can contemplate nuanced answers to simple questions such as, can AI drive cars? Or can AI assist with driving cars? The predictable answer is that it depends. It depends on the level of driving functionality. It depends on when the question is asked as technology to assist at any given level is rapidly changing. It depends on who is asking and their level of comfort with delegating functionality to AI.

We propose an analogous but more explicit two-part framework for strategic decision making. As with autonomous driving, answers to questions of whether AI can do strategy depend on task demands, the state of available AI technology, and the willingness of humans to delegate decision rights to AI. But, unlike the SAE schema, which bundles capability and autonomy into a single ladder, we separate these dimensions.

We, therefore, begin by outlining a hierarchy of cognitive and analytical tasks involved in strategic decision making. The causal ladder, developed in Section 2.1, characterizes

strategizing as a progression through increasingly sophisticated forms of causal cognition and a premise that AI will diffuse more deeply into strategy as it develops a greater capacity to reason causally. We then examine AI’s current contributions at each level as well as the role of humans. We then turn to developing a distinct ladder of strategic delegation that characterizes how and when organizations grant discretion to AI. The delegation ladder, developed in Section 2.2, starts from the premise that what matters for organizational delegation to AI is verifiable performance relative to a measure and not necessarily reasoning capacity.

The causal ladder takes strategy’s cognitive architecture seriously and asks what kinds of reasoning strategic tasks demand. The delegation ladder takes a performance-first view and asks when organizations will trust AI with strategic discretion regardless of how the AI arrives at its outputs. Whether these premises converge (AI expands its performance of strategy as it develops its capacity to causally reason) or diverge (AI expands its performance of strategy through means other than causal reasoning) is an empirical question that the field must answer. We maintain both ladders precisely because we do not yet know which premise is correct, and the research agenda benefits from taking each seriously.

2.1 Strategy as Nested Levels of Causal Cognition

Before dividing strategy making into a hierarchy of tasks, strategy needs a definition. Despite differences, most definitions of strategy share a common focus: strategy is a purposive process through which decision makers select actions designed to bring about a desired future, varyingly defined as the achievement of a goal (Chandler 1962), the occupation of a valuable position (Porter 1986, 1990), or the resolution of a problem (Rumelt 2011). As humans practice and theorize about it, strategic decision making is, therefore, a causal enterprise, one that presupposes beliefs—often tacit beliefs—about how particular initiatives, investments, or policies generate desired outcomes. Even when such beliefs are incomplete or uncertain, they provide the basis for comparing alternative courses of action and making choices that subsequently guide all others (Van den Steen 2018).

Viewed in this way, strategy approaches vary in the depth of causal reasoning and analysis that they demand. Some approaches emphasize prediction and planning within existing structures, treating strategy as forecasting and alignment with forecasts (e.g., Ansoff 1965). Others focus on estimating the causal effects of deliberate actions or resource commitments (e.g., Andrews 1971, Barney 1991). Still others highlight counterfactual reasoning and analogical extrapolation across contexts, perhaps framing strategy as the selection of different positions or configurations that are distinct from those occupied or composed at the present (Porter 1980, 1996; Gavetti et al. 2005). Another approach emphasizes strategy as the construction of new strategic models that redefine problems, model boundaries, and action spaces (e.g., Nickerson and Zenger 2004, Zott and Amit 2010, Felin and Zenger 2017, Teece 2018). Rather than treating these approaches as competing, we interpret these variations as emphasizing and demanding different levels of causal cognition.

To organize these differences, we depict strategic decision making as a hierarchy of increasingly sophisticated forms of causal cognition. We build on and extend Pearl’s (2000) ladder of causation, which distinguishes among reasoning based on observed associations, interventions, and counterfactuals, adapting this logic to the domain of strategy. In doing so, we distinguish counterfactual reasoning about new actions from reasoning about entirely new states and the causal models required to achieve them (Lee and Bettis 2023). This extension yields four levels of strategic decision making that progress in depth and scope of causal reasoning. Table 1 provides an overview of these four levels, summarizing the characteristic questions posed at each level and the corresponding object of strategic reasoning. At lower levels, strategizing is anchored in observed data and projections from existing patterns; at higher levels, strategy making increasingly departs from present data and requires model extensions and model building that enables projection beyond what has been directly observed. In the sections that follow, we elaborate each level in turn and provide illustrative examples, building in the process a more robust picture of AI’s capacity to contribute to strategic decision making as well as the role of humans, which we summarize in Table 2.

Table 1: Levels of Strategic Causal Reasoning

Level	Question/analysis	Object of reasoning
1	What will happen if X continues?	Associations/trends
2	What happens if I do X ?	Interventions on variables
3	What would happen if we did Y instead?	Counterfactual outcomes of a model
4	What model would make Z (e.g., the resolution of a problem) possible?	Causal models of a novel state

Notes. The table extends Pearl’s (2000) ladder of causation to illustrate four increasingly sophisticated forms of strategic decision making. Levels 1 and 2 operate within an existing causal model, progressing from simple prediction to intervention. Level 3 extends an existing model to consider counterfactual extrapolations. Level 4 builds entirely new models of a valuable future state.

2.1.1 Level 1: Analytic–Predictive Strategy.

Level 1 is analytic–predictive strategy, demanding diagnostic and predictive reasoning. Strategic decision making at this level asks, what should I do given that I observe X or expect to observe X ? The task is to identify regularities, estimate trends, and extrapolate their implications for action. Although future-oriented, this reasoning predicts outcomes from existing data within a given model structure and action space. For example, based on historical sales patterns and macroeconomic indicators, an organization may adjust capacity, inventory, or hiring. The strategist need not know the precise causal mechanism behind sales growth; simply recognizing forecast growth is sufficient.

Firms improve decision making here by aggregating diverse data sources to enhance predictive accuracy. AI is particularly powerful at this level. Machine learning systems detect regularities and generate probabilistic forecasts, whereas LLMs extract signals from unstructured data, synthesize qualitative information, and translate forecasts into actionable narratives. Together, these tools expand the informational basis of predictive analysis.

Humans remain central in defining which questions to ask, identifying relevant performance dimensions, judging which forecasts matter, and determining how predictions translate into concrete responses. Decisions about relevance, risk, and action continue—at least for now—to rest with human judgment.

Table 2: Levels of Strategic Causal Reasoning and the Human–AI Division of Cognitive Labor

Level	Core task	Strategic question	Example	AI role	Human role
Level 1: Analytic– predictive	Forecast outcomes within established models	What will happen if X continues?	Adjusting capacity, inventory, or hiring based on projected demand	Generates forecasts, detects patterns, quantifies uncertainty	Defines relevant outcomes, selects focal variables, decides how forecasts inform action
Level 2: Intervention-oriented	Estimate effects of familiar actions	What happens if we do X ?	Evaluating product expansion, marketing increases, or acquisitions	Estimates treatment effects, supports experiments, compares interventions	Defines objectives and candidate actions, assesses assumptions, weighs trade-offs
Level 3: Counterfactual– extrapolative	Project effects of novel but analogous actions	What would happen if we did Y instead?	Assessing a shift to a platform business model or entry into a new regulatory regime	Surfaces analogies, synthesizes external evidence, explores counterfactuals	Judges relevance of analogies, evaluates transferability, sets risk bounds
Level 4: Model-generative	Reconfigure the causal model and action space	What model would make Z possible?	Building a new-to-the-world product, service, or business model	Recombines concepts, stress tests assumptions, compares alternative models	Defines new problem frames, evaluates coherence and fit, authorizes commitment

Notes. The levels reflect increasing depth of causal reasoning. Across all levels, AI augments analysis (prediction, estimation, extrapolation, generation), whereas humans retain authority over problem framing and strategic commitment. At every level, humans play a central role in problem framing though the nature of framing shifts from relevance selection (Level 1) to intervention definition (Level 2) to analogical boundary setting (Level 3) to the construction of entirely new strategic problem frames (Level 4).

2.1.2 Level 2: Intervention-Oriented Strategy.

Level 2 focuses on establishing the causal effects of actions, choices, and interventions. The predictive capabilities of Level 1 are augmented by forming causal representations that support identification and estimation of the effects of deliberate actions still within an established and previously observed set of strategic options. Strategic problem framing takes the form of what happens if I do X ? Analytics estimate how outcomes change when familiar actions are taken under different configurations or intensities observable in data.

Whereas prediction remains the goal at Level 2, the effort is no longer forecasting trends but forecasting causal effects. The aim is to assess which actions by the firm or close analogs operating under comparable conditions will shape performance. A firm may draw on internal experience or industry cases to evaluate whether expanding a product line or increasing the intensity of marketing will improve performance.

Because moving from Level 1 to Level 2 requires identifying the causal links connecting actions to outcomes, it demands more sophisticated tools and estimation methods than Level 1. For instance, DoorDash’s analytics group may observe a strong correlation between discount coupons and purchase likelihood and conclude that aggressive couponing is merited. However deeper Level 2 analysis reveals a substantial selection bias. Coupons only cause a specific subset of consumers to purchase, and therefore, a very targeted couponing strategy is sufficient. A consumer products company, using purchasing about toothpaste and floss may simply observe that customers who buy toothpaste are also likely to buy floss (see Pearl 2021, p. 19) and therefore decide to place the two products adjacent on retail. But a deeper Level 2 analysis would want to know if such placement causes increased sales.

At this level, AI identifies and estimates the precise causal effects of available strategic actions. Analyses may use experimentation platforms, causal inference techniques, and machine learning to assess actions undertaken by the firm or close analogs. LLMs surface candidate interventions by synthesizing evidence from comparable organizations or external studies, thereby expanding the set of actions available for evaluation. Whereas AI improves

the precision and transparency of causal analysis, it does not determine which interventions to pursue.

At present, humans still act as intervention judges at Level 2, determining desired outcomes, evaluating candidate actions, assessing the plausibility of causal assumptions, interpreting estimated effects, and deciding which trade-offs are acceptable given organizational and ethical constraints.

2.1.3 Level 3: Counterfactual–Extrapolative Strategy.

Level 3 is counterfactual–extrapolative strategy. It is characterized by counterfactual reasoning that extrapolates the consequences of actions beyond the firm’s direct experience or data, holding the underlying causal model largely fixed or incrementally extended. Whereas Level 2 strategizing evaluates familiar interventions using observed or closely comparable data, Level 3 strategizing frames strategic questions about how actions not previously undertaken by the firm might perform if transported from other industries, geographies, or institutional contexts. The strategist does not abandon the effort to establish causality, but here, novelty arises from adapting existing causal models to imagine plausible, counterfactual futures generated by extrapolating from known causal logic into new domains.

By observing the actions of other firms or other parts of the firm, often in other industries or geographies, a strategist may assess the likely effects of shifting to a platform business model, adopting a novel pricing mechanism for the industry, or entering new markets operating under different institutional regimes. For instance, as Apple experiences performance limitations in outsourcing chips for its laptops, it extrapolates from its success integrating chip design into smartphones to conclude that a similar approach in laptops could elevate performance.

AI systems at this level can explore analogical extrapolations. They can organize and synthesize evidence from external firms, industries, geographies, and institutional settings; surface candidate analogies; and explore counterfactual action scenarios by simulating the implications of transporting causal relationships across contexts.

At this level, humans frame strategic problems, identify desired outcomes, and act as arbiters of AI-generated extrapolations. They judge which external analogies are strategically relevant and whether contextual differences constrain their valid transport.

2.1.4 Level 4: Model-Generative Strategy.

Level 4 is model-generative strategy. It marks a qualitative shift from applying or extending existing causal representations to actively constructing new causal models that define what actions, mechanisms, and outcomes are conceivable. Whereas Level 3 strategizing relies on counterfactual reasoning by adapting or extending existing models, Level 4 strategizing builds new models, targeting the resolution of entirely new strategic problems. Strategists at this level generate, recombine, and revise causal structures, articulating new mechanisms and causal elements that target an envisioned future state. Level 4 strategizing often focuses on the resolution of a novel problem or set of problems. It is Sam Walton composing Walmart as a novel geographic composition that solves a supply chain problem or Steve Jobs designing the Macintosh as a computing product that is elegant and easy to use. Level 4 strategizing is architecting entirely new futures, wholly distinct from the observable present.

Level 4 strategy subsumes the capabilities of the preceding levels: predictive reasoning informs the plausibility of newly composed models, interventional logic explores imagined actions within them, and counterfactual reasoning supports comparing alternative futures. What distinguishes this level is that these forms of reasoning are now organized in the service of constructing models rather than applying them.

At Level 4, AI can assist with model construction. LLMs may highlight critical steps to accomplish or subproblems that must be solved to achieve some future state. They may help recombine concepts, identify critical mechanisms, or perhaps surface implicit assumptions.

At Level 4, humans act as strategic theorists, envisioning novel future states and composing novel problem frames. Human judgment is particularly critical because the evaluation of models rests on data that are, at present, largely unavailable (by definition) and must be

generated by a costly launch of the strategy. The framing of problems and the composition and the composition of theories that solve them remains a distinctly human activity.

Taken together, these four levels depict a causal ladder as summarized in Table 2.

2.2 The Delegation Ladder: Verifiability and Strategic Autonomy

The causal ladder specifies what “doing strategy” entails as a progression of reasoning tasks; the delegation ladder explains when organizations will grant AI discretion over those tasks. Critically, the two rest on different premises. The causal ladder privileges reasoning as the vehicle for elevated performance; the delegation ladder privileges demonstrated performance regardless of its source, allowing for the possibility that AI systems may achieve strategic performance through methods that do not resemble human causal theorizing: prediction, search, or optimization under a scoring rule.

This performance-first view is motivated by AI’s own history. Moravec’s (1988) paradox reminds us that what humans find cognitively “basic” can be computationally hard, whereas what we find “high level” can sometimes yield to scale and optimization. Many of AI’s practical breakthroughs—from speech recognition to game playing—arose not because machines mirrored human reasoning but because they achieved measurable performance advantages under explicit scoring rules (Sutton 2019). If AI is defined as acting rationally—optimizing behavior with respect to an objective—without requiring that the system think like a human (Csaszar and Steinberger 2022), the relevant question becomes not when an AI “theorizes,” but when its owners grant it discretion. The quip attributed to Fred Jelinek, then director of IBM’s speech recognition laboratory—“Every time I fire a linguist, the performance of the speech recognizer goes up” (Moore 2005, p. 1)—captures the lesson that hand-imposed theoretical structure can sometimes constrain performance when data and feedback allow learning directly from outcomes (Shmueli 2010, Csaszar and Ostler 2020).

None of this implies that causal reasoning is irrelevant. Strategic environments are nonstationary, competitors react, metrics get gamed, and high-stakes moves generate sparse

feedback (Csaszar 2018). In such settings, causal models may be essential for robustness and governance even when not strictly necessary to produce a strong recommendation in a narrowly scoped task. But the key practical distinction is that performance is mandatory, whereas explanation is sometimes optional; and when explanation is required, it is often required by governance (auditability, accountability, legitimacy) rather than by optimization alone (Lee 2026).

The core claim of this section follows: “can AI do strategy?” is answered as much by institutions of verification and governance as by advances in causal cognition. Strategy will become more machine-driven first in domains in which performance can be scored, risk can be bounded, and accountability can be specified. As a result, advanced causal capability (Levels 3–4 on the Causal Ladder) is neither always necessary nor sufficient for delegation. This capability may be important for performing some strategic tasks, yet still fail to justify autonomy where verification is weak, whereas in other settings an AI can earn substantial discretion through verifiable performance-driven methods that look more like prediction, search, or optimization than like human causal theorizing. This motivates a delegation ladder that tracks how much discretion AI systems receive related to, but distinct from, the kinds of reasoning they perform. Whether AI’s strategic autonomy expands primarily by acquiring more causal cognition or by exploiting verifiable performance gradients through other cognitive routes remains an open empirical question.

2.2.1 Delegation Depends on Verifiability, Downside Bounds, and Accountability.

Rational delegation to AI is granted when there is a track record of performance against an accepted scoring rule. Algorithmic trading provides a useful analogy (Csaszar et al. 2024, p. 323). Once there was a verified metric demonstrating superior performance, competitive pressure rather quickly pushed decision rights toward the machine. Of course, strategy differs in material ways: feedback is slower, causal attribution is harder, and the objective is less well-defined. Yet a binding constraint on higher delegation in strategy is not just model

capability, but verification infrastructure: benchmarks, simulations, audits, and evaluation protocols that provide credible evidence of decision quality before years of outcomes arrive (Csaszar 2025). There is one caveat. Goodhart’s law reminds us that, when a measure becomes a target, it ceases to be a good measure; thus, defining objectives that resist gaming remains a central challenge in automating strategic tasks.

We believe the pace of delegation will be governed by four practical conditions, conditions that also order the tiers of autonomy below:

- (i) Feedback density/verifiability (can we score decisions credibly and repeatedly?).
- (ii) Downside boundedness/reversibility (can errors be contained by guardrails, budgets, or the ability to reverse course?).
- (iii) Objective specifiability (can we state what “good” means without Goodharting ourselves?).
- (iv) Accountability/liability (can we assign responsibility and audit the system’s behavior?).

2.2.2 The Delegation Ladder: Four Tiers of Strategic Autonomy.

This ladder parallels the career development of a human strategist, moving from an analyst with little discretion to a CEO with full executive authority. Table 3 summarizes the four tiers and the conditions governing progression across them. Critically, the tiers describe delegation regimes, not cognitive demands. A task at any causal ladder level can appear in theory at any delegation tier, if the four conditions above are met.

- Tier 1: Informational support (analyst). AI acts as sensor and summarizer, scanning markets, synthesizing competitor moves, extracting signals from unstructured text, drafting analyses. Errors are visible and correctable; governance is light. This level is already widespread in practice.

Table 3: The Delegation Ladder: A Matrix of Strategic Autonomy

Tier of delegation	Typical strategic role	Primary AI function	Accountability and governance
Tier 1: Informational support	Analyst	Sensing, summarizing, scanning, and organizing information	Light: humans retain full decision rights and responsibility
Tier 2: Tactical recommender	Associate	Generates options, scenarios, and structured arguments within a human-defined frame	Moderate: responsibility remains human; AI advises but does not decide
Tier 3: Bounded operator	Unit manager	Executes recurring strategic decisions under explicit constraints	Formalized: oversight, audits, and escalation rules assign responsibility
Tier 4: Strategic executive	CEO	Proposes or executes major strategic commitments and reallocations	Intensive: authority, liability, and governance arrangements become binding constraints

Notes. This matrix orders AI deployment in strategy by the degree of discretion granted, not by the sophistication of causal reasoning. Delegation increases as organizations develop clearer objectives, stronger oversight mechanisms, and greater confidence in assigning responsibility for outcomes. Whereas related to the causal ladder in Table 2, delegation can advance unevenly across strategy tasks and does not require human-like reasoning.

- Tier 2: Tactical recommender (associate). AI proposes options inside a human-defined frame: idea generation, scenario drafting, structured debate, premortems. Value comes from expanding and accelerating search, whereas humans retain decision rights. This is when virtual crowds (cheaply generating and aggregating diverse perspectives) could, in some settings, outperform committees by reducing social frictions (Boussioux et al. 2024).
- Tier 3: Bounded operator (unit manager). AI executes recurring decisions under explicit constraints. Examples include dynamic pricing adjustments in response to competitor moves, real-time resource allocation across a portfolio of projects, or automated responses to standard competitive signals. Whereas these tasks often sit within functional domains, they become strategic when they determine the firm’s competitive posture or resource allocation at scale. Delegation grows here because objectives are clearer, feedback is frequent, and downsides can be capped by budgets and rules. Verification is continuous and operational.
- Tier 4: Strategic executive (CEO). In this final tier, AI shapes the firm’s fundamental direction: proposing or enacting major moves such as mergers and acquisitions, business model pivots, or market entry. The core value of AI at this level is its ability to simulate complex scenarios and stress-test strategies against myriad variables, surfacing options that are robust rather than merely intuitive. However, achieving this level of autonomy faces significant technical and institutional hurdles. Because strategic feedback loops are slow and the stakes are existential, AI in this role requires rigorous validation mechanisms: stress tests to expose fragility, clear liability frameworks, and agent scorecards that allow boards to understand the AI’s risk appetite and biases before granting it authority.

2.2.3 Two Ladders, One Research Agenda.

The causal ladder and the delegation ladder are complements, not rivals:

- The causal ladder helps us specify which strategy tasks an AI system can perform and when human judgment remains essential, especially in Levels 3 and 4, in which extrapolation, problem framing, and model construction matter most.
- The delegation ladder helps us predict where adoption and autonomy will expand first—where goals are specifiable, feedback is dense, and risks are bounded.

Advances in either ladder accelerate the other. Better causal capabilities expand what can be delegated safely; better verification infrastructure makes delegation rational at higher levels and provides the performance gradients that allow systems to improve.

Ultimately, “can AI do strategy?” decomposes into two distinct problems: whether AI can perform the causal–cognitive work that strategizing requires and whether organizations can verify that work well enough to delegate it. These ladders may converge—if high-performing AI learns to reason causally—or they may diverge, with AI achieving strategic results through methods that look nothing like human theorizing. Regardless of the path, this framework clarifies the research agenda: the field must develop both the cognitive tools that enable deeper strategic reasoning and the measurement infrastructure that allows organizations to trust AI with consequential choices. The contributions to this special issue, reviewed in the next section, represent early efforts on both fronts—mapping current capabilities, prototyping new tools, and beginning to construct the intellectual infrastructure that cumulative progress requires.

3 Emerging Research Streams in Strategy and AI

Beyond our introductory essay, there are eight papers that comprise the special issue. These contributions can be organized around the three dimensions introduced in Section 1: what can

AI currently do in strategy, what could AI do as capabilities and infrastructure develop, and when (and whether) should AI be granted discretion over strategic tasks. Section 3.1 opens with a multiauthor dialogue that engages all three dimensions directly. Sections 3.2–3.5 then turn to the empirical and design-oriented papers organized by the emerging research streams they represent: assessing current capabilities, benchmarking strategic performance, reporting software experiences, and modeling market-level dynamics when AI enters competition.

3.1 A Dialogue and Debate

“Can AI Do Strategy? A Dialogue and Debate” (Chatterji et al. 2026, *Strategy Science* (this issue)) is a compilation of short essays that emerged from a two-panel dialogue during the AI–Strategy Conference, cosponsored by the ION Management Science Laboratory at the University of Utah and *Strategy Science*. The panelists, Aaron (Ronnie) Chatterji, Felipe Csaszar, James Evans, Teppo Felin, Jessica Hullman, Karim Lakhani, and Mari Sako, were asked to convert their transcripts from this dialogue into short essays that both present unique perspectives on the central question, “can AI do strategy?” and nicely speak to the question of what AI can, could, and should do.

Mari Sako grapples primarily with the “should” dimension: the conditions under which AI ought to be granted strategic discretion. Drawing on the professional model of diagnosis, inference, and treatment, she argues that, whereas AI may increasingly assist with diagnosis and aspects of inference, the treatment stage—arriving at and implementing context-specific decisions grounded in professional judgment—remains hardest to automate. Because strategy also starts with problem framing (Nickerson and Zenger 2004) and is a job function with workflows of tasks that are hard for AI to learn (Acemoglu 2025), even a full automation of strategic decision making leaves human strategists in the roles of framing problems and using their professional judgment to determine which tasks AI should perform.

Whereas Sako’s view is rooted in the primacy of human judgment, others emphasize the question of what AI “could” do. Felipe Csaszar draws on his research on foresight—the ability

to predict which courses of action will lead to better outcomes (Csaszar and Laureiro-Martínez 2018)—to argue that the decisive question today is whether we can make strategy verifiable enough for AI to learn the links between present actions and superior future value. He offers a set of testable propositions to guide the field of strategy toward building an intellectual infrastructure of verifiability with a clear performance gradient needed for AI systems to improve.

Karim Lakhani reframes the AI and strategy debate by focusing on “strategy in the wild”—how strategy is actually practiced in organizations. He advocates that the field of strategy embrace a new, evidence-based research agenda that involves conducting rigorous clinical trials of AI as a transformative force in practice.

Aaron (Ronnie) Chatterji emphasizes that strategy does not exist in a functional silo in organizations but is an integrative function that needs to align with finance, marketing, sales, and operations. He recommends that strategy researchers engage with AI researchers at frontier labs so as to shape the development of AI models that practitioners use across functions.

Whereas these “could” commentaries look toward expanding frontiers, three scholars foreground the “can” dimension’s present limitations. James Evans, a professor of sociology, suggests that AI will only do strategy effectively if it helps expose productive disagreement within an organization. He submits that strategy requires diverse perspectives, not a single recommendation. As such, AI should not be used to replace strategic judgment but to amplify the conflict that makes strategic judgment possible.

Jessica Hullman, a professor of computer science, argues that AI is unlikely to replace human strategists when context, goals, and values must be brought to bear in formulating the problem and identifying truly novel solutions. Her main concern are the contingencies around what future actions the firm should take in the face of uncertainty about the environment (e.g., market conditions) under competition from other firms. Problem framing under uncertainty is a critical strategy task with which AI currently struggles.

Finally, Teppo Felin draws from his research on the limits of prediction, arguing that AI cannot do strategy because it is—by construction—backward-looking, population-level, and based on a correlational recombination of past patterns. By contrast, strategic decision making demands forward-looking causal reasoning that constructs novel elements and combinations to achieve novel future states.

These commentaries, spanning “should,” “could,” and “can,” collectively map the conceptual terrain that the empirical and design-oriented papers in this special issue begin to cultivate. We turn now to those contributions, organized by the research streams they represent.

3.2 Assessing AI’s Current Strategic Capabilities

A natural starting point for delineating new research streams is to clarify the cognitive tasks involved in strategic decision making and to assess the extent to which AI can currently perform them. The paper “Mean Articulation Machines” (McBride 2026, *Strategy Science* (this issue)) offers precisely this conceptual groundwork. It highlights three foundational modes of intelligence: associative, causal, and interventional. Associative intelligence excels at identifying patterns and regularities across large information sets; causal intelligence models mechanisms and evaluates how interventions propagate through a system; interventional intelligence goes further still, constructing and comparing counterfactual futures. McBride’s “associative intelligence” corresponds closely to Level 1 of the causal ladder and “causal intelligence” maps onto Level 2, whereas “interventional intelligence” roughly maps to our Levels 3 and 4: the ability to reason about the effects of actions in untested contexts and, at the highest reaches, to construct new causal models altogether. These distinctions explain why strategic reasoning comprises multiple heterogeneous tasks—prediction, analogy, causal explanation, foresight, and intervention design—each with its own cognitive demands.

Within this framework, contemporary AI systems appear as highly capable associative reasoners with strong articulation abilities yet limited causal discrimination and minimal

interventional capacity. Humans, by contrast, possess richer causal reasoning and contextual judgment although they lack the breadth, speed, and consistency of machine-based pattern extraction. These contrasts provide a principled basis for identifying when AI can augment strategists and when human cognition remains indispensable.

Two empirical papers place these distinctions into practice. “The Strategic Value of Predictions in Acquisition Decision Making” (Kumar et al. 2026, *Strategy Science* (this issue)) evaluates AI performance in prediction tasks using stock market reactions to acquisitions as a setting. The study shows that machine learning models, trained on extensive historical data, can outperform human decision makers in domains in which noise structures and redundant cues favor associative inference. These findings illustrate high performance in Level 1 strategy and echo the strengths McBride attributes to contemporary AI: broad cue integration, efficient pattern recognition, and robust prediction from abundant data.

The next contribution, “AI-Augmented Strategic Decision-Making Under Time Constraints” (Kanis et al. 2026, *Strategy Science* (this issue)), probes the effectiveness of human–AI collaboration in shaping strategic judgment. The study examines how individuals form mental representations of strategic problems under time pressure and with access to an LLM. The results reveal that time constraints compress representations—reducing breadth and preserving depth—whereas LLMs expand breadth but dilute depth. Significantly, neither intervention improves strategic foresight. LLM-generated inputs risk inducing information overload and weakening psychological ownership, whereas time pressure heightens reliance on salient cues. These findings reinforce the logic of our causal ladder and McBride’s insight that the articulation of more information is not equivalent to better causal reasoning. At the same time, the authors illuminate the cognitive bottlenecks that humans face when relying on AI-generated representations.

Across these three opening papers, a coherent picture of what AI can do today emerges. Strategic tasks vary in the mix of associative and causal reasoning they require, and AI’s current strengths and limitations map systematically onto this spectrum. Human–AI comple-

mentarity, thus, becomes a natural organizing principle: machines contribute breadth, speed, and associative capability; humans contribute contextual interpretation, causal discrimination, and theoretical reframing. At the same time, the evidence cautions against naïve integration of AI into strategic work as misalignment between task demands and AI capabilities can distort rather than enhance decision quality.

3.3 Benchmarking AI’s Strategic Performance: Designs, Tests, and Simulations

To move beyond mapping current capabilities and explore what AI could do, the field requires rigorous methods for tracking progress. Benchmarks have aided and accelerated progress in LLMs across a variety of domains, including math, science, dialogue, coding, and professional certification exams. AI has now surpassed various human benchmarks in coding, reading comprehension, multimodal reasoning, and even PhD-level science questions. In other areas, LLM benchmarks provide an objective yardstick for tracking various dimensions of LLM performance to inform the limits of their use in specific domains (e.g., UC Berkeley SkyLab 2025). Importantly, such benchmarks significantly improve the quality of the systems they track, stimulating innovations in their respective fields by defining targets for further research and development; LLM industry leaders are actively calling for more domain-specific benchmarks to push forward model performance (e.g., OpenAI 2025). Yet, in the domain of strategic decision making, no comparable infrastructure exists. The next two papers begin to fill this gap.

“How Well Can AI Do Strategy? Empirical Benchmarking Using Simulations” (Allen and McDonald 2026, *Strategy Science* (this issue)) pioneers the development and demonstration of a benchmark for evaluating LLM capabilities for strategic decision making characterized by uncertainty, complexity, irreversible multiperiod moves, and delayed or noisy feedback. The paper reports impressive gains in LLM capabilities for strategic decision making and also reveals a concerning vulnerability. The capacity of the state-of-the-art models to manage

strategic uncertainty has declined. The decline points to the need for developing standardized LLM benchmarks and tracking domain-specific LLM capabilities continuously.

Importantly, the paper offers a strategy AI benchmark typology (see table 1 in Allen and McDonald 2026, *Strategy Science* (this issue)) to enable strategy researchers to track LLM progress in the domain of strategic decision making. The typology offers a strong direction for future research because existing studies in the strategy field typically examine human–LLM collaboration for just a single model. Results based on a single LLM at a time cannot truly reveal the evolution of LLM capabilities because of rapid changes in the technology as the underlying AI models improve.

“Can LLMs Aid Analogical Reasoning for Strategic Decisions? A Comparative Study” (Sen et al. 2026, *Strategy Science* (this issue)) develops task-level benchmarking for comparing the performance between humans and eight state-of-the-art LLMs on retrieving and matching analogies to the problems at hand. In contrast with the simulation-based benchmark by Allen and McDonald, the task-based benchmark by Sen, Workiewicz, and Puranam reports the first study mapping a cognitive frontier on analogical reasoning. Importantly, the paper examines the complementarities in the verbal analogical reasoning abilities between modern LLMs and humans in the context of business problems. The transformer architectures underlying the LLM models excel at detecting similarity and relevance between ideas expressed as text, that is, between a given prompt and patterns embedded in their training data. Yet it is an open question regarding LLMs’ ability to generate high-quality matches that show contextual fit with business problems. In this exploratory study that extends classic analogical transfer designs, the authors find a clear trade-off: humans frequently overlook valid analogies (low recall) but rarely misapply them (high precision); LLMs, in contrast, rarely miss valid analogies (high recall) but often surface spurious even if internally coherent matches (low precision).

Future research can continue to develop standardized LLM benchmarks across strategic domains. Allen and McDonald’s (2026, *Strategy Science* (this issue)) limitations provide

rich avenues for improvement: the simulation does not capture true multiplayer competitive dynamics; it abstracts from political dynamics, ethical considerations, and evolving stakeholder interests, and the LLMs evaluated lacked agentic capabilities. Beyond refining simulation-based approaches, a natural next step is to complement them with prospective designs that test whether LLMs can exercise strategic judgment under genuine real-world uncertainty, in which outcomes are unknown at the time of evaluation.

3.4 Software Experiences: What AI Could Do with Better Tools

If benchmarks provide the measurement infrastructure for tracking AI’s strategic capabilities, software tools provide the design infrastructure for expanding them. “AI and Theory-Based Strategy: Experimental Evidence” (Camuffo et al. 2026, *Strategy Science* (this issue)) demonstrates the potential for a new research stream on the practical experiences, lessons learned, and insights gained from developing AI-based software systems or tools. This study reports how a software tool shapes strategists’ reasoning, beliefs, and strategies. The paper documents the design journey of Aristotle, an agentic multiagent AI system for Level 4, theory-based strategic decision making. It highlights design choice trade-offs in composing such a strategy framework and demonstrates how humans interact with software and how those interactions lead to change. One of the key insights from this paper is captured in a five-dimensional taxonomy that maps the design space for agentic AI systems and a methodological road map that enables researchers and practitioners to experiment with, evaluate, and iteratively improve their AI system design choices. The taxonomy is rooted in a design view drawing from the rich literature of system complexity.

Future papers reporting on software experience can draw on the design tradition inspired by Simon (1969, 1993), who championed both the discovery of new alternatives and the design of “how things ought to be.” Such engineering approaches—designing, building, and testing software tools—will enable cumulative progress in strategy research and potentially launch a nascent research area. Bettis et al. (2016), Shaver (2020), and Lee (2024) initiate

discussions regarding the repeatability and cumulativeness of statistical research knowledge in strategic management. In our special issue, we extend their call by championing the reporting of software experiences as a new research practice in which sharing code is the norm (e.g., GitHub repository) and transparency in developing and testing novel AI tools is a vital contribution in itself.

3.5 Strategic Interaction and Market-Level Dynamics: What Should Happen When AI Enters Competition

Strategy is fundamentally relational: the outcomes that matter most for firm performance emerge not from isolated decisions but from interaction among rivals, partners, and stakeholders. As AI systems begin to participate directly in decision making, understanding how they reshape these interaction patterns becomes essential.

The closing paper in this special issue, “When AI Does Strategy: Learning, Good Times, Lock-in, and Human-Driven Strategic Renewal” (Neshenko and Ryall 2026, *Strategy Science* (this issue)), explores how finite-horizon, feedback-driven AI systems might behave under dynamic competition, how synthetic forms of innovation can accelerate learning and shape competitive outcomes, and how interaction among such agents may generate stabilizing and self-reinforcing patterns in which innovation stops. The analysis highlights the continuing importance of Level 4 human theorizing because of our ability to introduce new causal frames and de novo innovations capable of reshaping the strategic landscape. In doing so, the paper exemplifies the questions that are likely to animate future research on AI-enabled strategic interaction.

These questions represent a natural extension of long-standing concerns in competitive dynamics, evolutionary models, game theory, and behavioral strategy. What shifts in an AI-infused world are the underlying assumptions about foresight, adaptability, bounded rationality, and the informational environment. The implications for empirical research are equally significant. Understanding how humans, AIs, and hybrid agents coevolve in

competitive settings will be essential for explaining performance differences; anticipating new patterns of rivalry; and guiding how AI technologies are designed, governed, and integrated within firms (Krakowski et al. 2023).

Across these research streams, a consistent imperative emerges: the field must advance the cognitive architecture of strategizing, simultaneously building the measurement and design infrastructure that makes strategy verifiable and governable. We develop this dual mandate further in Section 4.

4 Reclaiming Strategy’s Core Purpose Through AI

The contributions to this special issue—conceptual frameworks, empirical benchmarks, software prototypes, and models of strategic behavior—collectively point toward a larger opportunity for the field of strategy. AI does not merely pose questions for strategy to answer; it provides tools with which strategy can renew its founding ambition. In this section, we argue that the field should seize this moment to complement its established strengths in explanation with a parallel emphasis on building, testing, and refining the decision tools and processes that make effective strategizing possible.

4.1 Artifacts as Method

The field of strategic management was founded on a normative ambition: to help decision makers navigate complex, high-stakes choices about the future. Much early work was engaged in developing logic, tools, and empirical insights directly focused on helping strategic actors with this navigation. Over time, strategy scholarship has drifted in many different directions, but frequently toward questions selected for empirical tractability—often rather narrow causal identification exercises rather disconnected from the cognitive challenges that define actual strategizing.

One measure of the field’s distance from its original normative agenda is that the content

of the average strategy class has seen limited change over the past 30+ years. Whereas strategy papers fill journals, surprisingly little of this research has migrated into reshaping the way strategists are taught to strategize. In a typical class, we define strategy and offer tools for analyzing environments, capabilities, and resources. We categorize and describe effective strategies and offer examples of success with an occasional example of failure. But what we fail to do is offer strategists a well-validated and well-researched process for actually doing strategy, particularly Levels 3 and 4 strategizing, for developing and testing causal models of valuable future states or composing and testing for the causal impact of novel actions and choices.

We believe the rapid progress of AI tools affords the strategy field an opportunity to renew its normative ambitions and return to addressing foundational strategy questions about improving strategy making. The emergence of LLMs reopens a path to the field’s original mandate, now with tools powerful enough to make progress on problems long deferred. The contributions to this special issue, discussed in Section 3, represent early steps along this path: they benchmark current capabilities, prototype new software tools, and map the cognitive terrain that such tools must navigate. The opportunity is not to abandon our field’s deep commitment to social science research and transform into something more like engineering or applied operations research. Rather, our hope is that AI will help strategy as a field complement a focus on explanation with a parallel emphasis on decision tools and processes that support strategic reasoning.

A focus on tool creation for strategy does not mean a departure from rigorous scholarship about strategy. When we design an AI system intended to help managers generate novel business models or evaluate counterfactual futures, we should aspire to normatively specify what effective strategizing requires. When we test whether that system improves human or organizational performance, we generate insight into cognitive limitations and the conditions under which AI or process augmentation succeeds. The learning process is reciprocal: building human-AI systems that strategize forces us to clearly articulate what strategizing entails,

and observing when those systems fail reveals what we had not yet understood.

More than this, when we embed strategic concepts into AI and interactions with it, we create manipulable instantiations of strategic reasoning: an artifact with which to experiment. We can ask, what happens when we alter the causal mechanism that the system uses to link actions to outcomes? When we change the information inputs, say, restricting competitor data or adding regulatory constraints? When we modify the feedback structure: immediate versus delayed, richer versus parsimonious? These variations constitute experiments on strategizing itself, made possible because the artifact renders reasoning measurable and partly controllable. By building and testing decision aids, strategy as a field can effectively move from merely observing strategizing to actively improving it.

4.2 The Frontier: Enhancing Levels 3 and 4

Abundant tools already exist to assist with Level 1 and Level 2 strategizing on the causal ladder. Many have been developed by scholars in adjacent fields. Statisticians and econometricians have developed statistical packages. Computer scientists have developed machine learning tools. Operations research scholars have developed various tools for optimization. Marketing scholars have built tools for demand forecasting and finance scholars for asset pricing. But all these tools address decision domains in which problems are well-defined, the action space is established, and the causal structure reasonably understood.

The strategy field’s opportunity for contribution lies elsewhere: in helping humans solve more difficult problems in which the potential for value creation is arguably orders of magnitude larger. Our opportunity is to help build tools and processes to support Levels 3 and 4, in which the challenge is not to optimize within well-understood problems and models but to construct, evaluate, and choose among them. Level 3 requires extrapolating the effects of novel actions by drawing on evidence from other contexts and judging which analogies hold. Level 4 demands the generative construction of causal paths to envisioned futures that do not yet exist in any data. Unlike Levels 1 and 2 in which data and data analysis are the

predominant currencies for contribution to strategic decision making, at Level 3 and Level 4, relevant data are sparse and more distant. Therefore, reasons, logic, comparison, and categorization are critical to elevating decision making.

Such tools could take two forms, reflecting the different premises of our two ladders. Some will help humans do better causal reasoning directly, aiding in analogy retrieval, counterfactual evaluation, or theory construction within the causal ladder’s framework. Others may achieve Level 3 or 4 outcomes through methods that do not closely resemble causal theorizing, working within the delegation ladder’s performance-first logic. Both are worth building, and comparing their outputs would itself be informative about the nature of strategic reasoning.

Levels 3 and 4 are the tasks with which human strategists struggle most and AI’s potential to contribute remains least understood. They are also the tasks that matter most for competitive advantage, precisely because they are rare, difficult, and resistant to routinization. If our field builds effective tools for Levels 3 and 4 reasoning, we occupy terrain that no adjacent disciplines can claim. If our field fails to do so, we risk ceding the design of strategizing aids to those who lack understanding of central elements in our field: choosing what not to do as contingent factors such as competitive dynamics, organizational politics, and institutional constraints shift over time.

A focus on developing strategy support tools also provides an opportunity for new dependent variables in strategy research. For decades, the field has organized itself around explaining variance in firm performance: a distant and noisy outcome that is often poorly matched to the mechanisms we study. Whereas the predictive accuracy of Level 1 strategizing can be rather easily verified against realized outcomes, evaluations of Level 4 strategizing may take years to materialize or may never materialize if the strategy is not pursued. The formation of AI-infused strategic decision-making tools may allow us to study strategic reasoning more directly to examine intermediate measures such as the breadth and novelty of options generated, the validity of analogies retrieved, the coherence of causal models constructed, or the calibration of confidence.

4.3 Infrastructure for Cumulative Progress

Progress in enhancing strategic decision making through AI requires infrastructure. Two forms seem essential: software systems that support and enhance Level 3 and Level 4 strategic reasoning and benchmarks that evaluate their performance.

We suggest that the development of strategic decision-making software and its effect on managers, firms, and competition should emerge as a new form of academic contribution in the field. In computer science, academic papers document the design, deployment, and refinement of significant software systems, treating the artifact and the knowledge extracted from building it as inseparable. For this form of research to become cumulative, critical norms must develop: transparency in code and design rationales, versioning as systems evolve, and open documentation of what failed and what improved.

What distinguishes strategy software experiences from computer science is the unique nature of theoretical grounding. The artifacts that scholars design must encode strategic constructs and processes—causal models, competitive interactions, organizational constraints—and be evaluated on the performance of tasks that matter to strategists. A system that simply generates fluent text is not a strategy contribution. A system that purports to help managers construct better theories of value creation and is evaluated on whether it does so is.

Benchmarks are, therefore, the second pillar. Progress in the broader field of AI is propelled by shared tests that track capabilities, reveal limitations, and define targets for improvement. Strategy currently lacks such infrastructure. We have no shared benchmarks for evaluating AI on the tasks that define strategic reasoning: generating options, validating analogies, framing problems, constructing causal theories, and adapting with feedback and failure. Careful benchmarking will enable the strategy community to accumulate knowledge and speak with authority about what AI can and cannot do. Strategy scholars are positioned to define what good strategizing means in computational terms. If we do not take this on urgently, other fields will, and their definitions may not reflect the complexity and judgment that characterize strategic reasoning at its best.

4.4 Governance, Training, and Field Identity

Although debates about the potential limits of AI remain active (McBride 2026, *Strategy Science* (this issue), Chatterji et al. 2026, *Strategy Science* (this issue), Felin and Holweg 2024), AI capabilities continue to expand, bringing continued evolution in the human role in strategizing. The human role will shift from performing analysis to framing and prompting analysis, evaluating resulting recommendations, and authorizing the commitments that follow. For the foreseeable future, humans will play a central role in Levels 3 and 4 strategizing, in model building and evaluation and in the composition and consideration of counterfactual reasoning. Yet the form and intricacy of AI collaboration will continue to evolve.

As AI's role in strategy expands, the human role will look increasingly like governance: specifying objectives and constraints, adjudicating among competing causal models, building commitment and legitimacy, and bearing accountability for outcomes that cannot be delegated to algorithms. Strategy research must explore not only AI-assisted tools but governance designs: override rights, escalation rules, accountability regimes that preserve judgment when recommendations emerge from opaque systems. Moreover, if, as Neshenko and Ryall suggest, AI optimizes so effectively within existing models that humans lose incentive to innovate, the result may be stability without progress: an equilibrium in which AI constrains rather than enhances long-run value creation. In any AI-augmented strategy future, the human role is not residual. It is the locus of judgment and responsibility that makes strategic decision making meaningful.

The agenda we propose carries implications for training scholars in strategy. Strategy doctoral programs need not produce engineers, but they should produce scholars capable of collaborating with technical colleagues, posing strategy-grounded design requirements for the AI-assisted tools, and evaluating sociotechnical systems credibly. Departments and journals must develop criteria recognizing benchmark creation, artifact building, and rigorous evaluation as legitimate academic outputs that can advance our understanding of effective strategic decision making.

The development of AI highlights an opportunity that the field has long ignored. Strategy has occupied an uncomfortable position: too applied for pure social science, too conceptual for engineering, too particular for economics, too focused on performance for organizational behavior. This position now becomes a source of advantage. A field that engages seriously with building and evaluating strategic decision-making tools—grounded in Levels 3 and 4 strategizing, attentive to rivalry and institutional constraint, rigorous about what works—has a distinctive role that no adjacent discipline can fill. If we as a field choose to watch from the sidelines, we will find our relevance diminished.

We close with an invitation and a responsibility. The invitation: Tools now exist to make progress on problems long deferred. We can build aids for counterfactual evaluation and theory construction, test whether they work, create benchmarks, and document software experiences. This is not a departure from strategy's heritage; it is a return to our founding ambition with enhanced capabilities to address the task.

The responsibility: If strategy scholars do not define what good strategizing means precisely enough to be encoded, tested, and refined, other disciplines will. They will embed thinner conceptions of strategy into the tools managers use, leaving the strategy field to critique from the sidelines what it did not help design.

The future of strategy will be shaped by those who build the tools strategists use. The tools that currently dominate strategy teaching, such as Porter's five forces or tools for resource assessment, reflect intense and extended scholarly work. Our hope is that the tools of tomorrow will be more powerful, providing more comprehensive strategic guidance. We invite the field to ensure that these tools reflect the depth and discernment that effective strategizing demands.

References

- Acemoglu, D. 2025. The simple macroeconomics of AI. *Economic Policy* **40**(121) 13–58.
- Allen, R., R. McDonald. 2026. How well can AI do strategy? Empirical benchmarking using strategy simulations. *Strategy Science* (this issue).
- Andrews, K. R. 1971. *The Concept of Corporate Strategy*. Dow Jones-Irwin, Homewood, IL.
- Ansoff, H. I. 1965. *Corporate Strategy*. McGraw-Hill, New York.
- Barney, J. 1991. Firm resources and sustained competitive advantage. *Journal of Management* **17**(1) 99–120.
- Bettis, R. A., S. Ethiraj, A. Gambardella, C. Helfat, W. Mitchell. 2016. Creating repeatable cumulative knowledge in strategic management: A call for a broad and deep conversation among authors, referees, and editors. *Strategic Management Journal* **37**(2) 257–261.
- Boussioux, L., J. N. Lane, M. Zhang, V. Jacimovic, K. R. Lakhani. 2024. The crowdless future? How generative AI is shaping the future of human crowdsourcing. *Organization Science* **35**(5) 1589–1607.
- Camuffo, A., S. Kazemi, A. Pandey. 2026. Beyond black boxes: Designing and testing agentic AI systems for strategy. *Strategy Science* (this issue).
- Chandler, A. D., Jr. 1962. *Strategy and Structure*. MIT Press, Cambridge, MA.
- Chatterji, A., F. A. Csaszar, J. Evans, T. Felin, J. Hullman, K. Lakhani, M. Sako. 2026. Can AI do strategy: A dialogue and debate. *Strategy Science* (this issue).
- Csaszar, F. A. 2018. What makes a decision strategic? Strategic representations. *Strategy Science* **3**(4) 606–619.
- Csaszar, F. A. 2025. Unbounding rationality: Why AI is a fundamental issue for strategy. Preprint, submitted September 8, <https://doi.org/10.2139/ssrn.5454634>.
- Csaszar, F. A., H. Ketkar, H. Kim. 2024. Artificial intelligence and strategic decision-making: Evidence from entrepreneurs and investors. *Strategy Science* **9**(4) 322–345.
- Csaszar, F. A., D. Laureiro-Martínez. 2018. Individual and organizational antecedents of strategic foresight: A representational approach. *Strategy Science* **3**(3) 513–532.
- Csaszar, F. A., J. Ostler. 2020. A contingency theory of representational complexity in organizations. *Organization Science* **31**(5) 1198–1219.
- Csaszar, F. A., T. Steinberger. 2022. Organizations as artificial intelligences: The use of artificial intelligence analogies in organization theory. *Academy of Management Annals* **16**(1) 1–37.
- Dell’Acqua, F., E. McFowland, E. Mollick, H. Lifshitz-Assaf, K. C. Kellogg, S. Rajendran, L. Kraye, K. R. Candelon, K. Lakhani. 2023. Navigating the jagged technological frontier: Field experimental evidence of the effects of AI on knowledge worker productivity and quality. Working Paper No. 24-013, Harvard Business School, Cambridge, MA.
- Doshi, A. R., J. J. Bell, E. Mirzayev, B. S. Vanneste. 2025. Generative artificial intelligence and evaluating strategic decisions. *Strategic Management Journal* **46**(3) 583–610.
- Felin, T., M. Holweg. 2024. Theory is all you need: AI, human cognition, and causal reasoning. *Strategy Science* **9**(4) 346–371.
- Felin, T., T. R. Zenger. 2017. The theory-based view: Economic actors as theorists. *Strategy Science* **2**(4) 258–271.
- Gavetti, G., D. A. Levinthal, J. W. Rivkin. 2005. Strategy making in novel and complex worlds: The power of analogy. *Strategic Management Journal* **26**(8) 691–712.

- Kanis, T., J. Mann, J. Stumpf-Wollersheim. 2026. AI-augmented strategic decision-making under time constraints: An experimental study on mental representations and strategic foresight. *Strategy Science* (this issue).
- Krakowski, S., J. Luger, S. Raisch. 2023. Artificial intelligence and the changing sources of competitive advantage. *Strategic Management Journal* **44**(6) 1425–1452.
- Kumar, S., X. Qu, T. Tong. 2026. The role of predictions in acquisition decision-making: The strategic value of AI-driven foresight. *Strategy Science* (this issue).
- Lee, G. K. 2024. Transportability analysis: How methods from the fields of machine learning and artificial intelligence can be useful to management research and cumulative theory testing. *J. Management Sci. Rep.* **2**(2) 179–197.
- Lee, G. K. 2026. Governance of artificial intelligence: Public policy and self-regulatory frameworks for consequential decision-making. F. A. Csaszar, N. Jia, eds., *Handbook of Artificial Intelligence and Strategy*. Edward Elgar Publishing, Cheltenham, UK, 378–387.
- Lee, G. K., R. Bettis. 2023. Structural causal modeling of managerial interventions: What if managers had not intervened by doing this? *Strategy Science* **8**(1) 24–43.
- McBride, R. 2026. Mean articulation machines. *Strategy Science* (this issue).
- Moore, R. K. 2005. Results from a survey of attendees at ASRU 1997 and 2003. *Proc. Interspeech*. https://www.isca-archive.org/interspeech_2005/moore05_interspeech.pdf.
- Moravec, H. 1988. *Mind Children: The Future of Robot and Human Intelligence*. Harvard University Press, Cambridge, MA.
- Neshenko, N., M. Ryall. 2026. When AI does strategy: Learning, good times, lock-in, and human-driven strategic renewal. *Strategy Science* (this issue).
- Nickerson, J. A., T. R. Zenger. 2004. A knowledge-based theory of the firm—the problem-solving perspective. *Organization Science* **15**(6) 617–632.
- OpenAI. 2025. Announcing OpenAI pioneers program. April 9, <https://openai.com/index/openai-pioneers-program/>.
- Pearl, J. 2000. *Causality: Models, Reasoning, and Inference*. 2nd ed. Cambridge University Press, New York.
- Pearl, J. 2021. Radical empiricism and machine learning research. *J. Causal Inference* **9**(1) 78–82.
- Porter, M. E. 1980. *Competitive Strategy*. Free Press, New York.
- Porter, M. E. 1986. *Competition in Global Industries*. Harvard Business School Press, Boston, MA.
- Porter, M. E. 1990. New global strategies for competitive advantage. *Planning Rev.* **18**(3) 4–14.
- Porter, M. E. 1996. What is strategy? *Harvard Business Review* **74**(6) 61–78.
- Rumelt, R. P. 2011. *Good Strategy/Bad Strategy: The Difference and Why It Matters*. Crown Business, New York.
- Sen, P., M. Workiewicz, P. Puranam. 2026. Can LLMs aid analogical reasoning for strategic decisions? A comparative study. *Strategy Science* (this issue).
- Shaver, J. M. 2020. Causal identification through a cumulative body of research in the study of strategy and organizations. *Journal of Management* **46**(7) 1244–1256.
- Shmueli, G. 2010. To explain or to predict? *Statistical Science* **25**(3) 289–310.
- Simon, H. A. 1969. *The Sciences of the Artificial*. MIT Press, Cambridge, MA.
- Simon, H. A. 1993. Strategy and organizational evolution. *Strategic Management Journal* **14**(S2) 131–142.

- Song, P., P. Han, N. Goodman. 2026. Large language models reasoning failures. *Trans. Machine Learn. Res.* Forthcoming.
- Sutton, R. 2019. The bitter lesson. March 13, <http://www.incompleteideas.net/IncIdeas/BitterLesson.html>.
- Teece, D. J. 2018. Business models and dynamic capabilities. *Long Range Planning* **51**(1) 40–49.
- UC Berkeley SkyLab. 2025. Hello from LMArena: The community platform for exploring frontier AI. June 23, <https://arena.ai/blog/hello-from-lmarena/>.
- Van den Steen, E. 2018. The strategy in competitive interactions. *Strategy Science* **3**(4) 574–591.
- Yang, M.-J., T. Felin, T. Zenger. 2026. Patchwork, misattention, and pandering: Theory and evidence on (generative) AI biases in strategy. Working paper, University of Colorado, Boulder.
- Zott, C., R. Amit. 2010. Business model design: An activity system perspective. *Long Range Planning* **43**(2–3) 216–226.