

Leadership After AI: What Will Still Make Humans Indispensable?

Ana Fernanda

Universidad Europea del Atlántico, Santander, Cantabria, Spain

ABSTRACT

Artificial intelligence now performs many of the analytical and operational tasks that once defined managerial work: forecasting, scheduling, performance evaluation, and even elements of strategic scenario planning. This paper asks a question that follows directly from that shift: once machines handle information processing at a scale and speed no human can match, what remains for human leaders to do? Rather than treating this as a question of job loss or job survival, the paper reframes it as a question about the composition of leadership itself. Drawing on and critically comparing agency theory, stewardship theory, upper echelons theory, dynamic capabilities, institutional theory, stakeholder theory, and several strands of leadership scholarship including transformational, authentic, adaptive, complexity, responsible, and servant leadership, the paper argues that these theories were built for an era in which cognition and judgment were bundled together in a single human actor. AI unbundles them. What remains once information processing is delegated to machines is a residue of ethical judgment, contextual interpretation, meaning-making, institutional stewardship, and the capacity to hold ambiguity long enough for a considered choice to emerge. The paper develops an original conceptual model, the Human Leadership Advantage Framework, which organizes sixteen interdependent leadership dimensions into four clusters: moral and ethical grounding, relational and cultural capacity, institutional and strategic stewardship, and adaptive and generative capability. The framework explains why these dimensions do not simply survive AI adoption but become more consequential as routine cognition is automated. The paper closes with implications for organizational design, executive development, board oversight, and leadership education, together with a research agenda for scholars studying leadership in AI-saturated organizations.

Keywords: Artificial intelligence, strategic leadership, executive judgment, responsible leadership, upper echelons theory, human-AI collaboration, leadership development, corporate governance

INTRODUCTION

For most of the last century, leadership theory and management practice shared an assumption so basic that it rarely needed to be stated: the person who processes the information is the person who decides, and the person who decides is the person who leads. Executives read reports, weighed evidence, and issued judgments. Managers scheduled, monitored, and corrected. The cognitive work of sensing an organization's environment and the social work of directing its people sat inside the same set of human minds. Artificial intelligence breaks that bundle apart. Large-scale prediction, pattern detection, and even generation of strategic options can now be performed by systems that never sit in a boardroom, never lose a night's sleep before a difficult call, and never have to look an employee in the eye after announcing a layoff (Krakowski et al., 2023; Raisch & Krakowski, 2021).

The immediate reaction to this shift, in both the popular press and a portion of the academic literature, has been to ask whether AI will replace managers and executives. That question, framed as a horse race between human and machine, tends to produce one of two unhelpful answers. Either AI is treated as a tool

that changes nothing essential about leadership, a position that ignores how thoroughly algorithmic systems already shape hiring, performance evaluation, resource allocation, and strategic forecasting (Hillebrand et al., 2025), or AI is treated as an inevitable successor to human authority, a position advanced provocatively by Van Quaquebeke and Gerpott (2023), who argue that algorithms may in some respects serve employees' psychological needs more consistently than human leaders do, precisely because they are not subject to fatigue, mood, or self-interest.

This paper takes a third position. The question worth asking is not whether AI will take over leadership, but which parts of leadership were never reducible to information processing in the first place, and how those parts change in character once the information-processing burden is lifted from human shoulders. Framed this way, the inquiry shifts from a contest over territory to an analysis of composition. Leadership has always contained at least two kinds of work: work that consists of gathering, weighing, and calculating, and work that consists of choosing what the calculation should serve, building the trust that makes a choice stick, and carrying responsibility for what happens after the choice is made. AI is extraordinarily good at the first kind of work. It has no established capacity for the second.

The paper proceeds as follows. Section 2 describes leadership as it was practiced and theorized before AI became embedded in managerial infrastructure, establishing a baseline against which change can be measured. Section 3 reviews the literature on AI and leadership that has accumulated rapidly since around 2018, and section 4 examines the theoretical traditions that inform, and are strained by, this literature. Section 5 traces how AI is reshaping the actual content of leadership work, and sections 6 and 7 separate the capabilities AI can credibly replicate from those that remain, on present evidence, distinctly human. Section 8 introduces the paper's central conceptual contribution, the Human Leadership Advantage Framework. Sections 9 through 12 draw out implications for organizational design, executive education, and governance, before section 13 proposes a research agenda and section 14 concludes.

LEADERSHIP IN THE PRE-AI ERA

Understanding what AI changes requires a clear picture of what came before it. Twentieth-century management theory, from Fayol's administrative functions through Drucker's writing on the effective executive, treated the manager as an information hub: someone who collected data from subordinates, customers, and markets, synthesized it, and issued directives. Mintzberg's later empirical work complicated this picture by showing that managerial work was fragmented, interrupted, and heavily oral rather than neatly analytical, but the underlying assumption persisted that a competent human mind, suitably placed in an organizational hierarchy, was the necessary vessel for both sensing and deciding.

Strategic leadership theory built on this assumption. Hambrick and Mason's (1984) upper echelons theory proposed that organizational outcomes are, in substantial part, a reflection of the cognitive frames, experiences, and values of top managers, because strategic choices under uncertainty are filtered through the idiosyncratic perceptual apparatus of the people making them. The theory's power came from treating executive cognition itself, with all its bounded rationality and bias, as a variable worth studying. Executives were not simply calculators; they were interpreters, and their interpretations mattered because no other mechanism in the organization performed that interpretive function at scale.

Leadership theory ran on a parallel track. Transformational leadership asked how leaders articulate vision and inspire followers toward outcomes beyond narrow self-interest. Authentic leadership asked how leaders' self-awareness and moral perspective shape follower trust (Avolio & Gardner, 2005). Servant leadership asked how a leader's orientation toward the needs of followers, rather than toward personal advancement, shapes organizational culture. Each of these frameworks, whatever their differences, located leadership's distinctive value in the relationship between a human leader's inner life, character, or moral stance and the behavior of followers who could sense and respond to that inner life. None of them

anticipated a world in which the analytical labor historically used to justify a leader's authority, the labor of having read more, calculated more, and forecast more accurately than anyone else in the room, could be delegated to a system with no inner life at all.

It is worth being honest about the limitations of this pre-AI baseline. Human leadership in this period was expensive, slow, and unevenly distributed. Executive judgment was shaped by cognitive biases that upper echelons research has documented extensively, from overconfidence to anchoring to motivated reasoning. Boards struggled to evaluate leadership quality using anything more rigorous than reputation and past performance. The romantic account of the wise, decisive human leader was always, to some degree, a story organizations told about themselves rather than an accurate description of how decisions were actually made. AI does not intrude on a pristine model of human leadership; it intrudes on a model that was already imperfect, unevenly applied, and dependent on organizational myth-making as much as on demonstrated capability.

LITERATURE REVIEW

Scholarly attention to AI's effect on management and leadership has grown quickly, though it remains fragmented across several distinct conversations that do not always speak to one another. One stream, exemplified by Raisch and Krakowski (2021), examines the tension between automation and augmentation, arguing that AI adoption in organizations tends to produce a paradox: firms that pursue automation to replace human judgment often find that residual human tasks become more complex and more consequential, not less, while firms that pursue augmentation strengthen human capability but face slower returns. This automation-augmentation paradox is a useful corrective to the assumption that AI adoption is a simple substitution process.

A second stream examines AI's effect on competitive advantage and strategy formulation. Krakowski, Luger, and Raisch (2023), using chess as a controlled empirical setting, show that AI adoption triggers substitution and complementation dynamics simultaneously: some traditional human capabilities lose their competitive value entirely, while new capabilities built around human-machine interaction become sources of durable advantage. Their finding that centaur players, humans working with AI assistance, developed new and persistent forms of skill differentiation is directly relevant to leadership, because it suggests that the future advantage of human leaders will not lie in outperforming AI at analysis but in developing skill at working with and around AI systems in ways that are difficult to imitate.

A third stream, more integrative in ambition, is represented by Hillebrand, Raisch, and Schad's (2025) review, which synthesizes research on human-AI collaboration and algorithmic management into a single systems-level account of what it means to manage with AI. Their central observation, that human-AI collaboration research has examined how humans use AI to manage while algorithmic management research has examined how humans are managed by AI, points to an underappreciated duality: leaders today are simultaneously users of AI and subjects of AI-mediated evaluation and control. This dual position, rarely addressed in a single framework before their review, has direct implications for questions of trust, autonomy, and legitimacy that earlier leadership theory did not need to consider.

A fourth and more provocative stream challenges the assumption that leadership itself is immune to automation. Van Quaquebeke and Gerpott (2023) argue directly against what they call the romanticization of human leadership, contending that algorithms, properly designed, could provide more consistent recognition, fairer evaluation, and more responsive support than fallible, inconsistent human leaders who are subject to fatigue, bias, and self-interest. Their argument is not that AI leadership is already superior in practice, but that the case for human leadership's permanence has rested on unexamined assumptions rather than demonstrated advantage. This challenge deserves to be taken seriously rather than dismissed, because much of the literature on AI and leadership, including some contributions reviewed in recent

systematic reviews of the field, has tended to assert human indispensability as a premise rather than defend it as a conclusion.

Jarrahi's (2018) earlier account of human-AI symbiosis offers a partial answer to that challenge by identifying where the complementarity actually lies. Organizational decision-making, Jarrahi argues, is typically characterized by uncertainty, complexity, and equivocality, and these three conditions call for different cognitive strengths. AI's computational capacity gives it an advantage in addressing complexity, meaning the sheer volume of variables and interactions in a decision problem. Humans, drawing on intuition and holistic judgment, retain an advantage in addressing uncertainty and equivocality, meaning situations where the relevant variables are unknown or where the same facts can be reasonably interpreted in more than one way. Strategic leadership, almost by definition, deals disproportionately in uncertainty and equivocality rather than in complexity alone, which is one reason the symbiosis argument, rather than the substitution argument, has more traction at the executive level than at the operational level.

Finally, systematic reviews of AI-driven leadership research point to a persistent gap: leadership theory has been slow to catch up with AI's practical penetration into organizational decision-making, and few studies have moved past documenting AI's technical capabilities or employees' behavioral reactions to actually revise leadership theory itself (Quaquebeke & Gerpott, 2023, as cited in subsequent systematic reviews). The literature, in other words, has described the disruption more thoroughly than it has resolved what leadership theory should become in response. That gap is the space this paper attempts to occupy.

Practitioner literature, produced outside the peer-reviewed academic process but influential in shaping how executives themselves think about the problem, tells a broadly consistent story, though it should be read as a distinct evidentiary category rather than treated as equivalent to peer-reviewed scholarship. Iansiti and Lakhani (2020) describe firms whose competitive advantage increasingly depends on algorithmic infrastructure that operates at a scale no human hierarchy could replicate, while cautioning that the firms succeeding with this infrastructure are not the ones that removed human leadership from the loop but the ones that redesigned human leadership around it. The World Economic Forum's (2025) Future of Jobs Report, drawing on survey data from more than a thousand employers, projects substantial disruption to existing skill requirements over the remainder of this decade, while simultaneously identifying leadership and social influence, resilience, and analytical thinking applied to non-routine problems as skills rising rather than falling in demand. Read carefully, this practitioner evidence does not contradict the academic literature reviewed above; it corroborates the pattern the automation-augmentation paradox predicts, in which the skills displaced by AI adoption are concentrated in routine analysis while the skills retained and increasingly rewarded are concentrated in judgment, relationship, and adaptive capability.

THEORETICAL FOUNDATIONS

No single existing theory is built to answer the question this paper poses, but several offer partial and sometimes contradictory resources.

Agency theory begins from the assumption that managers, as agents, may pursue interests that diverge from those of shareholders as principals, and that governance mechanisms exist to monitor and align these interests. AI complicates agency theory by introducing a third actor, the algorithmic system, whose interests cannot be specified in the same way a human agent's can, but whose outputs shape decisions attributed to human agents. When an executive defers to an algorithmic recommendation, agency theory has no settled answer to whether accountability shifts, dilutes, or remains fully with the human who acted on the recommendation.

Stewardship theory, developed by Davis, Schoorman, and Donaldson (1997) partly as a corrective to agency theory's pessimism about managerial motivation, assumes that managers can act as stewards of organizational and stakeholder interests rather than as self-interested agents requiring constant monitoring. This theory sits closer to the argument developed later in this paper, because stewardship depends on qualities, identification with the organization, long-term orientation, and intrinsic motivation toward collective rather than personal goals, that have no obvious algorithmic equivalent. An AI system has no stake in an organization's survival beyond the objective function it was built to optimize, and it cannot be said to identify with an institution in the psychological sense stewardship theory requires.

Upper echelons theory (Hambrick & Mason, 1984), as discussed above, treats executive cognition as a variable that shapes strategic choice. AI adoption does not eliminate this variable; it changes what it operates on. If routine analysis is delegated to algorithmic systems, the residual cognitive work performed by executives becomes concentrated in interpretation, framing, and the selection of which algorithmic outputs to trust, amplify, or override. Upper echelons theory therefore needs updating rather than replacement: the executive's cognitive frame still matters, but it now operates one level removed from raw data, filtered through the design choices embedded in whatever AI systems the organization has adopted.

Dynamic capabilities theory (Teece et al., 1997) describes how firms sense opportunities, seize them through resource reconfiguration, and transform their asset base over time. AI is often framed, correctly, as a resource that firms must integrate into their dynamic capabilities. Less often discussed is that dynamic capabilities theory's three core processes, sensing, seizing, and transforming, are themselves leadership functions before they are firm-level capabilities. An organization's capacity to sense weak signals in its environment, to commit to a course of action despite incomplete information, and to reconfigure itself once a new direction is chosen, still depends on human executives who can tolerate the ambiguity that precedes commitment. AI can supply better sensing in the narrow sense of data collection and pattern recognition; it cannot, on present evidence, supply seizing in the sense of committing an organization's identity and resources to an uncertain future.

Institutional theory adds a further dimension that the AI-and-leadership literature has largely neglected: organizations survive not only by being efficient but by being seen as legitimate within their institutional environment. Leaders perform a legitimacy-conferring function, standing before employees, investors, regulators, and communities as the visible, accountable face of decisions that were often made through distributed, partly automated processes. An algorithm cannot perform this legitimating function, because legitimacy is a social attribution that depends on the perceived presence of a responsible, accountable moral agent, not on the accuracy of a decision.

Stakeholder theory (Freeman, 1984) reinforces this point by insisting that organizations owe consideration to a range of parties beyond shareholders, each with different, sometimes conflicting claims. Balancing these claims requires judgment about competing values, not only optimization against a single objective function. AI systems can be built to weigh multiple objectives, but the weighting itself, deciding how much a community's interest should count against a shareholder's interest in a specific, contested situation, is a value judgment that stakeholder theory locates squarely in the domain of leadership responsibility.

Among leadership theories specifically, transformational and authentic leadership emphasize the leader's inner coherence and capacity to inspire, both of which depend on followers perceiving the leader as a fellow human subject to the same stakes they face. Adaptive leadership (Uhl-Bien & Arena, 2018) and complexity leadership theory (Uhl-Bien et al., 2007) offer perhaps the most directly applicable resources, because they already treat leadership as distributed across a system rather than concentrated in a single decision-maker, and they distinguish between administrative leadership, which manages existing

structures, and adaptive or enabling leadership, which fosters the emergence of new solutions under conditions of complexity. This distinction maps unusually well onto the AI-augmented organization: administrative leadership is precisely the terrain most exposed to automation, while adaptive and enabling leadership, the work of creating conditions under which genuinely new responses can emerge, remains a distinctly human contribution because it depends on judgment about which tensions in a system are productive and which are destructive, a judgment complexity leadership theory treats as inherently contextual and non-algorithmic.

Responsible leadership theory (Voegtlin et al., 2012) extends this line of thought by framing leadership as a relational and ethical process that must answer to a wide circle of stakeholders and to society more broadly, not merely to the organization's internal hierarchy. This framing anticipates much of what current AI governance debates are grappling with: who answers for an AI-influenced decision, and to whom. Human-centered AI scholarship (Shneiderman, 2020) approaches the same problem from the design side, arguing that AI systems should be built to keep humans in meaningful control, reliable, safe, and trustworthy, rather than autonomous in ways that erode human oversight. Read together, responsible leadership theory and human-centered AI design converge on a shared conclusion: the case for keeping humans in leadership roles is not merely descriptive, an observation that humans currently hold these roles, but normative, an argument that accountability structures require an identifiable human locus of responsibility.

These theories do not agree with one another in every respect. Agency theory's suspicion of managerial motive sits uneasily beside stewardship theory's more generous assumptions. Upper echelons theory's focus on cognitive bias sits uneasily beside complexity leadership theory's emphasis on distributed, emergent leadership rather than individual cognition. What they share, despite these tensions, is an assumption this paper treats as the central theoretical gap exposed by AI: each theory assumes that whoever performs the cognitive work of a decision is also the party who bears its social and moral weight. AI severs that assumption for the first time in the history of management thought, and none of these theories, built before the severance was conceivable, offers a ready-made account of what happens next.

AI AND THE TRANSFORMATION OF LEADERSHIP WORK

The practical effect of AI on managerial and executive work can be described along three dimensions: the tasks being automated, the tasks being augmented, and the entirely new tasks being created.

Automation is furthest along in functions with clear inputs, clear outputs, and abundant historical data: demand forecasting, workforce scheduling, credit and risk scoring, candidate screening, and increasingly, first-draft strategic scenario generation. These were once treated as core managerial competencies, tested in MBA curricula and used as proxies for executive readiness. Their delegation to algorithmic systems removes a set of tasks that used to function as a training ground and a credentialing mechanism for rising managers, which has consequences for leadership development discussed in section 10.

Augmentation is visible where AI extends human capacity without fully substituting for it. Executives now have access to scenario models that can process combinations of variables no unaided human could hold in mind simultaneously, and to natural-language systems that can summarize, translate, and draft at a pace that changes the rhythm of executive communication. Iansiti and Lakhani (2020) describe this shift as a move from firms organized around human decision hierarchies to firms organized around algorithmic processes that humans supervise and adjust at key decision points, a change in the executive's role from primary calculator to a combination of calibrator, interpreter, and final arbiter.

New tasks are also emerging that had no clear precedent in pre-AI organizations. Someone must decide which decisions an organization is willing to delegate to algorithmic systems and which it is not. Someone must audit algorithmic outputs for bias, drift, and failure modes that are not obvious from the outside. Someone must communicate to employees, customers, and regulators why a given decision was reached when the underlying process involved a model whose internal logic even its designers cannot fully explain. These tasks, oversight, boundary-setting, and explanation, did not exist in this form before AI, and they sit closer to the judgment and legitimacy functions described in section 4 than to the analytical functions AI is displacing.

Taken together, these three dimensions suggest that AI is not simply removing work from leaders; it is redistributing the ratio of analytical to judgment-based work within the leadership role, and it is adding an entirely new category of meta-level responsibility for governing the analytical systems themselves.

CAPABILITIES AI CAN REPLICATE

Honesty about AI's genuine capabilities is a precondition for a credible argument about what remains human. AI systems, at current and near-term levels of capability, can replicate or exceed human performance in several functions long associated with managerial competence.

Pattern detection across large, high-dimensional datasets is an area where AI's advantage is not marginal but categorical; no human manager can hold thousands of variables in working memory simultaneously, and few would want to try. Forecasting under stable, well-specified conditions, where historical patterns are a reasonable guide to future outcomes, is another area of clear AI strength, evident in applications ranging from demand planning to credit risk assessment. Consistency of application is a third area: an algorithmic evaluation process, once properly calibrated, applies the same criteria to every case without the day-to-day mood variation, fatigue, or interpersonal favoritism that affects even well-intentioned human evaluators. Speed of synthesis, the ability to summarize a large volume of documents, data, or communications into a workable brief, is a fourth area where generative AI systems already outperform most humans on raw throughput.

These are genuine capabilities, and organizations that fail to adopt them where appropriate will likely find themselves at a durable competitive disadvantage, a conclusion consistent with Krakowski et al.'s (2023) finding that substitution dynamics can permanently erode the value of purely human analytical skill in domains where AI performs reliably. The temptation to defend human leadership by understating what AI can do should be resisted, both because it is empirically indefensible and because it distracts from the more interesting and more defensible argument about what AI cannot do.

CAPABILITIES THAT REMAIN UNIQUELY HUMAN

Set against the capabilities described above, a distinct set of functions has, at least on current evidence, no credible algorithmic equivalent, not because AI is technically incapable of producing outputs that resemble these functions, but because the functions depend on properties, embodied stakes, social attribution of moral responsibility, and lived participation in the consequences of a decision, that are not reducible to computation regardless of how sophisticated the computation becomes.

Ethical judgment under genuine moral conflict, where two legitimate values cannot both be honored and a person must choose which to sacrifice, is a central example. AI systems can be trained to apply ethical frameworks consistently, but consistency is not the same as judgment, because judgment in these situations includes owning the choice, being accountable for having made it, and living with its consequences in a way that shapes future choices. Moral courage, the willingness to take a costly action because it is right rather than because it is optimal by some measurable criterion, similarly depends on a subject who can bear cost, and an algorithm bears no cost in any sense relevant to courage.

Contextual intelligence, the capacity to read a specific organizational, cultural, and historical situation and to recognize why a generally sound principle might not apply here, in this instance, draws on tacit, embodied experience that resists full codification. Trust formation between a leader and followers depends on followers perceiving that the leader has something at stake, reputation, relationship, or shared fate, that gives their commitments credibility; an entity that cannot lose anything cannot be trusted in the full sense that human trust relationships require, only relied upon in a narrower, transactional sense.

Meaning-making, the work of framing why an organization's activity matters beyond its measurable outputs, and vision creation, the work of describing a future state that does not yet exist and inviting others to help build it, both depend on a speaker who is understood to be personally invested in the future being described. Crisis leadership draws on similar ground: in a genuine crisis, followers look for a person who will stand publicly accountable for decisions made under extreme uncertainty, not for a well-calibrated recommendation engine.

Institutional stewardship, the sense of responsibility for an organization's identity and legacy that extends beyond any single decision or fiscal quarter, and long-term thinking that willingly trades short-term optimization for durability, both depend on an actor capable of caring about an institution's future in a way that survives changes in ownership, technology, or leadership personnel. Cultural leadership, shaping and transmitting the norms that hold an organization together between formal decisions, and conflict resolution among people with competing interests and interpretations, both depend on the social legitimacy that comes from being a fellow participant in the culture being led, not an external system applied to it.

Responsible innovation, weighing the second- and third-order social consequences of a new technology or product before those consequences are fully knowable, requires a form of anticipatory ethical imagination that current AI systems, trained on historical data, are structurally poor at, since by definition the harms of genuinely novel innovations have no historical precedent to learn from. Adaptive learning at the level of organizational identity, not merely updating a model's parameters but reconsidering what an organization is for, is likewise a function of reflective self-awareness that current AI architectures do not possess in any philosophically defensible sense.

The common thread running through this list is that each capability depends on the leader being a stakeholder in the outcome, someone who can be praised, blamed, trusted, or held to account in a way that changes the leader's own future. AI, however capable at generating outputs that mimic the surface form of these functions, a sympathetic-sounding message, a plausible ethical justification, a coherent vision statement, cannot occupy this position of stakeholdership. This is the foundation on which the framework in the next section is built.

THE HUMAN LEADERSHIP ADVANTAGE FRAMEWORK

The preceding sections identify sixteen leadership dimensions repeatedly cited across the literature as candidates for human indispensability: ethical judgment, moral courage, contextual intelligence, institutional stewardship, organizational legitimacy, vision creation, meaning-making, strategic imagination, empathy, trust formation, conflict resolution, cultural leadership, adaptive learning, crisis leadership, responsible innovation, and long-term thinking. Treated as an unstructured list, these dimensions risk becoming exactly what critics such as Van Quaquebeke and Gerpott (2023) warn against: a comforting inventory asserted rather than explained. The Human Leadership Advantage Framework organizes these sixteen dimensions into four interdependent clusters and specifies the mechanism by which each cluster becomes more, not less, consequential as AI assumes a greater share of analytical work.

Cluster One: Moral and Ethical Grounding

This cluster contains ethical judgment, moral courage, and responsible innovation. Its unifying mechanism is stakeholderism: these capabilities require an actor who bears the consequences of a choice in a way that can be recognized, praised, or sanctioned by others. As AI absorbs routine decisions, the decisions that remain for human leaders are disproportionately the ones that involve genuine value conflict, precisely because AI systems are typically deployed first in domains with clear, measurable objectives, leaving ambiguous, values-laden decisions in human hands by default. Moral and ethical grounding therefore does not merely persist under AI adoption; it becomes concentrated, since the population of decisions reaching human leaders becomes progressively weighted toward the hardest cases.

Cluster Two: Relational and Cultural Capacity

This cluster contains empathy, trust formation, conflict resolution, and cultural leadership. Its unifying mechanism is embodied social presence: these capabilities depend on followers perceiving the leader as a fellow participant in a shared fate, someone whose own wellbeing is bound up with the organization's outcomes. As algorithmic systems take over more of the organization's evaluative and coordinating functions, employees' need for a human counterweight, someone who can be reasoned with, appealed to, and held personally accountable, does not diminish; it intensifies, because the experience of being managed by opaque systems tends to increase, not decrease, the demand for legible human relationships elsewhere in the organization.

Cluster Three: Institutional and Strategic Stewardship

This cluster contains institutional stewardship, organizational legitimacy, long-term thinking, and strategic imagination. Its unifying mechanism is temporal and social continuity: these capabilities require a party who persists across time in a way an optimization process, tied to whatever objective function it was last configured to pursue, does not. Boards, regulators, and communities look to identifiable human leaders to answer for an organization's trajectory precisely because algorithmic systems can be replaced, retrained, or reconfigured without anyone being able to say the organization has taken responsibility for its past choices. As AI systems become more central to day-to-day operation, the felt need for a human party who can be asked, in effect, what does this organization stand for, and who answers for it, grows rather than shrinks.

Cluster Four: Adaptive and Generative Capability

This cluster contains vision creation, meaning-making, crisis leadership, and adaptive learning. Its unifying mechanism is generative imagination under genuine novelty: these capabilities are called upon precisely in situations that lack sufficient historical precedent for a data-trained system to model reliably, whether a genuinely new crisis, a genuinely new market, or a genuinely new technology whose social consequences have not yet unfolded. Where AI performs best on problems with abundant historical data, this cluster is defined by problems where historical data is, by definition, thin or absent, which is why organizations facing disruption continue to look to human leaders for direction rather than to their forecasting models.

Interdependence Across Clusters

The four clusters are not independent; failure in one typically undermines the others. A leader who possesses ethical judgment but lacks trust formation, for instance, may make defensible decisions that nonetheless fail to secure followers' cooperation. A leader who possesses institutional stewardship but lacks adaptive capability may preserve an organization's identity while allowing it to become irrelevant. This interdependence explains why the framework describes a system of leadership advantage rather than a checklist of separately valuable traits: AI adoption raises the stakes attached to each cluster simultaneously, because a failure that would once have been absorbed by distributed human judgment across many managers now falls more heavily on fewer human decision points once routine coordination

has been automated. Organizations that treat these four clusters as optional soft skills, secondary to technical and analytical competence, are likely to discover the cost of that assumption only after AI adoption has already stripped away the analytical work that used to compensate for weak judgment, trust, stewardship, or imagination.

IMPLICATIONS FOR ORGANIZATIONS

The framework carries direct implications for how organizations should structure leadership roles as AI adoption deepens. Job design at the managerial level should be explicitly reconsidered around the redistribution described in section 5: roles historically defined by analytical throughput, producing forecasts, compiling reports, monitoring compliance, should be redesigned around oversight, interpretation, and the four clusters of the framework, rather than assumed to disappear or to continue unchanged.

Succession planning and executive selection processes built around analytical credentials, quantitative test scores, technical expertise, case-competition performance, should be rebalanced toward evidence of judgment under ambiguity, demonstrated trust-building across diverse stakeholder groups, and a track record of ethical decision-making under pressure, since these are precisely the qualities the framework identifies as durable. Performance management systems for executives, similarly, should evolve beyond metrics AI can already optimize toward assessments of the qualitative clusters identified above, difficult to measure but increasingly determinative of organizational outcomes as analytical work is automated.

Organizations should also build explicit governance structures for the new category of oversight work identified in section 5: who decides which decisions are delegated to AI, who audits algorithmic outputs, and who communicates the reasoning behind AI-influenced decisions to affected parties. Leaving these questions to emerge informally risks concentrating unaccountable power in whichever technical function happens to control the organization's AI systems, a risk consistent with concerns raised in the algorithmic management literature reviewed by Hillebrand et al. (2025) about how algorithmic control can substitute for, rather than support, legitimate managerial authority.

Middle management deserves particular attention in this redesign, because it is the layer most exposed to both automation and algorithmic control simultaneously. Middle managers have traditionally performed the coordination and monitoring work that AI now automates most readily, while also serving as the human face of policies increasingly generated or informed by algorithmic systems above them. Organizations that simply thin out this layer, on the assumption that its analytical functions are no longer needed, risk losing the relational and cultural capacity, described in cluster two of the framework, that middle managers have historically supplied between senior leadership and frontline employees. A more durable redesign treats middle management roles as the organization's primary site for cultivating trust formation and conflict resolution, explicitly relieving these roles of routine analytical burden so that the human capacities the framework identifies can be exercised rather than crowded out by administrative load.

Organizations should further resist the temptation to treat AI-assisted decisions as self-legitimizing simply because they are data-driven. A recommendation generated by an algorithmic system still requires a human party willing to adopt it, explain it, and answer for it if it proves wrong. Building this requirement explicitly into decision protocols, rather than allowing algorithmic output to be treated as a neutral or final word, keeps the accountability structures that stakeholder theory and institutional theory both depend on intact even as the analytical labor behind those decisions shifts to machines.

IMPLICATIONS FOR LEADERSHIP EDUCATION

Business school curricula built around analytical mastery, financial modeling, quantitative strategy frameworks, operations optimization, face a similar reckoning. These skills remain necessary but are no longer sufficient to distinguish future leaders, since AI increasingly performs the calculations these courses once trained students to perform manually. Leadership education should shift emphasis toward the disciplined practice of judgment under moral conflict, structured exposure to genuine ethical dilemmas rather than case studies with clean analytical answers, and toward the deliberate cultivation of contextual intelligence through immersive, varied organizational experience rather than classroom simulation alone.

Case method teaching, historically built around identifying the single best analytical answer to a business problem, should give more room to cases without a defensible single answer, cases whose value lies precisely in forcing students to notice which values are in conflict and to practice owning a choice rather than calculating one. Programs should also teach students to work productively alongside AI systems, developing what Krakowski et al. (2023) describe as centaur-style skill, the capacity to combine algorithmic output with human judgment in ways that outperform either acting alone, since this hybrid competence, rather than either pure analytical skill or pure interpersonal skill, is likely to define competitive advantage among rising managers.

Executive education programs aimed at sitting leaders face a parallel task, since many current executives were trained and promoted under the older analytical model and may not have had reason to develop the judgment-centered competencies the framework identifies. Structured reflection on past decisions made under moral conflict, mentored exposure to institutional stewardship responsibilities before an executive is asked to carry them alone, and deliberate practice in communicating the reasoning behind AI-assisted decisions to skeptical stakeholders are all practical mechanisms through which existing executive development programs could begin to close this gap without waiting for a new generation of AI-native leaders to work its way up the pipeline.

POLICY AND GOVERNANCE IMPLICATIONS

Boards of directors face a distinct set of challenges in evaluating leadership quality once analytical performance can no longer serve as the primary proxy for executive competence. Board evaluation frameworks should incorporate explicit assessment of the four clusters described in section 8, seeking evidence of ethical grounding, relational trust, institutional stewardship, and adaptive capability, rather than relying on financial performance metrics that AI-assisted decision-making can increasingly deliver regardless of the human judgment quality behind them.

Regulatory attention to AI governance, reflected in frameworks such as the OECD's principles on trustworthy AI and the European Union's approach to AI regulation, has focused heavily on technical safety, data provenance, and algorithmic transparency. Less developed, but directly implied by this paper's argument, is regulatory and governance attention to the human accountability structures that must remain in place around AI-assisted decisions, ensuring that delegation of analytical work to algorithmic systems does not become delegation of moral and legal responsibility as well. Shneiderman's (2020) human-centered AI framework, emphasizing systems that remain reliable, safe, and trustworthy while keeping humans in meaningful control, provides a design-level complement to the organizational and governance-level argument developed here: both converge on the principle that human oversight must be substantive rather than symbolic if accountability is to remain intact.

Workforce transformation data reported by the World Economic Forum (2025) reinforces the practical urgency of this argument, projecting substantial disruption to existing skill sets over the remainder of this decade alongside rising demand for what the report categorizes as human-centered skills, including

leadership and social influence, resilience, and analytical thinking applied to ambiguous rather than routine problems. This external evidence is broadly consistent with the framework's prediction that demand for the human leadership clusters described in section 8 should rise, not fall, as AI adoption deepens.

A further governance question concerns disclosure. Investors, employees, and regulators increasingly want to know not only that a firm uses AI, but where, in which decisions, and with what degree of human review at each stage. Boards could reasonably require executives to maintain and periodically disclose a decision inventory, a record of which categories of decision are fully automated, which are AI-assisted with mandatory human sign-off, and which remain entirely human, together with the reasoning behind each classification. Such a practice would give the four clusters of the Human Leadership Advantage Framework a concrete governance anchor, turning what might otherwise remain an abstract claim about human indispensability into an auditable organizational commitment that can be reviewed, challenged, and revised as AI capability continues to evolve.

FUTURE RESEARCH AGENDA

Several questions raised by this paper's conceptual argument warrant empirical investigation. Longitudinal research is needed to track whether the four clusters identified in the Human Leadership Advantage Framework do in fact predict organizational resilience and legitimacy over time as AI adoption deepens, a claim this paper advances conceptually but which requires quantitative and qualitative testing across industries and organizational forms. Comparative research across sectors with differing degrees of AI penetration could examine whether the redistribution of leadership work described in section 5 follows a common pattern or varies by industry structure, regulatory environment, and organizational culture.

Research is also needed on the psychological experience of executives whose analytical authority has been partially displaced, examining whether this displacement produces identity threat, a reallocation of professional pride toward judgment-based competence, or some combination of both, and what implications this carries for executive wellbeing and retention. The question of algorithmic accountability, how organizations and legal systems should treat responsibility when a human leader acts substantially on an AI recommendation that later proves harmful, remains unresolved in both management theory and law, and deserves sustained interdisciplinary attention connecting management scholarship with legal and philosophical work on moral responsibility.

Finally, research comparing leadership development programs that explicitly teach the judgment-centered competencies described in section 10 against traditional analytically weighted programs would help establish whether the pedagogical shift argued for here produces measurably different leadership outcomes, or whether it remains, for now, a plausible but unverified conceptual claim.

CONCLUSION

Artificial intelligence is not making leadership obsolete. It is stripping away the analytical scaffolding that leadership theory, management education, and organizational practice have long used as a convenient, measurable proxy for leadership quality, and in doing so it is exposing what was underneath that scaffolding all along. What remains once forecasting, scheduling, and routine evaluation are delegated to algorithmic systems is not a diminished residue of leadership but its most demanding core: the capacity to make ethical judgments under genuine conflict, to build trust with people who can sense whether a leader has something at stake, to steward institutions across time in a way no optimization process can, and to imagine futures for which no historical data yet exists.

The Human Leadership Advantage Framework developed in this paper organizes these capacities into four interdependent clusters and explains, rather than merely asserts, why each becomes more

consequential as AI absorbs a larger share of organizational cognition. This account carries real implications for how organizations design leadership roles, how business schools prepare future leaders, and how boards and regulators think about accountability in AI-assisted organizations. None of these implications depend on underestimating what AI can do. They depend, instead, on taking seriously a distinction the twentieth-century management literature rarely had reason to draw sharply: the distinction between processing information and being answerable for what is done with it. That distinction, largely invisible while cognition and accountability were bundled together in the same human actor, is now the whole of what leadership, after AI, still is.

REFERENCES

- Avolio, B. J., & Gardner, W. L. (2005). Authentic leadership development: Getting to the root of positive forms of leadership. *The Leadership Quarterly*, 16(3), 315-338.
- Davis, J. H., Schoorman, F. D., & Donaldson, L. (1997). Toward a stewardship theory of management. *Academy of Management Review*, 22(1), 20-47.
- Freeman, R. E. (1984). *Strategic management: A stakeholder approach*. Pitman.
- Hambrick, D. C., & Mason, P. A. (1984). Upper echelons: The organization as a reflection of its top managers. *Academy of Management Review*, 9(2), 193-206.
- Hillebrand, L., Raisch, S., & Schad, J. (2025). Managing with artificial intelligence: An integrative framework. *Academy of Management Annals*, 19(1), 343-375.
<https://doi.org/10.5465/annals.2022.0072>
- Iansiti, M., & Lakhani, K. R. (2020). *Competing in the age of AI: Strategy and leadership when algorithms and networks run the world*. Harvard Business Review Press.
- Jarrahi, M. H. (2018). Artificial intelligence and the future of work: Human-AI symbiosis in organizational decision making. *Business Horizons*, 61(4), 577-586. <https://doi.org/10.1016/j.bushor.2018.03.007>
- Krakowski, S., Luger, J., & Raisch, S. (2023). Artificial intelligence and the changing sources of competitive advantage. *Strategic Management Journal*, 44(6), 1425-1452.
<https://doi.org/10.1002/smj.3387>
- Organisation for Economic Co-operation and Development. (2019). *Recommendation of the Council on Artificial Intelligence*. OECD.
- Raisch, S., & Krakowski, S. (2021). Artificial intelligence and management: The automation-augmentation paradox. *Academy of Management Review*, 46(1), 192-210.
- Shneiderman, B. (2020). Human-centered artificial intelligence: Reliable, safe & trustworthy. *International Journal of Human-Computer Interaction*, 36(6), 495-504.
- Teece, D. J., Pisano, G., & Shuen, A. (1997). Dynamic capabilities and strategic management. *Strategic Management Journal*, 18(7), 509-533.
- Uhl-Bien, M., & Arena, M. (2018). Leadership for organizational adaptability: A theoretical synthesis and integrative framework. *The Leadership Quarterly*, 29(1), 89-104.
- Uhl-Bien, M., Marion, R., & McKelvey, B. (2007). Complexity leadership theory: Shifting leadership from the industrial age to the knowledge era. *The Leadership Quarterly*, 18(4), 298-318.
<https://doi.org/10.1016/j.leaqua.2007.04.002>
- Van Quaquebeke, N., & Gerpott, F. H. (2023). The now, new, and next of digital leadership: How artificial intelligence (AI) will take over and change leadership as we know it. *Journal of Leadership & Organizational Studies*, 30(3), 265-275. <https://doi.org/10.1177/15480518231181731>
- Voegtlin, C., Patzer, M., & Scherer, A. G. (2012). Responsible leadership in global business: A new approach to leadership and its multi-level outcomes. *Journal of Business Ethics*, 105(1), 1-16.
- World Economic Forum. (2025). *The future of jobs report 2025*. World Economic Forum.