

Management, Leadership, and Governance in the Age of Artificial Intelligence: The Contributions of the FILE Odyssey

Author: Guillaume Mariani

ORCID: <https://orcid.org/0009-0003-1130-5061>

Affiliation: Independent management scholar, Paris, France

Date: July 2026 (first version); August 2026 (revised version)

Document Type: Working Paper

Subject Areas: Leadership; Leadership Studies; Leadership Theory; Leadership Development; Management; Management Theory; Organization Theory; Artificial Intelligence Governance; Human-Centered Artificial Intelligence; Management Education; Business Education

Abstract

Artificial intelligence (AI) is no longer only a technological instrument inside organizations. It is becoming a structural feature of management, leadership, strategy, decision-making, education, research, and governance. AI can automate tasks, classify situations, generate knowledge outputs, support decisions, organize work, optimize processes, and shape what leaders see before they decide. This working paper presents the FILE Odyssey as an intellectual journey and a long-term transdisciplinary research program on **management, leadership, and governance in the age of artificial intelligence and robotics**.

AI does not reduce the need for human leadership. It makes human judgment, responsibility, care, legitimacy, education, knowledge production, and governance more decisive. The central question is not only what AI can do, but what forms of human intelligence, judgment, responsibility, and institutional protection become more necessary when AI becomes more powerful.

The central question of the FILE Odyssey is: **What must human leadership become in the age of artificial intelligence and robotics?**

Leadership = Intelligence + Character + Judgment. Within the FILE Architecture, Intelligence is conceptualized through FILE, Character through SQL, and Judgment through LENS.

The paper synthesizes twelve original contributions of the FILE Odyssey to leadership, management, and the social sciences: Leadership = Intelligence + Character + Judgment; FILE: The Five Intelligences of Leadership Evolution; FILE, FILE³, FILE⁵, FILE⁷, and the 7E Cascade; SQL: The Six Qualities of Leadership; LENS: Leadership, Ethics, Nondelegation, and Sensemaking; RC: The Relational Commons; MAIL: Management, AI, and Leadership and MLT: Management, Leadership, and Technology; HACK: Human-AI Co-created Knowledge; Human Irreducibility; The Five Immutable Laws of AI and Robotics: HTTPS: Humanity, Truth, Transparency, Protection, and Security; BASIC: Basic Income, AI Access, Skills,

Inclusion, and Care, and its French equivalent, RULES : Revenu Universel, Lutte contre la pauvreté et les inégalités, Éducation et accès à l'IA et Solidarité; and Inalienable Leadership. Each contribution follows the same sequence: AI problem, literature review, research gap, and FILE positioning and contribution. Together, they form a conceptual, theoretical, and research-generating program for leadership, management education, responsible knowledge production, human judgment, and AI governance.

The paper's central management-theory argument is that AI should not be treated as a technical tool added to existing leadership models, but as a structural condition that changes how leaders exercise judgment, preserve relational conditions, form managers, produce knowledge, and govern institutions.

For leaders, managers, executives, entrepreneurs, educators, scholars, institutions, and organizations, FILE provides **a management theory and diagnostic framework for responsible leadership and decision-making in the age of artificial intelligence and robotics** while preserving human judgment, dignity, responsibility, care, trust, and the capacity to govern AI systems and robots under human leadership.

Keywords: FILE; Five Intelligences of Leadership Evolution; FILE Odyssey; artificial intelligence; AI; augmented intelligence; leadership; leadership studies; leadership theory; augmented leadership; AI leadership; management; management theory; organization theory; decision-making; management education; business education; leadership education; executive education; governance; AI governance; human-centered artificial intelligence; human judgment; Human Irreducibility; Human Flourishing; Humans + Machines; Relational Commons; SQL; LENS; MAIL; MLT; HACK; HTTPS; BASIC; RULES; Inalienable Leadership.

1. Introduction: Why the FILE Odyssey Matters

Artificial intelligence has moved from a managerial tool to a condition of management itself. The stronger AI becomes, the more leadership must remain human, responsible, and governed. This paper develops that claim across twelve linked contributions of the FILE Odyssey. Sections 1 and 2 define the problem and the twelve contributions. Sections 3 to 14 develop each contribution. Section 15 shows how they fit together. Section 16 develops the research agenda. The conclusion returns to the core claim of FILE: **Humans + Machines, under human leadership**.

FILE answers this problem through twelve linked contributions. They move from leadership intelligence to development, relational life, education, responsible knowledge production, judgment, Human Irreducibility, and AI governance.

The argument is simple: AI is no longer only a tool that managers use. It is becoming part of the conditions in which organizations work, learn, decide, and govern. Leadership therefore needs more than technical AI literacy.

1.1 The AI problem

AI is automating, accelerating, and mediating core functions of business, management, leadership, education, research, and organizational decision-making. It can generate text, classify situations, support decisions, organize work, optimize processes, and shape what leaders see before they decide. The central problem is no longer only what AI can do. The central problem is who governs the human consequences of what AI tools produce.

This shift is not limited to technical work. AI systems increasingly enter the core domains of managerial and organizational life: analysis, forecasting, coordination, evaluation, communication, recruitment, performance management, education, research synthesis, decision support, and institutional governance. Automation studies have already shown that technology transforms tasks rather than only replacing entire occupations (Autor, Levy, and Murnane 2003; Autor 2015; Acemoglu and Restrepo 2019). The current wave of generative AI extends that transformation into knowledge work, language work, professional judgment, and organizational decision-making (Brynjolfsson, Li, and Raymond 2023; Eloundou et al. 2023).

For management and leadership, the consequence is profound. AI technologies do not merely accelerate existing work. They mediate it. They influence what counts as relevant evidence, what appears as a problem, what categories are used, what options are generated, what risks become visible, and what forms of action appear legitimate. In this sense, AI does not simply assist decision-making. It can shape the conditions under which decisions become thinkable.

This creates a second-order problem for leadership. AI does not only support decisions. It can shape the frameworks, metrics, classifications, and decision architectures through which leaders perceive reality before they act. When leaders rely on AI tools, dashboards, predictive models, automated recommendations, or algorithmic classifications, they can believe they are exercising judgment over neutral information. Yet the information has already been structured by technical, organizational, political, and knowledge choices.

The risk is not only that AI can make errors. The deeper risk is that leaders can stop noticing how AI shapes the field of judgment before they decide. The problem is therefore not only the governance of AI outputs. It is the governance of the systems that shape what leaders can see, measure, compare, value, and decide.

1.2 Literature review

This paper anchors FILE in several bodies of scholarship: leadership studies, management theory, organization theory, sociotechnical systems, human-centered artificial intelligence, AI governance, decision-making, management education, epistemology, and moral philosophy.

Leadership theory already offers powerful accounts of influence, transformation, ethics, service, authenticity, adaptation, relational quality, distribution, and change. Transformational leadership explains how leaders mobilize people through vision and collective purpose (Burns 1978; Bass 1985). Servant leadership places care, development, stewardship, and humility at the center of leadership responsibility (Greenleaf 1977). Authentic leadership emphasizes self-awareness, values, transparency, and balanced processing (Avolio and Gardner 2005). Ethical leadership addresses normatively appropriate conduct, fairness, communication, and

accountability (Brown, Treviño, and Harrison 2005). Adaptive leadership explains how leaders mobilize people to confront difficult problems, losses, conflicts, and learning demands (Heifetz 1994; Heifetz, Grashow, and Linsky 2009). Complexity and distributed leadership approaches move beyond the individual leader to examine emergent, systemic, and collective forms of leadership (Gronn 2002; Uhl-Bien, Marion, and McKelvey 2007).

Management theory also provides essential foundations. It explains how organizations coordinate work, allocate resources, design structures, build strategies, define performance, shape cultures, and respond to uncertainty. Management frameworks help organizations act under complexity, but they can also become instruments of simplification, measurement, and control. In the age of AI, this ambiguity becomes sharper. Frameworks, dashboards, models, and algorithms do not merely describe organizations. They can reshape organizational attention and action.

Sociotechnical systems theory and related traditions are equally important because they reject the idea that technology acts alone. Technologies are embedded in human practices, organizational routines, institutional structures, professional cultures, incentives, and power relations. AI technologies therefore do not enter organizations as neutral tools. They are interpreted, adopted, resisted, domesticated, normalized, and governed by human systems. This is why leadership after artificial intelligence cannot be reduced to technical fluency. It requires judgment about the human, cultural, political, ethical, and institutional consequences of technological mediation.

Research on artificial intelligence and management has already identified crucial tensions. The automation-augmentation paradox shows that AI can both substitute for and augment human work (Raisch and Krakowski 2021). Human-artificial intelligence symbiosis in organizational decision-making shows that decision authority can become distributed between people and systems (Jarrahi 2018). Work on learning algorithms shows how organizational routines, expertise, and authority are reshaped when algorithmic systems participate in work (Faraj, Pachidi, and Sayegh 2018). Research on algorithmic management shows how control can be reorganized through data, metrics, platforms, and automated evaluation (Kellogg, Valentine, and Christin 2020). Studies of organizational decision-making in the age of AI show that decision structures themselves change when AI becomes part of organizational cognition (Shrestha, Ben-Menahem, and von Krogh 2019).

Management education literature also matters. Business schools have long been criticized for confusing analytical training, administrative technique, and managerial formation. Mintzberg argued that management education cannot be reduced to abstract analysis detached from experience, judgment, and practice (Mintzberg 2004). AI intensifies this critique. If many administrative, analytical, and coordinative functions can be automated or mediated by AI, then management education must ask a deeper question: what should human beings be formed to do when machines can perform many of the functions previously treated as the core of managerial competence?

Epistemology and moral philosophy complete the theoretical field of this paper. AI now participates in the production, synthesis, and circulation of knowledge. It also raises questions

about responsibility, dignity, meaning, care, trust, and moral agency. These questions cannot be answered by technical governance alone. They require reflection on what kind of being the human person is, what kinds of responsibility cannot be delegated, what forms of judgment cannot be reduced to prediction, and what limits must govern AI systems and robots at organizational and civilizational scale.

1.3 Research gap

Existing literature explains many parts of the AI transformation. Leadership theory explains influence, purpose, adaptation, ethics, relationships, and change. Management theory explains organization, strategy, coordination, culture, and institutions. Sociotechnical systems scholarship explains the co-constitution of technology, work, and organizational practice. AI governance literature explains risk, accountability, transparency, safety, security, and regulation. Management education literature critiques professional formation and business-school design. Epistemology explains knowledge, evidence, validity, and interpretation. Moral philosophy explains responsibility, dignity, care, and human meaning.

Yet these literatures often remain fragmented. The problem is not the absence of relevant scholarship, but its separation across leadership studies, management theory, sociotechnical systems, management education, epistemology, moral philosophy, and AI governance. Technology ethics, leadership psychology, management education, organizational theory, epistemology, and governance frequently address different parts of the same AI-driven transformation without connecting them in one research program.

This fragmentation matters because AI does not affect only one domain. It affects leadership, organization, education, judgment, knowledge production, relational life, governance, and human meaning at the same time. The transformation is simultaneous: AI reshapes work, education, knowledge, judgment, governance, and human meaning at once. A leader using AI does not face a purely technical question. The leader faces a leadership question, a management question, a sociotechnical question, a knowledge-related question, an educational question, a moral question, and a governance question simultaneously.

The broad knowledge gap is this: social science lacks a unified research program connecting leadership intelligence, development, relational life, education, methodology, judgment, philosophy, and governance under conditions shaped by AI. Many theories illuminate parts of the problem. Few bring them together around human responsibility in the age of AI. This fragmentation creates the need for a transdisciplinary program that connects leadership intelligence, organizational life, education, knowledge production, judgment, philosophy, and AI governance.

This is where FILE enters.

FILE enters this gap as a management theory and diagnostic framework for responsible leadership and decision-making in the age of artificial intelligence and robotics. It extends into a transdisciplinary research program on management, leadership, and governance. Its central question is simple: which human capacities must be developed and protected when AI becomes a powerful mediator of organizational life and everyday life?

Leadership = Intelligence + Character + Judgment. Within the FILE Architecture, Intelligence is conceptualized through FILE, Character through SQL, and Judgment through LENS.

1.4 FILE positioning and contribution

FILE is a management theory and diagnostic framework for responsible leadership and decision-making in the age of artificial intelligence and robotics. It identifies the forms of intelligence that become more necessary when machines become more powerful.

The foundational FILE Formula is expressed as:

$$\text{FILE} = \text{AI} + \text{EQ} + \text{CQ} + \text{PQ} + \text{AQ}$$

In this formula, AI means Augmented Intelligence (“Mind”), not artificial intelligence alone. EQ means Emotional Intelligence (“Heart”). CQ means Cultural Intelligence (“World”). PQ means Political Intelligence (“Society”). AQ means Adaptive Intelligence (“Evolution”). The formula is deliberately simple, but the problem it addresses is not. It states that responsible leadership in the age of artificial intelligence and robotics requires the integrated development of technological, emotional, cultural, political, and adaptive capacities under human responsibility.

FILE begins with leadership intelligence, then expands beyond the initial framework. It develops a wider research program connecting leadership, management education, relational life, responsible knowledge production, judgment, human irreducibility, and governance. The role of the paper is therefore not only descriptive. It argues that leadership in the age of AI requires a wider management vocabulary than digital transformation, AI leadership, leadership agility, or responsible AI alone can provide. FILE offers that vocabulary by integrating leadership intelligence, relational protection, management education, AI-assisted knowledge production, second-order judgment, human irreducibility, and AI governance into one transdisciplinary research program.

This paper presents the twelve original contributions of the FILE Odyssey: a human-centered definition of leadership, framework, developmental theory, ethos, judgment model, vision, education, methodology, philosophy, AI governance codex, public policy, and civilization. Leadership = Intelligence + Character + Judgment provides the human-centered definition. FILE names the five intelligences leaders need. FILE³, FILE⁵, FILE⁷, and the 7E Cascade explain how those intelligences mature through development, effectiveness, excellence, ecosystems, empowerment, execution, and embodiment. SQL: The Six Qualities of Leadership names Character through Vision, Intuition, Taste, Courage, Discipline, and Integrity. LENS: Leadership, Ethics, Nondelegation, and Sensemaking protects the conditions of human judgment when AI shapes what leaders see, measure, and decide. RC: The Relational Commons names the human world to protect: the shared workplace and world built on trust, dignity, psychological safety, and meaning. MAIL: Management, AI, and Leadership and MLT: Management, Leadership, and Technology name the educational response. HACK: Human-AI Co-created Knowledge names the method for responsible AI-assisted knowledge production. Human Irreducibility names the philosophical boundary. The Five Immutable Laws of AI and Robotics: HTTPS: Humanity, Truth, Transparency, Protection, and Security

provides the AI governance layer, five non-negotiable laws that must govern the design, deployment, and governance of AI systems and robots to protect human beings, human societies, and the future of human civilization. BASIC and RULES name the human-centered public-policy contribution. Inalienable Leadership names the human-centered civilization and the overarching metatheory of leadership in the age of artificial intelligence and robotics.

Its contribution is the management-centered integration of existing literature into one research program. The sharper question is this: what human intelligences, relational conditions, educational models, knowledge practices, judgment safeguards, philosophical boundaries, and AI governance laws become necessary when AI becomes a structural condition of both organizational life and everyday life?

The cumulative movement is clear: Leadership = Intelligence + Character + Judgment defines leadership; FILE names the intelligences; the 7E Cascade explains their development; SQL names Character; LENS protects judgment; RC names the human world they protect; MAIL and MLT translate the theory into education; HACK governs AI-assisted knowledge production; Human Irreducibility defines the philosophical boundary; HTTPS translates that boundary into human-centered AI governance; BASIC and RULES extend the program into public policy; and Inalienable Leadership provides the culminating civilizational and metatheoretical integration.

1.5 Purpose, scope, and method

This working paper is a conceptual, theoretical, and programmatic synthesis. It defines the FILE Odyssey, situates its twelve contributions within relevant literatures, clarifies the gaps those contributions address, and opens a research agenda for refinement, comparison, application, empirical research, organizational use, and teaching pilots.

Each contribution follows the same path: AI problem, literature review, research gap, FILE positioning and contribution. This sequence helps the reader compare the twelve contributions while keeping FILE as one integrated research program.

The paper advances two levels of argument. Each contribution addresses a specific problem created by AI, and the twelve contributions together form a transdisciplinary research program on management, leadership, and governance.

The research agenda is direct: clarify the constructs, compare them with established theories, operationalize them in education and organizations, examine their cross-cultural relevance, and examine their value for responsible leadership and decision-making.

2. Overview of the Twelve Original Contributions of the FILE Odyssey

Before developing the twelve human-centered contributions in detail, this section gives the reader the scope of the paper. Table 1 should not be read as a list of adjacent concepts. It presents the twelve contributions in their cumulative relation.

1. **A human-centered definition of leadership.** Leadership = Intelligence + Character + Judgment.
2. **A human-centered framework.** FILE: The Five Intelligences of Leadership Evolution.
3. **A human-centered developmental theory.** FILE, FILE³, FILE⁵, FILE⁷, and the 7E Cascade: Evolution → Effectiveness → Excellence → Ecosystems → Empowerment → Execution → Embodiment.
4. **A human-centered ethos.** SQL: The Six Qualities of Leadership.
5. **A human-centered judgment model.** LENS: Leadership, Ethics, Nondelegation, and Sensemaking.
6. **A human-centered vision.** RC: The Relational Commons.
7. **A human-centered education.** MAIL: Management, AI, and Leadership, a curriculum supporting proposed MLT: Management, Leadership, and Technology degrees.
8. **A human-centered methodology.** HACK: Human-AI Co-created Knowledge.
9. **A human-centered philosophy.** Human Irreducibility.
10. **A human-centered AI governance codex.** The Five Immutable Laws of AI and Robotics: HTTPS: Humanity, Truth, Transparency, Protection, and Security.
11. **A human-centered public policy.** BASIC: Basic Income, AI Access, Skills, Inclusion, and Care, and its French equivalent, RULES : Revenu Universel, Lutte contre la pauvreté et les inégalités, Éducation et accès à l'IA et Solidarité.
12. **A human-centered civilization.** Inalienable Leadership: the twelfth and culminating original contribution of the FILE Odyssey and the overarching metatheory of leadership in the age of artificial intelligence and robotics.

The sequence begins with a human-centered definition of leadership, then names the intelligences required for leadership in the age of AI and robotics. It then asks ten further questions: how those intelligences mature, what character must embody them, how judgment must be governed, what human world they must protect, how they should be taught, how AI-assisted knowledge should be governed, what human capacities machines must never replace, what laws should govern AI systems and robots, what public policy should accompany AI transformation, and what final human authority and responsibility cannot be surrendered.

The movement is cumulative. Leadership = Intelligence + Character + Judgment defines leadership; FILE provides the framework; the 7E Cascade develops it; SQL names Character; LENS governs Judgment; RC gives it a human destination; MAIL and MLT translate it into education; HACK extends it into responsible knowledge production; Human Irreducibility provides the philosophy; HTTPS provides the AI governance layer; BASIC and RULES provide the public-policy contribution; and Inalienable Leadership provides the culminating civilizational metatheory.

Several boundaries must remain clear. RC and Human Irreducibility are related, but not identical: RC names the shared relational and institutional ecosystem FILE seeks to protect, while Human Irreducibility names the philosophical boundary that explains why certain human capacities cannot be reduced to machine outputs. MAIL/MLT and HACK are also distinct: MAIL and MLT concern formation for management in the age of AI, while HACK

concerns responsible Human-AI Co-created Knowledge across research, education, R&D (research and development), innovation, strategy, organizations, and human life. LENS and HTTPS are distinct as well: LENS protects the conditions of human judgment, while HTTPS states the Five Immutable Laws of AI and Robotics.

None of these later contributions adds a sixth intelligence to FILE. Each develops, protects, educates, governs, or philosophically grounds the original five-intelligence foundation.

Table 1. The Twelve Original Contributions of the FILE Odyssey

No.	Contribution	FILE role	Core problem addressed	The FILE Odyssey response
1	Leadership = Intelligence + Character + Judgment	Human-centered definition of leadership	AI expands machine capability while leadership still requires a human integration of capability, character, and responsible judgment.	Leadership = Intelligence + Character + Judgment. Within the FILE Architecture, Intelligence is conceptualized through FILE, Character through SQL, and Judgment through LENS.
2	FILE: The Five Intelligences of Leadership Evolution	Human-centered framework	AI automates and reshapes management, decision-making, coordination, and communication, creating the need to identify the human intelligences that must govern leadership in the age of artificial intelligence.	FILE defines responsible leadership in the age of artificial intelligence through five integrated intelligences: Augmented Intelligence (AI, Mind), Emotional Intelligence (EQ, Heart), Cultural Intelligence (CQ, World), Political Intelligence (PQ, Society), and Adaptive Intelligence (AQ, Evolution).
3	FILE, FILE ³ , FILE ⁵ , FILE ⁷ , and the 7E Cascade	Human-centered developmental theory	AI accelerates complexity and instability, making static competence models insufficient for leadership development.	The 7E Cascade explains how leadership intelligence matures through Evolution, Effectiveness, Excellence, Ecosystems, Empowerment, Execution, and Embodiment.
4	SQL: The Six Qualities of Leadership	Human-centered ethos	AI can augment intelligence and execution without supplying the character of the human leader who remains responsible for action.	SQL defines Character through six qualities in this order: Vision, Intuition, Taste, Courage, Discipline, and Integrity.
5	LENS: Leadership, Ethics, Nondelegation, and Sensemaking	Human-centered judgment model	AI increasingly shapes the frameworks, metrics, classifications, and decision architectures through which leaders perceive and decide.	LENS protects the second-order governance of human judgment by insisting that leadership, ethics, nondelegation, and sensemaking remain human responsibilities when AI shapes decision-making conditions.

No.	Contribution	FILE role	Core problem addressed	The FILE Odyssey response
6	RC: The Relational Commons	Human-centered vision	AI increases efficiency while risking the erosion of trust, dignity, psychological safety, meaning, and human connection.	RC names the shared relational and institutional ecosystem FILE protects, cultivates, and renews: a workplace and world built on trust, dignity, psychological safety, and meaning.
7	MAIL: Management, AI, and Leadership / MLT: Management, Leadership, and Technology	Human-centered education	AI is automating or mediating many functions traditionally taught in business and management education.	MAIL argues that Management, AI, and Leadership must be taught and practiced as one integrated discipline. MLT translates this into a degree model: Management, Leadership, and Technology.
8	HACK: Human-AI Co-created Knowledge	Human-centered methodology	AI tools can now generate, summarize, translate, synthesize, critique, and draft knowledge outputs at scale, creating new questions of scholarly responsibility.	HACK establishes that AI-assisted knowledge production must be governed by human judgment, source discipline, interpretation, accountability, and final responsibility.
9	Human Irreducibility	Human-centered philosophy	AI creates pressure for organizations and societies to treat human capacities as replaceable, simulable, or functionally equivalent to machine outputs.	Human Irreducibility states that judgment, discernment, conscience, moral agency, dignity, meaning, responsibility, care, trust, and love cannot be replaced by a machine, no matter how powerful it becomes.
10	The Five Immutable Laws of AI and Robotics: HTTPS: Humanity, Truth, Transparency, Protection, and Security	Human-centered AI governance	AI technologies are deployed at organizational and civilizational scale, affecting truth, attention, security, autonomy, institutions, dignity, and the future of human societies.	HTTPS contributes a human-centered AI governance codex: The Five Immutable Laws of AI and Robotics: HTTPS: Humanity, Truth, Transparency, Protection, and Security, five non-negotiable laws that must govern the design, deployment, and governance of AI systems and robots to protect human beings, human societies, and the future of human civilization.

No.	Contribution	FILE role	Core problem addressed	The FILE Odyssey response
11	BASIC / RULES	Human-centered public policy	AI-driven and robotics-driven transformation reshapes work, access, skills, inclusion, inequality, and the social conditions of care.	BASIC: Basic Income, AI Access, Skills, Inclusion, and Care, and its French equivalent, RULES : Revenu Universel, Lutte contre la pauvreté et les inégalités, Éducation et accès à l'IA et Solidarité, provides the public-policy contribution.
12	Inalienable Leadership	Human-centered civilization	As AI and robotics become more capable, final authority and responsibility for the purposes, institutions, and direction of civilization must remain human.	Inalienable Leadership is the twelfth and culminating original contribution and the overarching metatheory of leadership in the age of artificial intelligence and robotics.

Table 1 shows the cumulative movement of the FILE Odyssey. It begins with the human-centered definition of leadership, then moves through the five intelligences required for leadership, development, character, judgment, relational life, education, methodology, philosophy, AI governance, public policy, and civilization. This progression matters because AI disruption is not one problem. It is a layered transformation. It affects what leadership is, what leaders must know, how leaders must develop, what character leaders must embody, how judgment must remain nondelegable, what organizations must protect, how education must change, how knowledge must be governed, what human capacities must remain irreducible, what laws must constrain AI systems and robots, what public policy must address, and what final human authority and responsibility must remain inalienable.

The paper now develops each human-centered contribution in turn. Each section follows the same sequence: AI problem, literature review, research gap, and FILE positioning and contribution. This common sequence helps readers compare the twelve contributions and understand how they form one integrated research program.

3. Contribution 1: A human-centered definition of leadership. Leadership = Intelligence + Character + Judgment

3.1 The AI problem

Artificial intelligence and robotics increasingly extend analysis, prediction, classification, recommendation, coordination, and execution. They enlarge what organizations can calculate and perform, but greater machine capability does not answer what leadership itself is. As AI systems and robots enter more organizational functions, the question becomes more precise: what must remain distinctively human in leadership when capability can increasingly be supplied or amplified by machines?

This is a definitional problem, not only a technological one. Leadership requires the capacities needed to understand and act, but capability alone cannot define leadership. A highly capable actor can pursue a destructive objective. Intelligence without Character can be directed toward unworthy ends. Intelligence without Judgment can produce sophisticated action without responsible discernment.

The stronger machine capability becomes, the more important this distinction becomes. Human leadership must integrate what the leader is capable of understanding and doing, the Character embodied by the person who acts, and the Judgment through which responsibility is exercised. Contribution 1 gives that integration a compact definition.

3.2 Literature review

Leadership scholarship has long treated effective action as dependent on differentiated human capabilities. Katz distinguished technical, human, and conceptual skills as different requirements of effective administration (Katz 1955). The skills approach was developed further in leadership research by Mumford and colleagues, whose model centers leadership performance on the capacity to solve complex social problems through knowledge, problem-solving skills, social judgment skills, and the construction of workable solutions (Mumford et al. 2000). Sophisticated capability-based accounts therefore long predate FILE. They establish that leadership performance involves more than undifferentiated intelligence or technical expertise.

Capability, however, does not exhaust what leadership scholarship asks of the person who leads. Bass and Steidlmeier located moral character, values, and ethical foundations within their distinction between authentic and pseudo-transformational leadership (Bass and Steidlmeier 1999). Crossan and colleagues place leader character even more directly within organizational and leadership research, treating character as a foundational component of leadership and connecting it to judgment and decision making (Crossan et al. 2017). Character therefore already has a substantial scholarly foundation.

A parallel intellectual tradition establishes the limits of calculation in organizational action. Simon's *Administrative Behavior* analyzes organizational decision-making under practical constraints and shows how decisions are shaped by organizational processes rather than arising from unlimited rational calculation (Simon 1947). Schön, writing about professional practice rather than leadership theory in the narrow sense, demonstrates the importance of professional knowing and reflection-in-action when uncertain situations exceed the reach of technical rationality (Schön 1983). These traditions establish that consequential human action cannot be understood as calculation alone.

Practical-wisdom scholarship deepens this challenge to purely technical or rule-bound action. Schwartz and Sharpe develop practical wisdom as context-sensitive judgment that brings knowledge, values, experience, and action together rather than relying mechanically on rules (Schwartz and Sharpe 2010). Shotter and Tsoukas provide a direct bridge to leadership by examining how people in leadership positions arrive at judgment in circumstances marked by ambiguity, surprise, and conflicting values, and by developing phronesis as situated practical

wisdom (Shotter and Tsoukas 2014). Leadership judgment therefore has explicit scholarly foundations.

Integrative leadership architectures also precede FILE. Sternberg's WICS model combines Wisdom, Intelligence, and Creativity within a model of leadership and argues that effective leadership depends on their synthesis (Sternberg 2003). WICS is therefore a serious prior comparator. Its existence rules out any claim that FILE is original merely because it integrates several human dimensions. FILE's contribution is different. It selects **Intelligence, Character, and Judgment** as the three dimensions of human leadership and organizes that architecture specifically for conditions created by artificial intelligence and robotics. The comparison makes the originality claim more precise: the contribution lies not in multidimensionality itself, but in the particular dimensions selected, their integration, and their function in defining human leadership when machine capability becomes increasingly powerful.

3.3 Research gap

Established scholarship therefore contains substantial accounts of leadership skills, character, organizational decision-making, practical wisdom, and multidimensional leadership. The gap is not the absence of these traditions. They are substantial, but differently organized. They do not provide FILE's specific compact three-dimensional architecture for defining human leadership under the conditions created by artificial intelligence and robotics.

Contribution 1 makes an architectural claim. Its constituent concepts have intellectual histories; their deliberate selection and integration into **Leadership = Intelligence + Character + Judgment** is the contribution. The AI-era function of that architecture is decisive. Machines can expand analytical and operational capability without supplying human Character, responsible Judgment, or human accountability for consequences. As machine capability grows, any definition of leadership that treats capability as sufficient becomes less adequate. FILE addresses this definitional and integrative gap by specifying the three dimensions that must remain joined in human leadership when intelligence-like capabilities are increasingly distributed across humans and machines.

3.4 FILE positioning and contribution

Leadership = Intelligence + Character + Judgment.

Within the FILE Architecture, Intelligence is conceptualized through FILE, Character through SQL, and Judgment through LENS.

Intelligence concerns the human capabilities necessary to understand, interpret, learn, engage, decide, and act. Character concerns the qualities embodied by the human leader who exercises those capabilities. Judgment concerns the responsible exercise of decision and action under ambiguity, uncertainty, competing values, and consequences, including conditions increasingly shaped by AI-mediated information, metrics, classifications, and recommendations.

The three dimensions are structurally and logically integrated. They are not stages. Intelligence, Character, and Judgment interact continuously. The plus signs do not mean first

Intelligence, then Character, then Judgment. Leadership requires capability to be joined to the qualities of the person who acts and to the responsible discernment through which action is chosen and owned.

This is Contribution 1 of the FILE Odyssey: a **human-centered definition of leadership**. It establishes the leadership architecture within which the subsequent contributions operate. FILE, SQL, and LENS then conceptually develop and give theoretical expression to Intelligence, Character, and Judgment within the wider FILE Architecture.

4. Contribution 2: A human-centered framework. FILE: The Five Intelligences of Leadership Evolution

4.1 The AI problem

AI increasingly mediates the ordinary work of management. It supports analysis, generates text, classifies situations, summarizes documents, compares options, forecasts risks, coordinates tasks, evaluates performance, and recommends actions. In many organizations, leaders no longer encounter reality only through direct observation, human conversation, professional experience, or managerial reports. They increasingly encounter it through data structures, dashboards, predictive models, automated summaries, and AI-assisted outputs.

This transformation changes the nature of leadership. It makes human judgment more decisive. It changes the conditions under which leadership must think, decide, communicate, and act. When AI processes more information, generates more options, and executes more functions, leaders face a new responsibility: they must govern not only people, organizations, and strategies, but also the AI technologies that mediate organizational attention and action.

The central problem is therefore not only technical adoption. The leadership question goes beyond tool knowledge, prompt literacy, or the supervision of digital transformation projects. These skills matter, but they address only a fraction of the leadership problem created by AI. A leader can use AI effectively and still fail to understand the emotional consequences of automation, the cultural biases embedded in data, the political effects of algorithmic control, or the adaptive demands created by uncertainty.

Leadership and management in the age of artificial intelligence and robotics require an integrated view of human intelligence. Augmented Intelligence (“Mind”) alone is not enough. Emotional Intelligence (“Heart”) is needed because AI affects fear, trust, motivation, identity, and belonging. Cultural Intelligence (“World”) is needed because AI operates across languages, norms, institutions, and knowledge traditions. Political Intelligence (“Society”) is needed because data, metrics, automation, and governance structures redistribute power. Adaptive Intelligence (“Evolution”) is needed because AI accelerates change, ambiguity, and learning demands.

The leadership problem, then, is one of integration. Organizations need a clear framework for the human intelligences that must govern work shaped by artificial intelligence. Without such a framework, leadership becomes fragmented. Some leaders focus only on technical fluency.

Others focus only on ethics, communication, or change management. Others respond through compliance or risk control. Each response is useful, but none is sufficient alone. The age of AI and robotics requires leaders who can integrate technological augmentation, emotional responsibility, cultural interpretation, political judgment, and adaptive learning.

FILE begins from this problem. It defines leadership in the age of artificial intelligence and robotics through the human intelligences that become more necessary, not less necessary, when machines (both AI systems and robots) become more powerful.

4.2 Literature review

FILE stands inside the leadership and management literature. It builds from several established traditions and reorganizes them around the conditions created by artificial intelligence described above.

Leadership theory has long examined how leaders influence people, interpret change, mobilize collective action, and hold responsibility under uncertainty. Transformational leadership emphasizes vision, motivation, and collective purpose (Burns 1978; Bass 1985). Servant leadership emphasizes stewardship, care, service, and the development of others (Greenleaf 1977). Authentic leadership emphasizes self-awareness, values, transparency, and moral consistency (Avolio and Gardner 2005). Ethical leadership emphasizes fairness, accountability, communication, and normatively appropriate conduct (Brown, Treviño, and Harrison 2005). Adaptive leadership emphasizes the capacity to mobilize people through difficult problems that cannot be solved by technical expertise alone (Heifetz 1994; Heifetz, Grashow, and Linsky 2009). Complexity leadership and distributed leadership approaches also show that leadership is not only an individual trait. It emerges within systems, interactions, networks, and changing contexts (Gronn 2002; Uhl-Bien, Marion, and McKelvey 2007).

These traditions remain essential. FILE integrates their central insight: leadership is not reducible to authority, charisma, technique, or position. Leadership requires judgment, relationships, context, responsibility, and the capacity to act under uncertainty. AI intensifies this point. When organizations become more technologically mediated, leadership must become more capable of integrating human and technical realities.

Management theory also provides an important foundation. Management frameworks help leaders diagnose organizations, coordinate work, allocate resources, design structures, make strategic choices, and translate ideas into action. Yet management frameworks can also narrow attention if they treat organizations only as systems of optimization. AI makes this risk more serious because it can turn managerial attention toward what is measurable, structured, and automatable. FILE responds by placing human intelligence at the center of management diagnosis.

The five intelligences of FILE also connect to established bodies of research. Augmented Intelligence draws from traditions of human-computer interaction, human-centered artificial intelligence, sociotechnical systems, and augmentation rather than replacement. The central idea is that technology should extend human capacity under human judgment, not displace human responsibility (Licklider 1960; Engelbart 1962; Shneiderman 2022).

Emotional Intelligence draws from research on emotion, perception, self-regulation, empathy, and social awareness (Salovey and Mayer 1990; Goleman 1995; Mayer, Salovey, and Caruso 2004). In organizations, emotions are not secondary to rationality. They shape trust, motivation, conflict, resilience, and decision-making. AI can generate efficient outputs, but it does not carry human care, accountability, or emotional presence.

Cultural Intelligence draws from research on cross-cultural interaction, intercultural competence, and the capacity to act effectively across different cultural contexts (Earley and Ang 2003; Ang et al. 2007). This is increasingly important because AI systems operate across cultures while being trained, designed, deployed, and interpreted within specific cultural, linguistic, institutional, and economic contexts. Leaders must therefore ask not only whether a system works, but for whom, according to which assumptions, and within which cultural frame.

Political Intelligence connects to work on power, influence, organizational politics, and political skill (Pfeffer 1992; Ferris et al. 2005; Pfeffer 2010). AI does not remove politics from organizations. It can hide it. Metrics, classifications, dashboards, and automated controls can appear neutral while redistributing authority, visibility, status, and voice. Leaders therefore need political intelligence to understand how technological systems reshape power.

Adaptive Intelligence connects to adaptive leadership, organizational learning, dynamic capabilities, complexity, and decision-making under uncertainty (Argyris and Schön 1978; Senge 1990; Teece, Pisano, and Shuen 1997; Heifetz 1994; Teece 2007). AI tools change rapidly. Their organizational consequences often emerge after adoption, not before. Leaders therefore need the capacity to learn, revise assumptions, interpret feedback, and adjust action under conditions that remain unstable.

FILE brings these literatures together through the five intelligences. It identifies the integrated human capacities that become necessary when AI mediates management, decision-making, coordination, and governance.

4.3 Research gap

Existing leadership theories explain many dimensions of leadership with depth and authority. They explain vision, motivation, ethics, service, authenticity, adaptation, relationships, distribution, and complexity. Emotional intelligence explains how leaders perceive and regulate emotions. Cultural intelligence explains how leaders act across cultural contexts. Research on power and political skill explains influence, legitimacy, and organizational politics. Adaptive leadership and dynamic capabilities explain learning under uncertainty. Sociotechnical systems theory explains the interaction between people, organizations, and technologies.

The gap is not a weakness. The gap is a separation. Organizations transformed by AI require leaders to bring these capacities together at the same time. A leader dealing with algorithmic performance management, for example, must understand the technology, the emotions it creates, the cultural assumptions it carries, the power it redistributes, and the adaptive

learning it requires. The age of artificial intelligence and robotics requires one diagnostic framework that brings these dimensions together.

Digital leadership literature contributes important insights, but it often remains too focused on transformation, innovation, digital strategy, or technological fluency. Ethics literature contributes important safeguards, but it often remains too focused on principles and risk. Leadership development literature contributes important formation models, but it often does not fully integrate artificial intelligence as a structural condition of decision-making. Human-centered artificial intelligence literature contributes important design principles, but it often focuses more on system design than on leadership formation. FILE enters this gap with a specific leadership question: what human intelligences must be integrated when AI becomes a powerful mediator of organizational life?

This local gap establishes the need for FILE as a management theory and diagnostic framework for responsible leadership and decision-making in the age of artificial intelligence and robotics. FILE integrates technological augmentation, emotional responsibility, cultural interpretation, political judgment, and adaptive learning under human leadership.

4.4 FILE positioning and contribution

Here, FILE organizes the human intelligences that leaders must integrate when AI becomes part of organizational life and everyday life.

The foundational FILE Formula is expressed as:

$$\text{FILE} = \text{AI} + \text{EQ} + \text{CQ} + \text{PQ} + \text{AQ}$$

In this formula, AI means Augmented Intelligence (“Mind”). It does not mean artificial intelligence alone. EQ means Emotional Intelligence (“Heart”). CQ means Cultural Intelligence (“World”). PQ means Political Intelligence (“Society”). AQ means Adaptive Intelligence (“Evolution”).

The five primary intelligences also contain a disciplined nested architecture. Augmented Intelligence includes Cognitive Intelligence and Complexity Intelligence; Political Intelligence includes Purpose Intelligence, Moral Intelligence, and Sustainable Intelligence; Adaptive Intelligence includes Decision Intelligence, Judgment Intelligence, and Learning Intelligence. These nested intelligences deepen AI, PQ, and AQ without creating additional primary intelligences.

FILE can also be understood through the Hand Metaphor: Augmented Intelligence (Mind) is the thumb; Emotional Intelligence (Heart) is the index finger; Cultural Intelligence (World) is the middle finger; Political Intelligence (Society) is the ring finger; and Adaptive Intelligence (Evolution) is the little finger. Each finger matters because each intelligence has a distinct function. Leadership happens through the coordination of the whole hand. The hand is a human hand, reminding us that leadership, however augmented by technology, remains a human responsibility. Leadership after AI requires a human hand capable of holding machine intelligence without being held by it.

The five intelligences correct one another. Augmented Intelligence asks whether a system is useful; Emotional Intelligence asks what it does to dignity, fear, care, and trust; Cultural Intelligence asks whose meanings are recognized or erased; Political Intelligence asks who benefits, who decides, and who can contest; Adaptive Intelligence asks whether people can learn, endure, revise, and exercise the Right to Refuse automation when adoption becomes harmful or irresponsible. This matters because **Political Intelligence without Purpose degenerates into Machiavellianism**, adaptation without limits becomes Adaptive Exhaustion, and resilience demanded without protection becomes Resilience Abuse.

The formula is simple by design and diagnostic in function. It helps leaders, educators, researchers, and institutions see when leadership development is too narrow for the conditions created by artificial intelligence. A technically fluent leader without emotional intelligence creates fear or disengagement. An emotionally intelligent leader without cultural intelligence risks misreading diverse contexts. A culturally sensitive leader without political intelligence fails to understand power, legitimacy, and institutional conflict. A politically aware leader without adaptive intelligence becomes trapped in old routines. A highly adaptive leader without Augmented Intelligence fails to govern AI systems effectively.

A management program can teach analytics while neglecting judgment, care, power, and adaptation. An executive team can invest in AI while ignoring emotional and cultural consequences. A public institution can automate processes while underestimating legitimacy and trust.

FILE therefore positions leadership as an integrated human responsibility. The five intelligences are not separate skills placed side by side. They are interdependent capacities connecting technological judgment, emotional responsibility, cultural interpretation, political legitimacy, and adaptive learning.

Augmented Intelligence requires Emotional Intelligence because AI adoption affects people. Emotional Intelligence requires Cultural Intelligence because emotions are interpreted within cultures. Cultural Intelligence requires Political Intelligence because cultures are shaped by institutions, histories, interests, and power. Political Intelligence requires Adaptive Intelligence because power relations shift under technological change. Adaptive Intelligence requires Augmented Intelligence because leaders must learn how to work with AI without transferring responsibility to it.

This positioning clarifies the role of artificial intelligence. FILE refuses to present machines as leaders. It treats AI systems as powerful tools, infrastructures, and mediators that must remain under human judgment. The leadership question is not how machines can lead. The leadership question is how human beings lead responsibly when AI shapes work, knowledge, attention, and decision-making.

FILE contributes The Five Intelligences of Leadership Evolution as a human-centered framework for responsible leadership in the age of AI and robots. Its contribution is integrative because it brings technological, emotional, cultural, political, and adaptive dimensions of leadership into one framework. It is diagnostic because it helps leaders and institutions see

when their leadership model is incomplete. It is educational because it asks which human intelligences should be formed when AI can perform many technical and administrative functions. It is protective because it keeps leadership as a human responsibility. It is programmatic because it opens a research agenda: clarify the five intelligences, compare them with adjacent constructs, examine their teachability, operationalize them across sectors and cultures, and examine their use in leadership education and organizational practice.

The distinction can be summarized through two positioning sentences from the FILE Corpus.

From FILE vs. Major Management Frameworks:

“Management frameworks help organizations coordinate action; FILE asks what human intelligences must govern that coordination.”

From FILE vs. Major Leadership Theories (V2):

“Major leadership theories explain dimensions of leadership; FILE asks what integrated human intelligences must govern leadership judgment under AI-mediated complexity.”

FILE adds to leadership theory by making AI a structural condition of leadership rather than a technical context around it. Its originality lies in integrating Augmented Intelligence, Emotional Intelligence, Cultural Intelligence, Political Intelligence, and Adaptive Intelligence as a diagnostic language for responsible leadership and decision-making.

Once these intelligences are named, the next question follows naturally. If leadership requires Augmented Intelligence, Emotional Intelligence, Cultural Intelligence, Political Intelligence, and Adaptive Intelligence, how do these intelligences develop over time? This question leads to FILE³, FILE⁵, FILE⁷, and the 7E Cascade: Evolution, Effectiveness, Excellence, Ecosystems, Empowerment, Execution, and Embodiment.

5. Contribution 3: A human-centered developmental theory. FILE, FILE³, FILE⁵, FILE⁷, and the 7E Cascade

5.1 The AI problem

AI accelerates change. It shortens cycles of analysis, production, communication, experimentation, and decision support. It also increases the speed at which organizations must adapt to new tools, new risks, new forms of work, new expectations, and new governance problems. Leaders therefore face a developmental challenge. It is not enough to possess a static set of competencies. Leaders must learn how their intelligences mature across time, scale, uncertainty, execution, and embodiment.

This problem is visible in many organizations. A leader can understand AI technologies at a technical level and still fail to translate that understanding into effective practice. A leadership team can adopt AI and fail to build the cultural and emotional conditions for responsible use. An organization can design an AI strategy and fail to empower people, coordinate workflows, or embed new habits. A school can teach AI literacy and fail to transform leadership

formation. A company can run pilots and experiments and fail to move from innovation discourse to institutional execution.

The complexity created by AI therefore produces a developmental gap. Leaders need to know what intelligences matter, but they also need to understand how those intelligences evolve. Leadership cannot remain a list of qualities. It must become a developmental process. It must move from awareness to effectiveness, from effectiveness to excellence, from excellence to ecosystems, from ecosystems to empowerment, from empowerment to execution, and from execution to embodiment.

FILE³, FILE⁵, FILE⁷, and the 7E Cascade address this developmental problem. FILE names the five intelligences. The 7E Cascade explains how leadership intelligence matures.

5.2 Literature review

The developmental dimension of FILE is situated within several established literatures: leadership development, adult development, organizational learning, dynamic capabilities, adaptive leadership, complexity leadership, professional formation, and ecosystem theory.

Leadership development research has long distinguished leadership as a set of individual traits from leadership as a developmental process. Day argued that leadership development should be understood not only as the acquisition of individual skills, but also as the expansion of social, relational, and organizational capacity (Day 2000). This distinction is important for FILE. The five intelligences cannot remain personal attributes. They must develop into practices, routines, relationships, and institutional capacities.

Adult development and professional formation also matter. Development is not only the accumulation of knowledge. It involves changes in how people make meaning, assume responsibility, relate to others, and understand themselves (Kegan 1982; Mezirow 1991). In professional life, formation means more than competence. It includes identity, judgment, ethical responsibility, and the internalization of standards. FILE⁷ uses this insight when it treats Embodiment not as an extra skill, but as the internalization of leadership intelligence into character, habit, and presence.

Organizational learning literature offers another foundation. Argyris and Schön distinguished between single-loop learning, which corrects errors within existing assumptions, and double-loop learning, which questions the assumptions themselves (Argyris and Schön 1978). Senge described learning organizations as systems capable of shared vision, mental-model reflection, team learning, and systems thinking (Senge 1990). AI intensifies the need for these forms of learning because it can reinforce existing assumptions at scale unless leaders question the categories, metrics, and routines that guide use.

Dynamic capabilities theory is also relevant. Organizations must sense opportunities and threats, seize possibilities, and transform resources and routines under changing conditions (Teece, Pisano, and Shuen 1997; Teece 2007). AI increases the importance of such capabilities because it changes markets, work processes, knowledge production, and competitive conditions. Yet dynamic capabilities are not only technical capabilities. They require leadership judgment, organizational learning, coordination, and adaptation.

Adaptive leadership provides a further bridge. Heifetz distinguishes technical problems from adaptive challenges. Technical problems can be solved through existing expertise. Adaptive challenges require learning, conflict, loss, experimentation, and changes in values or behavior (Heifetz 1994; Heifetz, Grashow, and Linsky 2009). Transformation driven by AI is often adaptive in this sense. It cannot be solved only by buying tools or hiring specialists. It requires people and institutions to change how they think, work, decide, and take responsibility.

Complexity leadership also supports the developmental logic of FILE. Leadership in complex systems cannot be reduced to linear control. It involves emergence, interaction, enabling conditions, and adaptation across networks (Uhl-Bien, Marion, and McKelvey 2007). FILE⁵ builds on this insight by moving from individual leadership capacity toward ecosystems and empowerment. In environments transformed by AI, leadership must operate across individuals, teams, organizations, platforms, institutions, and broader stakeholder systems.

Ecosystem theory reinforces this movement. Organizations increasingly depend on networks of partners, technologies, infrastructures, regulators, communities, and knowledge systems (Moore 1996; Iansiti and Levien 2004). AI makes this ecosystemic condition more visible. Data, models, vendors, platforms, standards, laws, users, employees, and institutions all shape what AI technologies do. Leadership must therefore move beyond individual competence and even beyond the single organization.

FILE³, FILE⁵, FILE⁷, and the 7E Cascade bring these developmental traditions into FILE. They establish that the five intelligences become meaningful when they mature across a full developmental trajectory.

5.3 Research gap

Many literatures explain leadership development, organizational learning, adaptation, capabilities, and ecosystems. Yet these literatures often address different levels of the problem. Leadership development often focuses on individuals and teams. Organizational learning focuses on routines, assumptions, and systems. Dynamic capabilities focus on strategic renewal. Adaptive leadership focuses on mobilizing people through difficult change. Complexity leadership focuses on emergence and enabling conditions. Ecosystem theory focuses on interdependence across organizations and environments.

These traditions are valuable. The age of AI creates a need to connect them more directly. Leadership in the age of artificial intelligence and robotics does not develop only inside the individual. It develops across tools, workflows, teams, organizations, ecosystems, and institutions. It also develops inwardly, through the leader's identity and responsibility. A leader must learn to understand AI, act effectively with it, build excellence over time, operate across ecosystems, empower human agency, execute responsibly, and embody the values that prevent technological capability from becoming irresponsible power.

The local gap is therefore developmental. FILE names the five intelligences. The five intelligences also require a theory of maturation. Without such a theory, FILE risks being read as a static list or a competency inventory rather than a theory of leadership evolution. The 7E Cascade resolves this risk by showing how the five intelligences move from recognition to

practice, from practice to excellence, from excellence to ecosystems, from ecosystems to empowerment, from empowerment to execution, and from execution to embodied leadership.

This gap is especially important for management education and organizational application. Educators cannot teach FILE only as a formula. Organizations cannot use FILE only as a checklist. Leaders cannot embody FILE only by naming the five intelligences. Development requires sequencing, practice, reflection, feedback, institutionalization, and internalization.

5.4 FILE positioning and contribution

FILE positions leadership as developmental. The five intelligences are not fixed possessions. They must mature through experience, reflection, education, practice, and responsibility. FILE³, FILE⁵, FILE⁷, and the 7E Cascade extend FILE from a framework into a human-centered theory of leadership evolution.

The 7E Cascade is:

Evolution → Effectiveness → Excellence → Ecosystems → Empowerment → Execution → Embodiment

FILE³ marks the first developmental extension. It connects the framework to Evolution, Effectiveness, and Excellence. FILE⁵ extends this movement toward Ecosystems and Empowerment. FILE⁷ adds Execution and Embodiment. Together, these stages explain how the five intelligences mature from awareness into practice, ecosystemic leadership, responsible action, and embodied leadership.

Each stage defines a different leadership requirement in the age of AI and robotics.

Evolution names the need for leadership to change when AI changes the conditions of work, knowledge, coordination, and decision-making. Effectiveness names the practical capacity to act responsibly in those conditions by combining Augmented Intelligence, Emotional Intelligence, Cultural Intelligence, Political Intelligence, and Adaptive Intelligence in concrete managerial situations. Excellence names the discipline through which this integration becomes consistent, rigorous, and sustainable over time.

Ecosystems name the requirement that leadership move beyond the individual leader and the single organization. Leadership now operates across interdependent systems of people, technologies, institutions, cultures, stakeholders, platforms, vendors, regulators, users, data infrastructures, labor markets, educational systems, and publics. Empowerment names the normative direction of this ecosystemic leadership. The goal is not only organizational efficiency. It is the expansion of human agency, dignity, autonomy, participation, and responsible capacity as AI becomes part of organizational life and everyday life.

Execution translates the theory into organized practice. It concerns the routines, workflows, decision protocols, training processes, governance practices, and organizational habits through which the five intelligences become operational. Embodiment turns the theory into lived leadership. It concerns character, presence, responsibility, and the internalization of judgment. A leader can know FILE conceptually and still fail to embody it. Embodiment means that the

five intelligences become part of how the leader perceives situations, relates to people, uses AI, and accepts responsibility.

These stages are developmental stages, not additional intelligences. They do not create a sixth FILE intelligence. FILE remains grounded in the five-intelligence formula. The 7E Cascade explains how that formula develops across time, scale, and practice.

The contribution of the 7E Cascade is therefore developmental and diagnostic. It enables researchers, educators, and organizations to examine where leadership development is incomplete. Some organizations are strong in Evolution because they understand that leadership must change, but weak in Execution because they cannot translate theory into routines. Others are strong in AI adoption, but weak in Empowerment because people experience the change as surveillance or loss of agency. Others teach the language of leadership intelligence, but fail to cultivate Embodiment.

The 7E Cascade prevents FILE from remaining a static framework. It shows that the five intelligences must become effective, excellent, ecosystemic, empowering, executable, and embodied. In this sense, Contribution 3 deepens Contribution 2. FILE names the intelligences required for leadership. FILE³, FILE⁵, FILE⁷, and the 7E Cascade explain how those intelligences mature into leadership evolution.

After FILE names the intelligences and the 7E Cascade explains their development, the paper turns to SQL: The Six Qualities of Leadership, the Character dimension of the human leader.

6. Contribution 4: A human-centered ethos. SQL: The Six Qualities of Leadership

6.1 The AI problem

AI can increase analytical capacity, automate execution, and accelerate organizational action. It can compare alternatives, detect patterns, generate scenarios, predict outcomes, and carry out instructions at a speed no human leader can match. None of these capabilities supplies the character of the human leader who decides what should be pursued, which trade-offs are acceptable, when action is warranted, and what responsibility requires.

Leadership therefore confronts a category mistake if machine performance is treated as a substitute for leader character. A system may help identify a direction without possessing Vision, surface weak signals without possessing Intuition, rank alternatives without exercising Taste, automate action without Courage, sustain a process without Discipline, and reproduce rules without Integrity. The point is not that AI cannot generate language or behavior associated with these qualities. It is that simulation is not embodiment, and output is not responsibility.

The stronger AI becomes, the more leadership needs a clear character dimension that remains human and embodied. SQL addresses that dimension. It asks what qualities must be formed in

the person who exercises intelligence, acts under uncertainty, and remains answerable for the consequences of leadership.

6.2 Literature review

Character already has a substantial place in leadership scholarship. Bass and Steidlmeier (1999) placed moral character and ethical values at the heart of their distinction between authentic and pseudo-transformational leadership. Crossan et al. (2017) developed a broader framework of leader character in organizations and connected character to leadership effectiveness, judgment, and decision-making. Their work is especially important because it treats character not as an ornamental virtue but as a substantive part of how leaders act. Neither account reduces leadership to capability alone.

Vision has an equally established genealogy. Burns (1978) and Bass (1985) made purposive direction and transformation central to influential accounts of leadership. Their theories extend far beyond Vision as SQL uses the term, but they establish a durable point: leadership involves orienting people toward a future that does not yet exist. Direction is not a technical afterthought. It is part of what distinguishes leadership from the competent administration of the present.

Other leadership demands arise before formal analysis is complete. Dane and Pratt (2007) conceptualize intuition in managerial decision-making as rapid, nonconscious, and holistic judgment shaped by experience and domain knowledge. Organizational aesthetics opens a different but related territory. Strati (1992) established aesthetic awareness as a legitimate way of understanding organizational life, while Hansen, Ropo, and Sauer (2007) brought aesthetic experience directly into leadership through sensory perception, felt meaning, and tacit knowledge. These literatures provide a serious basis for recognizing forms of sensing and qualitative discrimination that cannot always be reduced to explicit criteria. Taste, in this sense, is not personal preference or social distinction. It concerns the capacity to discriminate among qualities, forms, meanings, proportions, and possibilities when metrics alone cannot determine what is worth choosing.

Seeing what should be done is still not doing it. Detert and Bruno's (2017) synthesis of workplace courage shows that courage is a complex organizational phenomenon involving action in the presence of perceived risk. Leadership often reaches precisely this threshold. A direction may be clear and an option may appear worth choosing, yet action can still require commitment under uncertainty, opposition, fear, or potential loss.

Action must also endure. Manz's (1986) theory of self-leadership gives an important adjacent foundation through self-influence, self-direction, and individual self-control. Bandura (1991) provides a stronger account of self-regulation through self-monitoring, standards, judgment, self-reaction, and personal agency. This scholarship helps distinguish sustained self-governance from obedience or externally imposed control. Initiating action and sustaining it across time are different leadership demands.

Integrity closes a different problem: coherence. Palanski and Yammarino (2007) show how conceptual confusion around integrity can be reduced by focusing on consistency between

words and actions. That conception is narrower than an all-purpose theory of ethical leadership, which is precisely its value here. Leadership requires more than choosing and acting. What leaders say, value, decide, and repeatedly do must remain sufficiently coherent for responsibility to be credible.

The scholarly picture is therefore already rich. Vision, intuition, aesthetic understanding, courage, self-regulation, and integrity all have serious intellectual antecedents, while character scholarship provides an overarching foundation. What these streams do not provide is one protected six-quality sequence organized around a common Character function. That is the point at which the literature review ends and SQL's theoretical contribution begins.

6.3 Research gap

The research gap is architectural and integrative, not an absence of scholarship on character or virtue. Existing research provides sophisticated but dispersed accounts of the qualities that SQL brings together. What is absent from the literature reviewed here is this particular configuration: Vision, Intuition, Taste, Courage, Discipline, and Integrity, in that order, as the Character dimension within **Leadership = Intelligence + Character + Judgment**.

AI makes this integration more consequential. Machine capability can increasingly support analysis, prediction, choice architecture, and execution without embodying the human qualities through which a leader establishes direction, senses before complete proof, discriminates among possibilities, commits to action, sustains effort, and remains coherent over time. The problem is therefore not whether machines can perform more leadership-related tasks. It is what Character must mean when more of the surrounding cognitive and operational work can be technologically augmented.

SQL is selective rather than exhaustive: it does not claim that these six qualities are the only virtues or qualities relevant to good leadership.

6.4 FILE positioning and contribution

SQL stands for The Six Qualities of Leadership. The six qualities are Vision, Intuition, Taste, Courage, Discipline, and Integrity, in that order.

SQL is the human-centered ethos of the FILE Odyssey and the Character dimension within **Leadership = Intelligence + Character + Judgment**. FILE develops the intelligences and capabilities required for leadership. SQL describes the character of the embodied human leader who exercises them. LENS addresses Judgment, including its ethical and nondelegable dimensions. Character and Ethics therefore remain distinct.

The six qualities perform different functions. **Vision provides Direction:** it concerns where leadership seeks to go. **Intuition provides Sensing:** it enables leaders to register patterns, tensions, and possibilities before complete proof is available. **Taste provides Discrimination:** it concerns recognizing what is genuinely worth choosing when formal criteria are insufficient.

Courage provides Commitment: it moves leadership from perception toward action despite uncertainty or risk. **Discipline provides Sustained Execution:** it keeps action coherent across

time through self-governance rather than external compulsion. **Integrity provides Coherence:** it aligns identity, principle, words, and conduct.

The functional arc is:

Direction → Sensing → Discrimination → Commitment → Sustained Execution → Coherence

This is an architecture, not a rigid chronology. Leadership can move backward and forward across these functions as circumstances change. New information may reopen sensing, failed execution may require renewed discrimination, and integrity may force a leader to reconsider a direction already chosen.

SQL's contribution is therefore not the invention of Vision, Intuition, Taste, Courage, Discipline, or Integrity individually. It lies in their deliberate selection, protected order, integration, functional interpretation, and common role as the Character dimension of the FILE Architecture.

SQL therefore connects leadership development to embodiment: intelligence must be joined by character before it can be exercised consistently under responsibility.

7. Contribution 5: A human-centered judgment model. LENS: Leadership, Ethics, Nondelegation, and Sensemaking

7.1 The AI problem

AI increasingly shapes the conditions under which leaders perceive, evaluate, and decide. It organizes information, classifies people, ranks risks, recommends actions, predicts outcomes, flags anomalies, summarizes evidence, and structures attention. In many organizations, leaders do not encounter reality directly. They encounter reality through dashboards, metrics, models, classifications, alerts, reports, rankings, and AI-assisted recommendations.

This creates a second-order problem. First-order leadership asks how leaders decide. Second-order governance asks who or what governs the conditions under which leaders perceive, evaluate, and decide. The issue is not only whether a leader accepts or rejects an AI output. The deeper issue is whether human judgment still governs the systems, metrics, classifications, outputs, and decision architectures that shape what the leader sees before the decision is even made.

AI can make this problem invisible because it often enters leadership through support functions. It appears as decision support, performance monitoring, workflow optimization, risk analysis, recruitment screening, customer segmentation, compliance assistance, or strategic forecasting. These uses may be valuable. Yet they can also shape judgment in advance. They can decide what counts as relevant, what becomes visible, what remains hidden, who is ranked as risky, which options appear rational, and which questions are never asked.

The danger is algorithmic capture: the progressive shaping of human perception, evidence, and choice by predictive systems before leaders recognize that judgment has already been shaped. Algorithmic capture does not require machines to make final decisions. It can occur when leaders retain formal authority but lose control over the conditions of attention. They still decide, but they decide inside frames, rankings, categories, and recommendations that they did not govern.

This is why the age of AI and robotics does not merely require better decision-making. It requires the governance of judgment itself. Leaders must ask not only whether a recommendation is useful, but how the recommendation was produced. They must ask not only whether a metric is accurate, but what it excludes. They must ask not only whether a classification is efficient, but whether it is legitimate. They must ask not only whether an output is explainable, but whether the decision condition remains humanly accountable.

LENS enters this problem as FILE's judgment model. It names the deeper governance problem already present in FILE before the acronym was introduced. It asks who governs the conditions of judgment when AI increasingly shapes what leaders see, measure, compare, value, and decide.

7.2 Literature review

LENS is grounded in decision theory, managerial judgment, sensemaking, ethics, accountability, nondelegation, AI governance, algorithmic management, explainability, audit society, and organizational responsibility.

Decision theory provides the first foundation. Simon showed that decision-making is bounded by limited information, limited attention, and limited cognitive capacity (Simon 1947). March emphasized ambiguity, rules, learning, and the organizational conditions under which decisions are made (March 1994). Tversky and Kahneman showed that judgment is shaped by heuristics and biases (Tversky and Kahneman 1974), while prospect theory demonstrated systematic departures from expected-utility reasoning under risk (Kahneman and Tversky 1979). These traditions matter because AI does not remove bounded rationality. It reorganizes it. It changes what information is selected, how alternatives are framed, and which predictions appear authoritative.

Managerial judgment adds a second foundation. Management has never been only calculation. It requires interpretation, prudence, experience, timing, practical wisdom, and responsibility. Mintzberg's work on managerial practice shows that management is embedded in action, interruption, communication, and context rather than abstract analysis alone (Mintzberg 1973). Schön's account of reflective practice shows that professionals think in action when situations are uncertain, complex, and unstable (Schön 1983). LENS extends these insights into management in the age of artificial intelligence and robotics. When machines support analysis, leaders must become more responsible for judgment, not less responsible.

Sensemaking is central. Weick showed that organizations act through processes of interpretation, enactment, selection, and meaning-making (Weick 1995). Leaders do not only receive facts. They make sense of ambiguous environments. AI can produce summaries,

predictions, and classifications, but it does not possess human sense. It does not carry the lived responsibility of deciding what a situation means for people, institutions, dignity, legitimacy, and action.

Ethics and accountability deepen the argument. Jonas's ethics of responsibility places responsibility at the center of action under technological power (Jonas 1984). Bovens describes accountability as a relationship in which actors must explain and justify conduct to others (Bovens 2007). Leadership in the age of AI intensifies both concerns. When systems shape recommendations, rankings, and evidence, leaders must remain answerable not only for final decisions but also for the conditions that made those decisions appear reasonable.

Nondelegation provides another foundation. In law and political theory, nondelegation marks the boundary of responsibility that cannot be transferred without damaging legitimacy. LENS translates this logic into leadership judgment. AI may assist perception, analysis, and comparison, but it cannot own moral agency. It cannot carry responsibility. It cannot answer for consequences. Human beings may delegate tasks, but they cannot delegate the burden of judgment.

AI governance and algorithmic management also belong in the field. Work on algorithmic management shows that data, metrics, and automated control can reshape authority, evaluation, surveillance, and worker autonomy (Kellogg, Valentine, and Christin 2020). Work on opaque algorithmic power shows that technical systems can classify people, organize opportunities, and shape institutional outcomes while hiding the political and moral choices embedded in their design (Pasquale 2015; O'Neil 2016). LENS takes these concerns into leadership theory by asking how human judgment governs the conditions that AI creates.

Explainability and the audit society complete the theoretical field. Explainability research asks how machine outputs can be made understandable to human users (Doshi-Velez and Kim 2017). Power's account of the audit society shows how systems of measurement and verification can reshape organizational behavior (Power 1997). These literatures are necessary but not sufficient. Explanation alone does not guarantee judgment. Audits alone do not guarantee responsibility. LENS asks whether human beings retain the capacity and authority to govern what explanations, metrics, classifications, and audits mean.

7.3 Research gap

Existing AI governance approaches often focus on fairness, transparency, privacy, explainability, safety, security, accountability, and compliance. These concerns are necessary. They ask whether AI is lawful, fair, robust, explainable, and safe. They also ask whether systems can be audited, monitored, documented, and constrained.

Yet these approaches often leave a deeper leadership problem underdeveloped. They do not always ask who governs the conditions under which human judgment itself is formed. They may improve the technical reliability of an AI system while leaving leaders dependent on its categories. They may require transparency while leaving metrics unchallenged. They may demand explainability while leaving the frame of decision untouched. They may audit outputs while leaving the underlying structure of attention intact.

The local gap is the second-order governance of human judgment. The central question is not only: is the AI output fair, transparent, or accurate? The central question is also: who or what governs the systems, metrics, classifications, outputs, and decision architectures that increasingly shape what leaders see, measure, compare, value, and decide?

This gap matters because judgment can be weakened without being formally removed. A leader may remain the official decision-maker while becoming dependent on frames created by AI. A committee may retain authority while accepting the categories produced by a dashboard. A manager may claim accountability while treating predictive scores as reality. An institution may preserve human approval while allowing automated rankings to define what counts as merit, risk, performance, or need.

The deeper risk is capture of knowledge and judgment: AI can shape the field of judgment before leaders recognize that their perception has already been shaped. This capture can happen quietly. It occurs when leaders no longer ask why a metric matters, how a classification was built, what a dashboard excludes, who can challenge a recommendation, or what moral responsibility remains when prediction becomes persuasive.

LENS fills this gap by naming and modeling the second-order governance of human judgment when AI shapes the conditions of perception and decision. It does not replace AI governance. It does not duplicate ethics or compliance. It asks a leadership question that must stand at the center of responsible management: how does human judgment remain sovereign over the systems that increasingly shape judgment itself?

7.4 FILE positioning and contribution

LENS is FILE's human-centered judgment model.

LENS stands for Leadership, Ethics, Nondelegation, and Sensemaking.

It proposes that human judgment, moral agency, and responsibility cannot be delegated to machines, and that the capacity for sense, meaning, and discernment must remain irreducibly human in every organization that uses AI.

This makes LENS distinct from HACK, HTTPS, and the five-intelligence formula. HACK governs responsible knowledge production with AI. LENS governs how human judgment interprets knowledge, evaluates evidence, and makes decisions when systems, metrics, classifications, outputs, and decision architectures shape what leaders see, measure, and decide. HTTPS states the Five Immutable Laws of AI and Robotics. LENS is not another intelligence. It protects the nondelegable human responsibility required by all five intelligences.

LENS formalizes a second-order problem already present in FILE: leaders must govern not only decisions, but also the conditions that shape decision-making. Leadership can fail before a decision is made, when data are selected badly, metrics hide meaning, categories erase persons, or predictive scores become moral shortcuts.

Leadership means that human beings remain responsible for direction, meaning, accountability, and action. AI may support analysis, coordination, prediction, and decision

conditions, but it does not carry leadership responsibility. Leaders cannot hide behind dashboards, models, vendors, procedures, or machine outputs. They remain responsible for what is done in the name of the organization.

Ethics means that judgment must be governed by responsibility, dignity, care, fairness, legitimacy, and attention to consequences. Ethical judgment cannot be reduced to compliance, risk control, or technical accuracy. An output can be accurate and still be unjust. A process can be efficient and still be degrading. A recommendation can be explainable and still be irresponsible. LENS insists that decisions shaped by AI must remain answerable to human moral agency.

Nondelegation means that the burden of judgment cannot be transferred to artificial intelligence. AI can assist perception, analysis, comparison, prediction, and drafting. It can support leaders and expand their field of attention. But it cannot become the moral subject of leadership. It cannot bear guilt, responsibility, courage, conscience, or accountability. The burden of judgment cannot be outsourced because **accountability cannot be outsourced**.

Sensemaking means that leaders must preserve the capacity to interpret meaning, context, ambiguity, and consequence. AI can process information, but leaders must make sense of what matters. They must ask what is being measured, what is being missed, whose voice is absent, which assumptions are embedded, which consequences are likely, and what the situation means for human beings. Sensemaking is the human capacity to hold meaning where data are incomplete, contested, or morally charged.

LENS asks leaders to examine the metrics that structure attention, the classifications that define persons, the dashboards that select reality, the outputs that frame alternatives, and the decision architectures that make some choices appear natural. It does not ask leaders to reject AI. It asks them to govern the conditions under which AI shapes judgment.

LENS protects Human Irreducibility before the paper explicitly turns to that philosophical boundary. It states that human judgment, moral agency, responsibility, sense, meaning, and discernment cannot be delegated to machines. These capacities are not ornamental. They are the conditions under which leadership remains human.

LENS adds to decision ethics and AI governance by shifting attention from decisions to the conditions that shape judgment before decisions appear.

After defining the judgment model that governs decision-making conditions shaped by AI, the paper turns to the human ecosystem that leadership must protect, cultivate, and renew: The Relational Commons.

8. Contribution 6: A human-centered vision. RC: The Relational Commons

8.1 The AI problem

AI can make organizations faster, more measurable, and more efficient. It can automate reporting, rank performance, summarize conversations, monitor activity, allocate tasks,

predict risks, and recommend decisions. It can also make organizational life more fragile and less human when it is deployed without responsible leadership.

This is the central problem addressed by RC: The Relational Commons. AI can improve coordination while weakening trust. It can increase visibility while damaging dignity. It can support decision-making while reducing voice. It can optimize workflows while eroding psychological safety. It can measure output while hiding care, meaning, invisible work, fatigue, conflict, and resentment.

The danger is not only technological error. The deeper danger is relational depletion. People can become data points, productivity units, risk profiles, compliance objects, or sources of extractable information. Managers can begin to see dashboards before they see people. Organizations can confuse measurable activity with meaningful work. Leaders can celebrate efficiency while the relational conditions that make responsible leadership possible are being damaged.

AI makes this problem more urgent because it scales managerial attention. What was once a local managerial practice can become an automated norm. What was once a subjective evaluation can become a data-driven classification. What was once a conversation can become a metric. What was once a human judgment can be displaced by an automated recommendation that appears objective because it is technical.

The Relational Commons names what is at stake. It names the shared workplace and world built on trust, dignity, psychological safety, and meaning. It also names the human ecosystem that FILE seeks to protect, cultivate, and renew through the integrated exercise of all five intelligences. When this commons is strong, people can speak, contest, learn, cooperate, disagree, repair trust, and act with responsibility. When it erodes, leadership becomes formal authority without human legitimacy.

RC therefore gives FILE its vision and leadership its destination. The question is not only whether AI makes organizations more efficient. The question is whether its use protects or depletes the human conditions that make responsible leadership possible.

8.2 Literature review

The Relational Commons is grounded in several established literatures: commons theory, trust, psychological safety, dignity, care ethics, social capital, legitimacy, meaningful work, sociotechnical theory, and responsible leadership.

Commons theory is essential because it shows that some resources are shared, fragile, and dependent on stewardship. Ostrom's work demonstrated that commons are not automatically destroyed by collective use. They can be governed through norms, rules, participation, monitoring, accountability, and shared responsibility (Ostrom 1990). RC extends this insight into the relational field of leadership. Trust, dignity, psychological safety, meaning, legitimacy, and mutual recognition are not owned by one person. They are shared social and moral resources. They can be built, damaged, depleted, repaired, protected, and cultivated.

Trust research also provides a foundation. Trust shapes cooperation, vulnerability, commitment, and risk-taking in organizations (Mayer, Davis, and Schoorman 1995; Rousseau et al. 1998). In work transformed by AI, trust becomes both more important and more difficult. Employees must trust leaders, but they must also understand how AI technologies are used, what data are collected, how decisions are supported, and where human judgment remains accountable.

Psychological safety is equally central. Edmondson defined psychological safety as a shared belief that a team is safe for interpersonal risk-taking (Edmondson 1999). It enables learning, voice, error reporting, experimentation, and dissent. AI tools can support learning, but they can also weaken psychological safety when workers feel monitored, ranked, replaceable, or unable to challenge automated outputs.

Dignity and recognition deepen the argument. Human beings do not only need efficient processes. They need to be treated as persons whose voice, worth, experience, and moral standing matter. Recognition theory shows that social life depends on forms of respect and mutual recognition (Honneth 1995). Work on dignity emphasizes that organizations can honor or violate the personhood of workers through everyday practices of evaluation, control, communication, and participation. In organizations transformed by AI, the protection of dignity becomes a leadership responsibility because automated classifications can reduce people to categories and scores.

Care ethics also belongs in the theoretical field of RC. Noddings and Tronto showed that care is not only a private sentiment. It is an ethical and political practice involving attention, responsibility, competence, responsiveness, and relational obligation (Noddings 1984; Tronto 1993). RC carries this insight into leadership. A workplace without care can still function, but it cannot sustain the human conditions of responsible leadership.

Social capital and legitimacy add further depth. Social capital research shows that relationships, norms, and trust support cooperation and collective action (Coleman 1988; Putnam 2000). Legitimacy research shows that institutions require more than formal authority. They must be perceived as appropriate, credible, and justified (Suchman 1995). AI technologies can disturb both. They can weaken social capital when they isolate workers or intensify surveillance. They can weaken legitimacy when decisions appear opaque, imposed, or impossible to contest.

Finally, sociotechnical theory explains why RC cannot be separated from technology. AI enters organizations through routines, incentives, workflows, authority relations, cultures, and practices. It is never only a technical instrument. It reshapes the relational field in which leadership happens. The Relational Commons therefore brings the human consequences of sociotechnical change into the center of leadership theory.

8.3 Research gap

The literatures on trust, psychological safety, dignity, care, social capital, legitimacy, and sociotechnical work are rich and necessary. They explain many of the conditions that allow people to cooperate, learn, speak, and belong in organizations. Yet these literatures often

remain separated. Trust is studied as trust. Psychological safety is studied as team climate. Dignity is studied as moral status. Care is studied as ethics. Social capital is studied as networks and norms. Legitimacy is studied as institutional acceptance. Sociotechnical theory studies the relation between technology and work.

Organizations transformed by AI need these literatures to speak together. A workplace shaped by AI does not face one relational problem at a time. Surveillance affects trust and dignity. Automated evaluation affects psychological safety and legitimacy. Productivity optimization affects care and meaning. Opaque metrics affect voice and ability to challenge. Algorithmic classification affects recognition and belonging.

The local gap is therefore conceptual and practical. Leadership needs a concept that names the shared human condition beneath responsible leadership in the age of artificial intelligence and robotics. It needs a concept that connects trust, dignity, psychological safety, meaning, mutual recognition, voice, care, ability to challenge, responsibility, legitimacy, and belonging without reducing them to culture, morale, or employee engagement.

The Relational Commons fills this gap. It gives FILE a human-centered vision that is stronger than a value statement and more precise than a general appeal to culture. RC names the commons that responsible leadership must protect. It also names the destination toward which the five intelligences are directed.

This point is decisive for FILE. The five intelligences are not mobilized only for performance, competitiveness, or adaptation. They are mobilized to protect, cultivate, and renew the Relational Commons. When Augmented Intelligence, Emotional Intelligence, Cultural Intelligence, Political Intelligence, and Adaptive Intelligence are correctly mobilized together, RC emerges as the relational condition that FILE seeks to protect, cultivate, and renew: the human ecosystem in which responsible leadership remains possible.

8.4 FILE positioning and contribution

RC gives FILE its vision.

RC is FILE's human-centered vision. It stands for the Relational Commons, the shared workplace and world built on trust, dignity, psychological safety, meaning, respect, emotional recognition, mutual recognition, human connection, belonging, care, reciprocity, collaboration, voice, ability to challenge and contest, autonomy, responsibility, legitimacy, harmony, wellbeing, and creative possibility. It is the North Star and final destination of FILE: the human ecosystem that the five intelligences seek to protect, cultivate, and renew.

RC also requires safeguards. FILE's Sovereign Ecosystem framework means that organizations must retain the human capacity to deliberate, contest, refuse, repair, and govern the systems that shape work and judgment. This requires anti-instrumentalization safeguards, contestability principles, and non-data sanctuaries: spaces where metrics, dashboards, and algorithmic outputs are deliberately suspended to protect human judgment, psychological safety, moral reasoning, and voice. It also includes a Right to Opacity: people and communities should not be made fully data-legible whenever legibility becomes a condition of control, reduction, or domination.

This fixes RC's place in FILE. The Relational Commons is not another intelligence, a governance model, or a methodology. It names the human world that responsible leadership must protect, cultivate, and renew.

RC emerges when the five intelligences are used together: Augmented Intelligence keeps artificial intelligence under human judgment; Emotional Intelligence protects trust and psychological safety; Cultural Intelligence protects meaning across contexts; Political Intelligence protects voice and legitimacy; Adaptive Intelligence protects learning and repair. The aim is therefore not merely to produce smarter leaders. It is to form leaders who can protect the human ecosystem that makes responsibility, judgment, cooperation, and meaning possible.

RC contributes a North Star: a shared workplace and world built on trust, dignity, psychological safety, and meaning. Without it, AI adoption can become a narrow pursuit of speed, scale, automation, and control. RC also contributes a condition. It names the shared human ground beneath responsible leadership: trust, dignity, legitimacy, mutual recognition, voice, care, the ability to challenge, responsibility, meaning, and belonging. When these conditions are absent, leadership loses its human ground.

RC also contributes a commons. It treats trust, dignity, psychological safety, meaning, and human connection as shared social and moral resources. They can be built, depleted, enclosed by power, extracted by surveillance, damaged by domination, weakened by neglect, and repaired through voice, care, responsibility, and accountable leadership.

RC contributes a stewardship task. Leaders protect the Relational Commons through spaces for human deliberation, norms of voice, mechanisms for challenging decisions, protection against retaliation, accountability when trust is damaged, and responsible rules for AI use. The Relational Commons cannot be protected by values statements alone. It must live through daily practices that keep people visible as persons.

RC contributes a diagnostic test. It is an early warning signal of leadership failure. When employees stop contesting automated evaluations because they believe objections will not be heard, the Relational Commons is already weakening. Trust, voice, dignity, meaning, and human connection often erode before performance indicators reveal the damage.

RC is therefore decisive for management and leadership in the age of artificial intelligence. It asks whether AI remains subordinate to the human ecosystem that responsible leadership must protect, rather than making organizations more efficient while making them less human.

RC adds to work on psychological safety, dignity, trust, and care by naming the shared relational world that leadership must protect when AI enters work, knowledge, and decision-making.

RC: The Relational Commons names the human world FILE seeks to protect. MAIL: Management, AI, and Leadership and MLT: Management, Leadership, and Technology then ask how leaders and students must be formed to protect it.

9. Contribution 7: A human-centered education. MAIL: Management, AI, and Leadership / MLT: Management, Leadership, and Technology

9.1 The AI problem

AI is transforming the foundations of management education. It can perform or support many functions that business schools have traditionally treated as central to managerial competence: analysis, reporting, presentation, coordination, market research, scenario generation, decision support, operational planning, and knowledge synthesis.

AI is already automating or taking over many of the core functions taught in MBA programs: accounting, auditing, finance, management control, business analytics, information systems, marketing, sales, customer relationship management (CRM), human resource management (HRM), recruitment, procurement, logistics, supply chain management, operations, project management, and strategic management. When machines can run the functions, the most important thing left for humans is to lead.

This transformation demands a deeper rethinking of management education. It makes traditional business administration education insufficient. If AI can automate or mediate many administrative, analytical, and coordinative functions, then business schools must ask a sharper question: what should human beings be formed to do when machines can perform many tasks once used to define managerial competence?

The answer cannot be limited to AI literacy. Leaders and students need more than tool familiarity, prompt technique, or technical awareness. They need judgment, responsibility, ethics, culture, power analysis, adaptation, care, and the capacity to govern AI systems and robots. They must learn when to use AI technologies, when to contest them, when to override them, when to slow down, when to deliberate, and when to protect the human conditions that AI can weaken.

This is an educational problem because leadership after AI is a formation problem. Managers cannot be prepared only through functional silos. They cannot be trained only in accounting, finance, marketing, HRM, operations, business analytics, data, or strategy. These domains remain important, but the age of artificial intelligence and robotics requires a different center of gravity. Management education must form leaders capable of integrating management, AI, and leadership under human responsibility.

MAIL and MLT enter this problem directly. MAIL names the curriculum logic: Management, AI, and Leadership must be taught together. MLT names the proposed degree model: Management, Leadership, and Technology. Together, they translate FILE into education.

9.2 Literature review

MAIL/MLT is grounded in management education, leadership education, adult learning, reflective practice, responsible management education, and human-centered artificial intelligence.

Management education has long faced criticism for separating analysis from practice, technique from judgment, and business training from responsibility. Mintzberg argued that

management cannot be taught as abstract analysis detached from experience and practice (Mintzberg 2004). Bennis and O'Toole criticized business schools for becoming too detached from the real practice of management (Bennis and O'Toole 2005). Pfeffer and Fong questioned whether business-school education delivers the managerial and professional outcomes it claims (Pfeffer and Fong 2002). Ghoshal warned that bad management theories can shape bad management practices when they normalize narrow assumptions about human behavior and organizational purpose (Ghoshal 2005).

These critiques become stronger in the age of artificial intelligence. If AI can perform many analytical and administrative functions, then management education must move beyond the reproduction of tasks that machines can now support. Its purpose must be formation. It must form judgment, responsibility, leadership presence, ethical reasoning, cultural interpretation, political understanding, and adaptive capacity.

Leadership education also supports this move. Leadership is not learned only through concepts. It requires practice, reflection, feedback, identity formation, ethical awareness, and the development of relational capacity. Experiential learning emphasizes the cycle of experience, reflection, conceptualization, and experimentation (Kolb 1984). Reflective practice emphasizes the capacity of professionals to think in action and learn from complex situations (Schön 1983). These traditions support MAIL/MLT because management in the age of AI requires leaders who can reflect, judge, and act under conditions that remain uncertain.

Responsible management education adds another foundation. Management education cannot be neutral about the kind of leaders it forms. It shapes how students understand organizations, workers, markets, technology, power, responsibility, and the public good. The age of artificial intelligence and robotics increases this responsibility. Business schools must prepare students not only to use AI tools, but to govern their effects on people, institutions, and society.

Human-centered artificial intelligence completes the theoretical field. Human-centered artificial intelligence emphasizes that technology should support human agency, dignity, oversight, and well-being rather than displace them. For management education, this means that AI cannot be taught as a separate technical topic placed beside leadership. It must be integrated into the formation of responsible managers and leaders.

MAIL/MLT brings these literatures into one educational response. It treats management, AI, and leadership as inseparable in the education of leaders for the age of artificial intelligence.

9.3 Research gap

Existing management education still often reflects older divisions. Strategy, finance, marketing, operations, accounting, organizational behavior, ethics, entrepreneurship, technology, and leadership are commonly taught as separate domains. This division can help organize knowledge, but it also fragments the formation of leaders. Students learn functions, tools, and cases, but they do not always learn how to govern organizations transformed by AI as human, cultural, political, ethical, and adaptive realities.

AI education can also remain too narrow. It often focuses on tools, data, automation, analytics, prompt use, or technical literacy. These are necessary, but they do not answer the leadership

question. Knowing how to use AI tools is not the same as knowing how to govern their effects on trust, dignity, power, culture, voice, work, and meaning.

Leadership education can remain equally incomplete when it ignores AI. A leadership curriculum that teaches vision, communication, ethics, or change management without integrating work transformed by artificial intelligence and robotics prepares students for a world that no longer exists. AI now shapes knowledge, coordination, decision support, evaluation, and organizational attention. Leadership education must therefore treat AI as a structural condition of management.

The local gap is educational integration. Management education needs a clear way to teach management, AI, and leadership together. It needs a curriculum logic that treats AI not as an elective add-on, but as a condition of responsible management. It also needs a degree model that signals a new educational center of gravity.

MAIL/MLT fills this gap. MAIL establishes the curriculum logic: Management, AI, and Leadership must be taught together. MLT establishes the proposed degree model: Management, Leadership, and Technology. The goal is not to create technologists with a thin layer of leadership language, or leaders with a thin layer of AI literacy. The goal is to form responsible managers and leaders capable of governing AI systems and robots while protecting the Relational Commons.

9.4 FILE positioning and contribution

MAIL/MLT translates FILE into education.

MAIL/MLT also concerns identity formation: management education must form provisional leadership identities, responsible confidence, and the capacity to inhabit judgment before certainty is available, while remaining open to francophone and cross-cultural management-education scholarship rather than reproducing a narrow Anglo-American model.

MAIL means Management, AI, and Leadership. It states that these three domains must be taught and practiced together. Management provides the organizational and strategic field. AI names the artificial intelligence reality that now reshapes work, knowledge, and decision-making. Leadership provides the human responsibility that must govern AI and organizational action.

MLT means Management, Leadership, and Technology. It names the proposed degree model that gives institutional form to the MAIL logic. MLT signals that management education can no longer treat technology as a support function, leadership as a soft supplement, and management as a set of technical or administrative tasks. These domains now belong together.

MAIL/MLT is not a replacement for FILE. It is the educational contribution of this research program. It asks how leaders and students must be formed to protect the human ecosystem named by RC. Education is central because the five intelligences do not develop by declaration. They require formation: responsible AI use, care and trust, cross-cultural interpretation, power and legitimacy analysis, and adaptive learning under uncertainty.

MAIL/MLT contributes a human-centered education for the age of artificial intelligence and robotics. Its contribution is curricular because it establishes that Management, AI, and Leadership must be taught together. Management without AI is incomplete because AI now shapes organizational work. AI without leadership is dangerous because technical capability can outrun judgment, care, legitimacy, and responsibility. Leadership without management is abstract because responsibility must be exercised through organizations, institutions, resources, decisions, and action.

Its contribution is institutional because MLT gives this curriculum logic a degree model: Management, Leadership, and Technology. The name matters because it changes the center of gravity. It places leadership between management and technology. It refuses both technocratic reduction and leadership romanticism. It states that the future of management education must form people who can act responsibly where organizations, technology, and human beings meet.

Its contribution is developmental because MAIL/MLT forms the five intelligences. It teaches Augmented Intelligence through responsible AI use under human judgment. It teaches Emotional Intelligence through fear, trust, motivation, care, and dignity in work transformed by AI. It teaches Cultural Intelligence through interpretation across languages, institutions, and cultural assumptions. It teaches Political Intelligence through power, legitimacy, purpose, governance, and the ability to challenge. It teaches Adaptive Intelligence through learning, revision, experimentation, and responsible action under uncertainty.

Its contribution is protective because MAIL/MLT forms leaders capable of protecting, cultivating, and renewing the Relational Commons. Education is not only about employability, productivity, or innovation. It is about forming people who can govern AI systems and robots while preserving trust, dignity, psychological safety, meaning, voice, and human connection. RC names the human world FILE seeks to protect. MAIL/MLT teaches leaders how to protect it.

Its contribution is programmatic because MAIL/MLT opens a research agenda for management education. Educators can examine how students develop the five intelligences, how AI reshapes learning, how leadership judgment can be taught under conditions created by AI, and how curricula can protect the Relational Commons. Institutions can test course designs, learning outcomes, pedagogical methods, assessment practices, and cross-cultural applications.

MAIL/MLT therefore gives FILE its educational instrument. It moves FILE from theory into formation. It makes clear that responsible leadership and decision-making in the age of artificial intelligence and robotics must be taught, practiced, embodied, and institutionalized.

After defining the relational world FILE protects and the education required to sustain it, the paper turns to HACK: Human-AI Co-created Knowledge, the methodology of responsible knowledge production with AI.

10. Contribution 8: A human-centered methodology. HACK: Human-AI Co-created Knowledge

10.1 The AI problem

AI can now generate, summarize, translate, synthesize, compare, critique, and draft knowledge outputs at scale. It can produce literature summaries, outline arguments, identify patterns, propose categories, reformulate concepts, simulate objections, translate across languages, and accelerate editorial work. This creates a central methodological and ethical challenge for responsible knowledge production: what counts as responsible knowledge when AI contributes to its production?

HACK is not limited to academic research or scholarship. It applies wherever human beings use AI to produce, organize, test, interpret, transmit, or apply knowledge: education, management, organizational learning, research and development, strategy, public administration, creative work, and everyday decision-making.

The problem is not only that AI may produce errors. Error is serious, but it is not the deepest issue. The deeper issue is that AI can make knowledge production feel complete before human inquiry has matured. It can fill silence before a question has been clarified. It can impose structure before the researcher has struggled with the material. It can generate plausible synthesis before the human has tested sources, tensions, concepts, and counterarguments. It can make the first draft of thought feel like thought itself.

AI also changes the speed and scale of scholarship. A researcher can now generate summaries, comparisons, outlines, bibliographic leads, conceptual maps, counterarguments, translations, and synthetic drafts in minutes. This power can support intellectual work, but it can also weaken the habits of attention that knowledge requires. Scholars can become dependent on fluent outputs. Students can mistake generated coherence for understanding. Professionals can accept plausible claims without source discipline. Organizations can treat AI-assisted synthesis as evidence without asking who verified it, who interpreted it, and who remains accountable for its use.

The danger is therefore knowledge-related and formative. AI can strengthen human inquiry when it is used under disciplined judgment. It can weaken inquiry when it replaces the human acts of questioning, reading, comparing, doubting, verifying, interpreting, and taking responsibility. The decisive question is not whether AI can assist research. It can. The decisive question is whether human judgment governs the whole process from question formation to source selection, interpretation, synthesis, revision, and final accountability.

This is the problem addressed by HACK: Human-AI Co-created Knowledge. HACK begins from a simple claim: AI-assisted knowledge production needs a methodology, not only tools. Without such a methodology, AI becomes either a shortcut, an oracle, or an invisible intellectual infrastructure. With such a methodology, AI outputs become material for questioning, verification, and interpretation within a process governed by human judgment, human authorship, and human accountability.

10.2 Literature review

HACK is grounded in epistemology, knowledge creation, research methodology, human-computer interaction (HCI), augmented intelligence, collective intelligence, responsible knowledge production, tacit and situated knowledge, social epistemology, and research integrity.

Epistemology provides the first foundation because HACK concerns the conditions under which knowledge claims become responsible. Knowledge is not only the possession of information. It requires justification, interpretation, evidence, and judgment. AI can generate outputs that appear knowledgeable, but the appearance of knowledge is not the same as warranted knowledge. A fluent summary is not proof. A coherent interpretation is not verification. A plausible citation is not a source. HACK therefore places AI inside a knowledge process governed by human inquiry.

Knowledge creation theory also matters. Nonaka's work on organizational knowledge creation distinguishes tacit and explicit knowledge and shows that knowledge emerges through processes of articulation, conversion, social interaction, and meaning-making (Nonaka 1994). Polanyi's account of tacit knowledge shows that human beings know more than they can fully state (Polanyi 1966). These traditions are essential for HACK because AI can manipulate explicit language, but it does not possess the embodied, historical, experiential, moral, and contextual understanding through which human beings decide what matters.

Research methodology adds another foundation. Responsible scholarship depends on questions, methods, sources, evidence, interpretation, limits, and transparency. AI can assist many of these activities, but it cannot replace the methodological responsibility of the researcher. The human scholar must define the object of inquiry, identify what counts as evidence, distinguish source from synthesis, test claims, disclose uncertainty, and take responsibility for the final argument. HACK brings these methodological obligations into the age of AI-assisted research.

Human-computer interaction (HCI) and augmented intelligence provide a fourth foundation. The tradition of augmentation argues that computing should extend human capability rather than replace human agency (Licklider 1960; Engelbart 1962). Human-centered artificial intelligence develops the same concern by asking how technological systems can remain reliable, safe, trustworthy, accountable, and supportive of human agency (Shneiderman 2022; Stanford Institute for Human-Centered Artificial Intelligence n.d.). HACK applies this logic to knowledge work. AI should augment inquiry, not displace judgment. It should expand the field of thought, not close it prematurely.

HACK also belongs to the lineage of collective intelligence and distributed cognition associated with Thomas Malone, Pierre Lévy, Edwin Hutchins, James Surowiecki, and MIT's Center for Collective Intelligence (Hutchins 1995; Lévy 1997; Surowiecki 2004; Malone and Bernstein 2015; Malone 2018; MIT Center for Collective Intelligence n.d.); it adds a stricter governance condition for generative AI: distributed knowledge production remains

responsible only when human judgment governs the question, sources, interpretation, verification, and final claim.

Situated knowledge deepens the argument. Haraway's account of situated knowledges emphasizes that knowledge is produced from locations, perspectives, histories, and responsibilities rather than from nowhere (Haraway 1988). This matters because AI outputs can appear placeless, neutral, and universal. HACK rejects that illusion. It asks the human author to locate the question, clarify the standpoint, understand the context, and give meaning to the output. AI may help produce language, but only human beings can situate knowledge within responsibility.

Social epistemology also belongs in the field. Knowledge is not produced by isolated minds alone. It depends on testimony, trust, institutions, interpretive communities, peer review, credibility, and responsibility for knowledge (Goldman 1999; Fricker 2007). AI complicates this field because it can produce outputs that resemble testimony without being a responsible knower. HACK therefore asks who is accountable when machine-generated outputs enter scholarly, managerial, or educational discourse. Its answer is direct: accountability remains human.

Research integrity completes the theoretical field. Authorship, citation, verification, transparency, and responsibility are not administrative details. They protect the credibility of knowledge. In responsible knowledge production with AI, these disciplines become more important, not less important. HACK treats verification, source discipline, and final responsibility as part of human judgment itself.

10.3 Research gap

Existing debates about artificial intelligence and knowledge production often move between three reductive positions. The first treats AI as a neutral tool. In this view, the tool has no deeper methodological significance. It simply helps users work faster. The second treats AI as a threat to scholarship. In this view, the tool contaminates intellectual work because it weakens authorship, originality, and responsibility. The third treats AI as a replacement for human labor. In this view, the machine becomes the producer of knowledge, and human judgment becomes secondary.

Each position sees something real, but each remains incomplete. AI is not neutral because it shapes how questions are framed, what answers appear plausible, what sources are surfaced, how complexity is simplified, and how language is organized. AI is not only a threat because it can support critique, comparison, accessibility, translation, synthesis, and intellectual exploration. AI is not a replacement for human knowledge because it cannot own responsibility, inhabit meaning, verify stakes, or answer morally for consequences.

The local gap is therefore methodological. Management and social science research need a language for responsible AI-assisted knowledge production. They need a way to distinguish assistance from authorship, fluency from knowledge, synthesis from verification, and output from responsibility. They also need a method for governing the process through which AI contributes to research, education, leadership reflection, and organizational analysis.

The same gap appears outside scholarship. Organizations also need a method for responsible knowledge production when AI participates in innovation, product development, strategic analysis, internal reports, training materials, policy drafts, customer knowledge, and research and development. HACK therefore concerns knowledge work broadly, not only academic writing.

Within that broader scope, management scholarship remains especially important. Management knowledge is not abstract information alone. It affects organizations, workers, students, leaders, institutions, and public decisions. A flawed interpretation can shape a curriculum. A distorted synthesis can guide a strategy. An unverified claim can influence a policy. A false citation can weaken a scholarly argument. A plausible but superficial analysis can make leaders more confident and less responsible.

The gap also concerns formation. The danger of AI is not only that it may be wrong. The danger is that it may become useful enough to weaken the habits of thought that human beings most need to preserve. When users outsource too much synthesis, they may lose the capacity to identify what matters. When they accept fluent outputs too quickly, they may stop reading deeply. When they use AI to avoid uncertainty, they may weaken judgment. When they let the machine supply the first structure, they may never discover their own.

HACK fills this gap by giving responsible knowledge production with AI a disciplined human-centered methodology. It does not reject AI. It refuses intellectual surrender. It treats AI as a powerful contributor of outputs within a process governed by human authorship, human interpretation, human verification, and human accountability.

10.4 FILE positioning and contribution

HACK gives FILE its methodology for responsible knowledge production.

HACK is FILE's human-centered methodology for responsible AI-assisted knowledge production. It stands for Human-AI Co-created Knowledge. It explains how human judgment must govern, curate, verify, and give meaning to AI outputs, and how the quality of that governance shapes the quality of the resulting knowledge.

HACK applies FILE's Augmented Intelligence to knowledge production. AI can extend inquiry, but it does not replace authorship, verification, interpretation, or responsibility.

In HACK, "co-created" does not mean AI co-authorship, equal agency, or delegated responsibility. It means that AI provides outputs within a knowledge process governed by human judgment, verification, authorship, and accountability. The machine may assist. It does not become the author, the scholar, or the accountable party.

HACK belongs after MAIL/MLT in the FILE Odyssey because education forms leaders, students, researchers, and professionals for work transformed by AI, while HACK defines how knowledge production can remain responsible when AI participates in research, education, organizational analysis, innovation, research and development, synthesis, critique, drafting, and revision. MAIL/MLT asks how people must be formed. HACK asks how they must produce knowledge once AI enters the intellectual process.

HACK contributes a method for using AI in intellectual work without confusing output with knowledge. AI can assist exploration, drafting, critique, comparison, organization, translation, and synthesis. Yet knowledge becomes responsible only when human beings govern the question, evaluate the evidence, interpret the meaning, verify the claims, and take responsibility for the final result.

HACK insists that human inquiry must begin before machine assistance. Before asking AI for help, the human should ask: What do I already think? What do I need to know? What is the problem? What assumptions am I making? What would a good answer need to include? What are the stakes? This discipline protects human agency. It ensures that the human question comes first, and that AI enters only afterward as support.

HACK treats AI outputs as material to challenge and verify, not as an oracle. An oracle gives answers. HACK asks leaders, scholars, educators, and professionals to use AI outputs to generate alternatives, expose assumptions, produce counterarguments, clarify tensions, identify blind spots, and improve expression. The question is not: what does AI say? The question is: how does this output help human thought become clearer, and where might it mislead?

HACK makes verification part of judgment. Verification is not a bureaucratic step. It applies to factual claims, citations, quotations, data, summaries, interpretations, and claims that affect decisions. The more consequential the output, the stronger the verification burden. A scholar who uses AI to summarize a source remains responsible for checking it. A leader who uses AI to generate analysis remains responsible for testing it. An author who uses AI to support drafting remains responsible for every claim in the final text.

HACK protects human voice and final responsibility. AI can make writing smoother, but it can also flatten voice. It can remove tension, courage, rhythm, and intellectual personality. HACK asks whether the output clarifies human thought or replaces it. It asks whether the final text still carries the author's judgment, meaning, and responsibility.

HACK also protects the Relational Commons. Knowledge production is not private technique only. It shapes trust, legitimacy, education, public understanding, and institutional responsibility. When AI-assisted knowledge is careless, unverified, or falsely attributed, it damages trust. When it is governed by human judgment and transparent responsibility, it can strengthen the conditions under which people learn, deliberate, contest, and act together.

HACK is therefore not a prompt recipe or an efficiency model. It is a human-centered method for governing AI-assisted knowledge production through judgment, verification, meaning, authorship, and accountability. It adds to responsible AI and augmented-intelligence debates by focusing on knowledge production itself.

HACK shows how AI can contribute to intellectual work without being mistaken for authors, authorities, or moral agents.

The next question is which human capacities must remain irreducibly human when AI becomes more capable.

11. Contribution 9: A human-centered philosophy. Human Irreducibility

11.1 The AI problem

AI can now imitate many outputs associated with human capacities. It can produce emotionally appropriate language, generate moral arguments, summarize ethical dilemmas, simulate care, adapt tone across cultures, support judgment-like recommendations, create persuasive narratives, and respond to human distress with fluent reassurance. These capacities are powerful. They can support leaders, managers, educators, researchers, and organizations. They can also create a philosophical danger: societies may begin to confuse simulated human outputs with morally equivalent human capacities.

The problem is not only that AI may be wrong. Error matters, but error is not the deepest issue. The deeper issue is that AI can become persuasive enough to make replacement appear natural. A generated expression of empathy can resemble care. A moral analysis can resemble conscience. A culturally adapted response can resemble belonging. A recommendation can resemble judgment. A fluent answer can mimic understanding without possessing it. In each case, AI can imitate the visible output of a human capacity while lacking the human condition that gives that capacity moral weight.

Human Irreducibility begins from a simple but decisive claim: artificial intelligence can simulate human outputs, but **simulation is not equivalence**. AI can simulate care, generate moral language, imitate cultural sensitivity, and produce emotionally appropriate text. But it does not suffer consequences, owe responsibility, belong to a culture, carry moral agency, or stand answerable before another human being. It can produce the form of a response without inhabiting the human reality from which responsibility, conscience, vulnerability, and meaning arise.

This distinction is central for leadership. Leaders do not merely process information. They assume responsibility for decisions that affect human beings. They interpret ambiguity. They protect dignity. They decide under uncertainty. They carry consequences. They answer to others. AI can support these acts, but it cannot become the human subject who bears them. Usefulness does not create moral responsibility. A machine can generate an output, but it cannot answer for that output as a moral subject.

The AI problem is therefore philosophical before it is operational. If organizations treat simulated care as equivalent to care, simulated judgment as equivalent to judgment, or generated moral language as equivalent to conscience, they weaken the human foundation of leadership. Human leadership remains necessary not because humans are always more efficient, but because **responsibility cannot be automated**.

11.2 Literature review

Human Irreducibility belongs to a long intellectual field concerned with human dignity, moral agency, responsibility, embodiment, care, recognition, capability, relationality, and meaning. It

can be situated in philosophical anthropology, phenomenology, ethics, care theory, recognition theory, humanistic management, philosophy of mind, and leadership ethics.

Human dignity provides the first foundation. Modern moral philosophy has often treated the human person as more than an instrument, object, or unit of calculation. Kant's account of dignity, for example, insists that persons must be treated as ends and not merely as means (Kant 1785). Human rights traditions develop this insight institutionally by grounding law and politics in the protection of persons. Human Irreducibility extends this concern into the age of artificial intelligence and robotics. It asks what happens when machines can produce outputs that resemble human capacities while lacking personhood, responsibility, and moral standing.

Moral agency provides the second foundation. Moral agency requires more than output generation. It requires responsibility, answerability, intention, judgment, and the capacity to be addressed by moral claims. Philosophers of responsibility have long shown that action cannot be separated from accountability (Jonas 1984; Ricoeur 1992). Human Irreducibility applies this logic to leadership in the age of artificial intelligence. AI may support a decision process, but it does not become the moral agent of the decision. It cannot bear guilt, remorse, courage, conscience, or obligation.

Embodiment and phenomenology provide a third foundation. Human beings do not encounter the world as abstract processors. They live through bodies, histories, vulnerability, mortality, memory, culture, and relation. Phenomenological traditions emphasize that human perception and meaning are embodied and situated (Merleau-Ponty 1945). This matters because AI can process language about vulnerability without being vulnerable. It can speak about grief without mourning. It can classify risk without living under risk. It can produce words about care without being bound by the human exposure that care requires.

Care ethics and relational theory deepen the argument. Care is not merely a sentiment. It is a practice of attention, responsiveness, responsibility, and relation (Gilligan 1982; Noddings 1984; Tronto 1993). Recognition theory adds that human beings need to be seen, respected, and acknowledged as subjects, not merely processed as data points (Honneth 1995; Taylor 1992). Human Irreducibility therefore rejects the reduction of human beings to inputs, profiles, scores, prompts, or predicted behaviors. Leadership after AI must protect the conditions under which persons remain visible as persons.

Capability theory also belongs here. Sen and Nussbaum argue that human development concerns what people are able to be and do, not only what systems can produce (Sen 1999; Nussbaum 2011). Human Irreducibility aligns with this concern. It asks whether AI helps protect human agency, dignity, judgment, and flourishing, or whether it narrows the human being into a function optimized by machines.

Philosophy of mind clarifies another boundary. A system can manipulate symbols, generate language, and produce appropriate responses without possessing consciousness, lived understanding, or moral experience. The point is not to settle every debate about machine intelligence. The point is narrower and more important for leadership: output similarity does not establish moral equivalence. Human Irreducibility therefore gives leadership theory a

boundary against technological reductionism. It refuses to confuse performance with personhood.

Human Irreducibility can be deepened through a broader anthropological account of the human being. Its ten pillars (judgment, discernment, conscience, moral agency, dignity, meaning, responsibility, care, trust, and love) name the irreducible capacities FILE protects, and this anthropology shows why these capacities belong to human beings as responsible, vulnerable, embodied, relational, cultural, political, moral, imaginative, mortal, and capable of refusal. This broader anthropology supports FILE without replacing it. It explains why the ten pillars are not abstract virtues: they belong to the human condition itself.

11.3 Research gap

Many approaches to AI ethics use the language of human-centered design, human oversight, safety, fairness, transparency, explainability, privacy, and accountability. These concerns are necessary. They protect important dimensions of responsible AI. Yet they often leave a deeper philosophical gap underdeveloped. They do not always specify which human capacities cannot be replaced, delegated, simulated into equivalence, or made morally interchangeable with machine outputs.

Human Irreducibility matters because care, judgment, responsibility, and moral agency must remain human. AI may support human action, but it must not replace the human presence and accountability that these capacities require.

Existing AI governance often asks whether systems are reliable, safe, fair, transparent, and accountable. These questions are indispensable, but they do not fully answer the philosophical question: what must never be replaced? A system can be accurate and still degrade dignity. It can be efficient and still weaken responsibility. It can be explainable and still have a narrow meaning. It can be helpful and still encourage human beings to surrender capacities they should preserve.

Leadership theory also needs this boundary. Traditional leadership theories often emphasize influence, vision, motivation, transformation, ethics, adaptation, and decision-making. These remain important. But AI introduces a new problem. Leaders must now protect the human capacities that make leadership morally possible in the first place. A leadership theory for the age of artificial intelligence and robotics cannot only ask how leaders perform. It must ask what leadership must refuse to delegate.

Human Irreducibility addresses this gap. It states that some capacities are not merely difficult to automate. They are morally non-substitutable. They may be supported, extended, informed, or challenged by artificial intelligence. They cannot be replaced by it. Human Irreducibility is not yet an empirically validated measurement scale. It is a philosophical boundary and research-generating proposition. Research can examine how leaders preserve responsibility, judgment, dignity, care, trust, meaning, moral agency, and discernment under algorithmic pressure.

11.4 FILE positioning and contribution

Within the FILE Odyssey, Human Irreducibility gives FILE its philosophical foundation.

Human Irreducibility is FILE's human-centered philosophy. It recognizes that certain human capacities, such as judgment, discernment, conscience, moral agency, dignity, meaning, responsibility, care, trust, and love, cannot be replaced by a machine, no matter how powerful it becomes. Human Irreducibility is the philosophical foundation of the FILE Odyssey and the line that AI must never cross.

These capacities are not decorative values added to leadership after technical decisions have already been made. They are the conditions under which leadership remains human.

This definition positions Human Irreducibility as the moral boundary of FILE. It is not anti-AI. It is not nostalgia. It is not sentimental language placed against technology. It defines the point at which augmentation must not become substitution, support must not become replacement, and simulation must not be confused with equivalence.

Human Irreducibility also clarifies the relation between the earlier contributions. LENS protects human judgment; Human Irreducibility explains why that judgment must remain human. HACK governs responsible knowledge production with AI; Human Irreducibility explains why the final act of meaning, verification, and responsibility cannot be transferred to the machine. Human Irreducibility names what cannot be replaced in the human person; the Relational Commons names the shared human world in which those irreducible capacities can survive.

Judgment is the capacity to decide under conditions of ambiguity, uncertainty, conflict, and consequence. AI can provide information, patterns, predictions, and recommendations. It cannot become the human subject who weighs meaning, assumes responsibility, and answers for a decision.

Discernment is the capacity to distinguish what matters from what merely appears relevant. AI can rank, classify, and retrieve. It cannot know the human weight of a situation from within responsibility, vulnerability, and context.

Conscience is the inner moral demand that asks not only what works, but what is right. AI can generate moral language. It cannot be bound by conscience.

Moral agency is the capacity to act under responsibility and to be answerable for action. AI can participate in processes that affect moral outcomes. It cannot become the moral agent of those outcomes.

Dignity names the worth of the human person beyond utility, productivity, classification, or performance. AI can process information about persons. It must not become the measure of their worth.

Meaning is the human capacity to interpret life, work, suffering, relation, purpose, and consequence. AI can produce language about meaning. It cannot inhabit meaning as a human being does.

Responsibility is the anchor of Human Irreducibility. A machine can generate an output, but it cannot answer for that output as a moral subject. It cannot suffer consequences, feel remorse, bear guilt, owe repair, or stand accountable before another human being.

Care is attention to another human being as vulnerable, real, and worthy of response. AI can simulate caring language. It cannot care as a human being cares.

Trust is not only confidence in performance. It is a relational condition created through responsibility, reliability, vulnerability, and recognition. AI can be reliable. It cannot enter trust as a moral subject.

Love is the deepest and least reducible human capacity in FILE. It names the human power to affirm another person beyond utility, calculation, transaction, and performance. AI can imitate affectionate language. It cannot love.

Human Irreducibility gives FILE its deepest answer to the age of artificial intelligence. Leadership is not human merely because humans still perform some tasks better than machines. Leadership is human because responsibility, dignity, meaning, care, trust, conscience, moral agency, discernment, judgment, and love cannot be transferred to a machine. The question is not only what AI can do. **The question is what human beings must never stop being.**

Human Irreducibility explains why artificial intelligence can support leadership, knowledge, analysis, education, governance, and decision-making, but cannot become morally equivalent to a human being. It also explains why FILE insists on Humans + Machines, under human leadership. The machine may assist. It may extend. It may accelerate. It may reveal. It may challenge. But it must not replace the human capacities that make leadership worthy of the name.

The future of leadership depends not only on what AI becomes, but on what human beings refuse to stop being.

Once this philosophical boundary is established, the next question is how AI must be governed so that it never crosses the line Human Irreducibility defines. That question leads to the Five Immutable Laws of AI and Robotics: HTTPS: Humanity, Truth, Transparency, Protection, and Security.

12. Contribution 10: A human-centered AI governance codex. The Five Immutable Laws of AI and Robotics: HTTPS: Humanity, Truth, Transparency, Protection, and Security

12.1 The AI problem

AI is now embedded in decision-making, communication, education, knowledge production, management, public life, and governance. It shapes what people see, what institutions measure, what organizations classify, what leaders compare, what citizens believe, and what

societies treat as true. It affects dignity, trust, security, autonomy, attention, institutional legitimacy, and the future of human societies.

The danger extends far beyond isolated misuse. Misuse matters, but it is not the whole problem. The deeper danger is that AI systems and robots may be designed, deployed, and governed without non-negotiable human-centered laws. Technical performance can improve while human dignity weakens. Optimization can increase while truth erodes. Automation can scale while transparency declines. Efficiency can expand while vulnerable people become more exposed. Capability can grow while security, freedom, and democratic trust become more fragile.

AI governance therefore cannot be reduced to technical compliance. It must address the human and civilizational stakes of design, deployment, and governance. Systems that classify people, generate information, recommend action, influence attention, or support decisions must be judged not only by performance, accuracy, or speed. They must also be judged by whether they protect humanity, truth, transparency, protection, and security.

This problem follows directly from Human Irreducibility. If certain human capacities cannot be replaced by a machine, then artificial intelligence and robotics must be governed so that they do not reduce, displace, manipulate, or violate those capacities. Human Irreducibility defines what must never be replaced; HTTPS defines the laws that must govern AI systems and robots so that this boundary is protected.

The AI problem is therefore one of governance at scale. Organizations and societies need more than principles that can be negotiated away under pressure. They need a compact, memorable, human-centered set of laws that places human beings, human societies, and the future of human civilization at the center of AI governance.

12.2 Literature review

HTTPS belongs to the fields of AI governance, responsible AI, digital ethics, human rights, human dignity, safety, transparency, accountability, security, democratic governance, sociotechnical risk, and civilizational risk.

Responsible AI provides the immediate field. It asks how AI can be designed and deployed in ways that are fair, safe, transparent, accountable, reliable, privacy-protecting, and aligned with human values. Floridi and Cowls (2019) identify substantial convergence around beneficence, non-maleficence, autonomy, justice, and explicability. Jobin, Ienca, and Vayena (2019) likewise find a large international landscape of AI ethics guidelines, with recurring principles but important divergence in their interpretation and implementation. Shneiderman (2022) places human needs, reliable systems, and human responsibility at the center of human-centered AI. Mittelstadt (2019), however, demonstrates why agreement at the level of principles does not by itself guarantee ethical practice when professional norms, implementation mechanisms, incentives, and accountability remain insufficient. HTTPS accepts the seriousness of this existing field but reorganizes the governance demand around five non-negotiable laws: Humanity, Truth, Transparency, Protection, and Security.

Human rights and human dignity provide the normative foundation. AI affects human beings before it affects markets, systems, or organizations. It can influence access to education, employment, healthcare, credit, public services, information, mobility, and recognition. Sen's account of substantive freedom and human agency and Nussbaum's capabilities approach place what persons are actually able to be and do at the center of human development (Sen 1999; Nussbaum 2011). Shneiderman (2022) carries a related human-centered priority into technological development. These traditions do not supply HTTPS, but they reinforce the proposition that technology must remain ordered toward human agency, capability, dignity, and flourishing. **Humanity is not one value among others. It is the first condition of legitimate AI governance.**

Truth gives HTTPS its knowledge foundation. AI can generate language, images, summaries, classifications, recommendations, and synthetic media at scale. This power can support knowledge, but it can also intensify deception, fabrication, manipulation, and disorder in shared knowledge. Bouilloud, Deslandes, and Mercier (2019) establish truth-telling as an ethical responsibility of organizational leadership. Chesney and Citron (2019) show how deepfakes can scale deception and create serious harms for privacy, democracy, and national security. Democratic societies, universities, organizations, and public institutions depend on shared conditions of truth. Without them, leadership loses the ground on which responsible decision-making depends.

Transparency gives HTTPS its institutional foundation. Human beings must be able to understand when and how AI shapes decisions, outputs, classifications, or recommendations. Transparency does not mean that every technical detail will be simple, nor that visibility alone creates accountability. Ananny and Crawford (2018) demonstrate the limits of treating transparency as sufficient for understanding and governing algorithmic systems. Pasquale (2015) exposes the power and informational asymmetries created by black-box systems, while Raji et al. (2020) show the importance of institutional auditing mechanisms across the AI system lifecycle. AI must remain intelligible, accountable, and contestable enough for human beings to govern its role in social life. A system that cannot be questioned becomes a structure of power without adequate human answerability.

Protection gives HTTPS its social foundation. AI affects people unequally. Vulnerable communities, workers, students, patients, citizens, and marginalized groups can be harmed by classification, exclusion, surveillance, manipulation, or automated indifference. Selbst et al. (2019) show why fairness cannot be abstracted from the sociotechnical systems in which decisions operate. O'Neil (2016) documents how opaque quantitative systems can scale asymmetric harm, and Eubanks (2018) shows how automated public systems can intensify exclusion and burdens for poor and vulnerable populations. Protection therefore requires attention to foreseeable harm, exploitation, exclusion, manipulation, dehumanization, and institutional vulnerability. It is not enough to build powerful systems. Institutions must ask whom those systems may expose, silence, or damage.

Security gives HTTPS its civilizational foundation. AI can create technical risks, organizational risks, political risks, psychological risks, and civilizational risks. Brundage et al. (2018) provide

the most direct security anchor by examining malicious uses of AI across digital, physical, and political domains. Russell (2019) addresses the broader AI control problem and the systemic risks created when increasingly capable systems pursue objectives that are inadequately aligned with human purposes. Security therefore concerns more than cybersecurity or technical protection. It concerns institutional resilience, democratic stability, information integrity, human autonomy, and the conditions under which human societies remain able to govern themselves responsibly.

Sociotechnical theory completes the field. Trist and Bamforth (1951) established that technical and social arrangements must be understood together, while Orlikowski (2000) shows how technologies and organizational structures are constituted through practice. AI is never only technical. It is embedded in organizations, incentives, norms, institutions, cultures, markets, laws, and forms of power. HTTPS therefore treats governance as a human and institutional obligation. It asks not only whether AI works, but whether AI systems and robots are designed, deployed, and governed under laws that protect human beings.

12.3 Research gap

Existing AI governance approaches offer substantial bodies of principles, recommendations, guidelines, checklists, risk categories, audits, documentation practices, and compliance systems. These instruments are useful. They help institutions identify risks, standardize practices, improve accountability, and create oversight. Jobin, Ienca, and Vayena (2019) demonstrate both significant convergence and continued variation across the international principles landscape, while Mittelstadt (2019) shows why principles can remain weak when implementation and accountability are underdeveloped.

The first gap is fragmentation. AI governance discussions often separate fairness, transparency, accountability, safety, privacy, security, human rights, and human oversight into distinct categories. This can be analytically useful, but it can also obscure the unity of the human-centered governance problem. These concerns are differently named, differently operationalized, and unevenly institutionalized. HTTPS addresses this fragmentation by compressing the governance demand into five laws that are simple enough to remember and serious enough to guide institutional action.

The second gap is negotiability. Principles may coexist without establishing which boundaries are non-negotiable when values conflict with speed, profit, convenience, or competition. HTTPS uses the language of laws because some boundaries should not depend on managerial preference. Humanity, Truth, Transparency, Protection, and Security are not optional values. They are non-negotiable conditions for legitimate AI governance.

The third gap is bureaucratization. Governance can become a documentation ritual. Institutions may complete forms, perform audits, or produce policy language without changing the design, deployment, and governance of AI systems and robots. HTTPS resists this reduction. It does not ask only whether procedures exist. It asks whether technology remains ordered to Humanity, Truth, Transparency, Protection, and Security.

Governance too often becomes reactive, with institutions responding only after harm is visible. HTTPS is preventive. It asks institutions to govern AI systems and robots before deception, opacity, exposure, exclusion, or insecurity become normalized. It brings the philosophical boundary of Human Irreducibility into the practical domain of technological design, deployment, and governance.

The originality claim is therefore architectural. HTTPS does not claim to have invented each normative concern separately. Its contribution lies in their exact selection, integration, protected order, and formulation as five immutable laws governing artificial intelligence and robotics: Humanity, Truth, Transparency, Protection, and Security.

12.4 FILE positioning and contribution

HTTPS is a **human-centered AI governance codex**. It extends FILE into human-centered AI governance and translates Human Irreducibility into that governance domain.

Its exact architecture is **The Five Immutable Laws of AI and Robotics: HTTPS: Humanity, Truth, Transparency, Protection, and Security**. These five non-negotiable laws must govern the design, deployment, and governance of AI systems and robots to protect human beings, human societies, and the future of human civilization.

This positions HTTPS as Contribution 10 in the twelve-part presentation. Human Irreducibility defines what AI must never replace. HTTPS defines the laws that must govern AI systems and robots so that this boundary is protected. Human Irreducibility is the philosophical boundary. HTTPS is the AI governance that follows from it.

HTTPS is distinct from LENS and RC. LENS protects the second-order governance of human judgment under decision-making conditions shaped by AI. RC names the human world FILE seeks to protect. HTTPS protects that world at the level of AI design, deployment, and governance. It does not replace existing regulations, standards, or institutional governance regimes. It gives FILE five non-negotiable laws for human-centered AI governance.

Humanity means that AI must remain ordered to the protection of human beings, human dignity, human agency, and human flourishing. It must serve persons, not reduce them to data, targets, profiles, scores, or instruments. H is the first letter in HTTPS because Humanity is the first, the highest, and the supreme law of AI governance: no technical gain, organizational benefit, or strategic advantage can justify the diminishment, instrumentalization, or endangerment of human beings. Humanity must be protected at all costs.

Truth means that AI must not be used to normalize deception, fabrication, manipulation, or disorder in shared knowledge. Responsible leadership requires shared reality. Without truth, there can be no trustworthy judgment, education, governance, or democracy. Truth means that AI must not become an engine of synthetic falsehood, strategic confusion, counterfeit expertise, or informational domination.

Transparency means that AI must remain intelligible, accountable, and contestable enough for human beings to understand when and how it shapes decisions, outputs, classifications, or

recommendations. People must not be governed by invisible systems they cannot identify, question, understand, or challenge. Transparency protects the human capacity to contest power.

Protection means that AI must protect vulnerable people, communities, institutions, and societies from foreseeable harm, exploitation, exclusion, and dehumanization. AI governance must begin with those most likely to be harmed by automation, classification, surveillance, manipulation, or institutional indifference. A system that benefits the powerful while exposing the vulnerable violates the human-centered logic of HTTPS.

Security means that AI must be designed, deployed, and governed in ways that protect human beings and human societies from technical, organizational, political, psychological, and civilizational threats. Security includes cybersecurity, institutional resilience, democratic stability, information integrity, psychological safety, and protection against systemic misuse. Security recognizes that powerful AI can create risks that exceed any single organization.

Together, these five laws translate Human Irreducibility into human-centered AI governance. They also protect the Relational Commons because trust, dignity, psychological safety, and meaning cannot survive in an environment of dehumanization, deception, opacity, exposure, or insecurity.

HTTPS gives FILE its final AI governance contribution. FILE names the intelligences required for leadership. The 7E Cascade develops them. RC defines the shared human world FILE seeks to protect, cultivate, and renew. MAIL/MLT forms leaders to act within that world. HACK governs responsible knowledge production. LENS formalizes the second-order governance of human judgment. Human Irreducibility defines the philosophical boundary. HTTPS states the immutable laws that must govern AI systems and robots so that the boundary is never crossed.

The result is a clear final claim: artificial intelligence and robotics must be governed by Humanity, Truth, Transparency, Protection, and Security because leadership after AI is responsible not only for performance, but for the human future that performance must never betray.

After defining the philosophical boundary of FILE and the laws that must govern AI systems and robots, the paper turns to BASIC and RULES, the human-centered public-policy contribution of the FILE Odyssey.

13. Contribution 11: A human-centered public policy. BASIC: Basic Income, AI Access, Skills, Inclusion, and Care, and its French equivalent, RULES : Revenu Universel, Lutte contre la pauvreté et les inégalités, Éducation et accès à l'IA et Solidarité

13.1 The AI problem

Artificial intelligence changes not only organizations but also the distribution of work, access, skills, opportunity, and security. A human-centered management theory therefore reaches a public-policy boundary when technological transformation affects who can participate in economic and social life.

Robotics makes this boundary equally concrete. AI systems and robots can reorganize tasks, alter skill demand, affect employment and wages, change access to productive technologies, and reshape opportunities for participation. They can also widen the distance between people who can use new technological capabilities and those who cannot. None of these effects is mechanically negative or positive. The issue is distribution and participation: who gains new capabilities, who absorbs disruption, who remains able to learn and work, and who is excluded from institutions increasingly mediated by technology.

The public-policy problem is to ensure that AI-driven transformation does not leave human beings without the material, educational, technological, social, and caring conditions required to participate with dignity.

Leadership cannot remain human-centered inside the organization while treating these surrounding conditions as somebody else's problem. When AI and robotics transform work and participation at scale, material security, meaningful technological access, learning, inclusion, and care become part of the institutional environment in which leadership itself is exercised.

13.2 Literature review

Public policy enters first through work. Autor (2015) shows why automation cannot be reduced to a simple story of machines eliminating jobs: technology substitutes for some tasks while complementing others and changing the composition and rewards of work. Acemoglu and Restrepo (2019) sharpen this task-based view by distinguishing displacement from the creation or reinstatement of tasks in which labor retains a comparative advantage. Robotics makes the distributional issue tangible. In U.S. commuting zones with greater exposure to industrial robots, Acemoglu and Restrepo (2020) identify negative effects on employment and wages. This is not evidence that robots always reduce employment. It is evidence that technological change can impose geographically and socially uneven labor-market consequences that public policy cannot treat as incidental.

Once disruption can affect income and participation, material security becomes a legitimate part of the architecture. Van Parijs and Vanderborght (2017) provide a mature philosophical, economic, and institutional defense of unconditional basic income as an instrument connected to freedom and economic security. Their work does not establish basic income as the necessary

response to AI, nor does it anticipate BASIC. It establishes the seriousness of the prior Basic Income tradition. FILE's move is different: it places Basic Income inside a wider architecture in which material security matters because technological disruption should not automatically become exclusion from economic and social participation.

Material security, however, is not technological participation. The digital-divide literature shows why access must be understood as layered rather than binary. Warschauer (2003) shifts attention from mere possession or connectivity toward the conditions that allow people to use technology meaningfully in social practice. Hargittai (2002) demonstrates substantial differences in people's effective online skills even among users, while van Dijk (2005) connects access, skills, usage, participation, and persistent inequality. These works predate contemporary AI and should not be retroactively described as AI-access theories. They provide the intellectual foundation for FILE's extension to AI Access: meaningful access to increasingly consequential AI capabilities, considered together with availability, affordability, skills, literacy, agency, safety, privacy, social support, and the institutional conditions of effective participation.

Access without capability can remain formal. Research on technological change has long shown that the content of work and the demand for skills move with technology. Autor, Levy, and Murnane (2003) document how computerization changes the allocation of routine and non-routine tasks. Deming (2017) shows the growing labor-market value of social skills and interaction. For BASIC, the implication is deliberately bounded: Skills are a public-policy condition of participation. People need opportunities to learn, adapt, work, and use new technologies meaningfully as tasks and institutions change. This is not a second theory of FILE's human capabilities; it is a claim about access to the learning conditions required for participation.

Even income, access, and skills do not guarantee inclusion. Sen (1999) defines development through the expansion of substantive freedoms, while Nussbaum (2011) asks what people are actually able to do and to be, making capability and dignity central to substantive opportunity. Eubanks (2018) shows the other side of the problem: automated public systems can classify, burden, and exclude vulnerable populations rather than simply deliver neutral administrative efficiency. Inclusion therefore cannot mean physical presence, nominal eligibility, or possession of a technological tool. It concerns the real possibility of benefiting, exercising agency, accessing institutions, and avoiding systematic exclusion.

A human-centered public policy also fails if it treats paid work as the whole of human life. Tronto (1993) establishes care as an ethical and political practice structured by attentiveness, responsibility, and responsiveness, and Tronto (2013) extends care into questions of democracy, equality, institutions, and the distribution of caring responsibilities. Folbre (2001) adds an economic and public-policy dimension by examining the value and organization of care work and the tensions between market incentives and caring responsibilities. Care is therefore not a sentimental final item in BASIC. It recognizes that human beings remain dependent, responsible for others, embedded in families and communities, and reliant on institutions of support even as AI and robotics transform productivity and employment.

These literatures converge on a harder policy question than any single instrument can answer. A secure income can coexist with technological exclusion. Access can exist without the skills or agency required for meaningful use. Skills can exist while institutions continue to exclude. Formal inclusion can coexist with care systems that leave dependency and responsibility unsupported. BASIC treats these conditions as mutually enabling rather than interchangeable. Its unit of concern is meaningful human participation: the ability to remain materially secure, technologically capable, able to learn, substantively included, and supported within a society in which care still matters.

BASIC and RULES do not replace these traditions or existing public-policy systems.

13.3 Research gap

The relevant scholarship is substantial. The gap addressed here is not absence but fragmentation. Labor economics, basic-income scholarship, digital inequality research, skills research, capability approaches, studies of automated exclusion, care ethics, and the economics of care have developed through different disciplinary and policy traditions.

Fragmentation matters because AI and robotics increasingly connect problems that policy institutions still address through separate instruments, mandates, and disciplinary lenses. The weakness is not that any one field is missing, but that material security, technological participation, capability formation, inclusion, and care can become disconnected precisely when technological change makes their interaction more consequential. The policy problem is therefore one of coordination around meaningful human participation, not simply the presence or absence of individual instruments.

BASIC integrates these five conditions into one human-centered public-policy architecture for the age of artificial intelligence and robotics. Its originality lies in that specific integration, not in claiming ownership of the constituent literatures. The contribution is therefore synthetic and architectural: it connects material security, AI Access, Skills, Inclusion, and Care around a common problem of participation under technological transformation.

13.4 FILE positioning and contribution

BASIC means Basic Income, AI Access, Skills, Inclusion, and Care.

Its French equivalent is RULES : Revenu Universel, Lutte contre la pauvreté et les inégalités, Éducation et accès à l'IA et Solidarité.

Basic Income provides a material-security floor so that technological disruption does not automatically become exclusion from economic and social life. Yet material security is not technological participation: meaningful AI Access also depends on affordability, agency, safety, privacy, and the skills and institutional conditions required for effective use. Skills matter because people must be able to learn, adapt, work, participate, and use technological capabilities as tasks and institutions change. Inclusion concerns the further question of whether that capability translates into substantive participation in economic, institutional, technological, and social life. Care keeps the architecture from reducing human participation

to labor-market productivity by recognizing the relationships of dependency, responsibility, support, and caring labor on which societies continue to depend.

The five dimensions are designed to work together. None is sufficient in isolation, and none substitutes for the others. BASIC is therefore not a list of benefits but a compact architecture for the conditions under which people can remain capable of participating with dignity as AI and robotics reshape work and social institutions.

Together, BASIC and RULES form the eleventh original contribution of the FILE Odyssey: a human-centered public policy for the age of artificial intelligence and robotics.

Through this contribution, the FILE Odyssey crosses a bounded public-policy frontier. Leadership and management take place inside social and institutional arrangements that shape whether people can work, learn, access technology, participate, care and be cared for, and live with dignity. When AI and robotics transform those conditions at scale, a human-centered theory of leadership cannot treat their distribution as external to the human consequences of technological change.

This contribution remains connected to RC because public policy shapes the conditions under which trust, dignity, inclusion, meaning, care, and participation can be sustained beyond the boundaries of a single organization.

14. Contribution 12: A human-centered civilization. Inalienable Leadership: the overarching metatheory of leadership in the age of artificial intelligence and robotics

14.1 The AI problem

As artificial intelligence and robotics become more capable, the central leadership question eventually reaches beyond organizations. The issue becomes who retains final authority and responsibility for the purposes, institutions, and direction of civilization when cognition, analysis, coordination, creativity, and execution can increasingly be distributed across humans and machines.

The distinction is fundamental: **Operational autonomy $\neq$ normative authority**. A machine may generate actions, recommend decisions, classify, coordinate, execute, operate autonomously, control physical processes, or act through robots. These capabilities can constitute significant operational autonomy. They do not, by themselves, create legitimate authority over human purposes, democratic legitimacy, moral standing, final responsibility, the right to determine institutional ends, or the right to determine civilizational ends.

Operational delegation can expand dramatically without making final human authority and responsibility delegable. **Capability changes the scale of the problem; it does not remove the human obligation to answer for direction and consequence.**

14.2 Literature review

The arguments developed in this paper already connect leadership theory, management, sociotechnical systems, human-centered AI, Human Irreducibility, judgment, and AI governance. Inalienable Leadership integrates these strands at the highest conceptual level of the FILE Odyssey. Metatheory provides a legitimate way to relate distinct theories without dissolving them into one model. Abrams and Hogg (2004) show how a metatheoretical perspective can clarify relationships among theoretical levels and organize different explanatory traditions around higher-order questions.

Leadership itself has long exceeded technical administration. Selznick (1957) places institutional leadership at the level of purpose, character, and the responsibilities through which an organization becomes an institution rather than only an efficient instrument. Weber's account of rational-legal authority locates legitimate authority within an institutional order of offices and rules (Weber 1978). This distinction matters because technical competence and legitimate institutional authority are not the same category. Weber does not establish the Inalienable Leadership claim that final civilizational authority remains human. That extension is the theoretical move developed here.

Technological power enlarges both the reach of action and the horizon of responsibility. Jonas (1984) argues that modern technological capacity creates responsibilities extending toward future human life. Weizenbaum (1976), from another direction, distinguishes what computers can perform from what human beings should decide to entrust to computation. Greater computational capability can therefore enlarge the field of possible action without deciding which purposes that action should serve.

Automation complicates this relationship in practice. Bainbridge (1983) showed that the more routine control is automated, the more human operators may be left responsible for exceptional intervention precisely when their routine involvement and opportunities to maintain expertise have been reduced. Automation can therefore increase while effective human control declines. Formal retained responsibility is not the same as substantive control or practical competence.

The problem becomes sharper when autonomous and learning systems behave in ways that are difficult to predict or attribute. Matthias (2004) identifies a responsibility gap when conventional forms of responsibility attribution encounter learning automata whose behavior cannot be fully foreseen. Elish (2019) shows a related sociotechnical danger through the concept of moral crumple zones, in which nearby human operators can absorb blame despite having limited control over the wider automated system. Inalienable Leadership cannot mean that the nearest human operator should always be blamed. Final human authority and responsibility must operate at the appropriate leadership, institutional, and civilizational levels, where purposes are authorized and systems are governed.

Work on meaningful human control provides a direct bridge to autonomous systems and robotics. Santoni de Sio and van den Hoven (2018) connect meaningful control to the responsiveness of systems to relevant human reasons and to the traceability of outcomes to

responsible human agents. Mindell (2015) challenges simplistic narratives of complete robotic autonomy by showing how advanced machines remain embedded in human technical, organizational, and operational systems. Operational autonomy can therefore increase without erasing the human institutions that design, authorize, deploy, maintain, and depend upon it.

Pasquale (2020) extends this concern to the preservation of human expertise and to boundaries on inappropriate machine substitution in AI and robotics. Rahwan et al. (2019) take a different step by treating machine behavior itself as an increasingly consequential object of scientific inquiry within social systems. Their contribution is descriptive and analytical: machines increasingly participate in environments shared with people and institutions and therefore require systematic study. They do not establish that machines lack legitimate final authority. That normative conclusion belongs to Inalienable Leadership.

Inalienable Leadership does not erase the distinct functions of FILE, SQL, LENS, RC, Human Irreducibility, HTTPS, or the other contributions. It provides the metatheoretical principle under which their human-centered direction can be understood together. The literature makes clear that machine capability, operational autonomy, human control, institutional authority, and responsibility are distinct problems. A system can act without possessing the right to define the purposes it serves. This is the boundary between operational autonomy and normative authority.

14.3 Research gap

The relevant scholarship is substantial. It already addresses institutional leadership and purpose, legitimate authority, technological responsibility, human judgment, automation paradoxes, responsibility gaps, moral crumple zones, meaningful human control, robotics, human expertise, and machine behavior.

Operational autonomy $\neq$ normative authority. Increasing machine autonomy changes what machines can do. It does not itself answer who has legitimate authority, who defines collective purposes, who determines institutional ends, or who remains finally answerable for civilizational direction.

The gap is the absence of an overarching leadership metatheory specifically organized around the principle that final human authority and responsibility for purposes, institutions, and the direction of civilization remain inalienable under conditions of increasingly autonomous artificial intelligence and robotics.

Inalienable Leadership addresses this gap by making that inalienability the governing principle at the institutional and civilizational level. Machine capability may transform delegation. It does not confer normative authority.

14.4 FILE positioning and contribution

Inalienable Leadership is the twelfth and culminating original contribution of the FILE Odyssey and the overarching metatheory of leadership in the age of artificial intelligence and robotics.

It holds that humanity may share work and intelligence with machines but cannot delegate, transfer, surrender, or relinquish its final authority and responsibility for the purposes, institutions, and direction of civilization.

Inalienable Leadership integrates, grounds, governs, and gives civilizational direction to FILE and the complete FILE contribution system without replacing their distinct identities.

It places the development, deployment, and governance of AI systems and robots under human leadership, ensuring that technological progress remains accountable to human dignity, democratic choice, the common good, and the long-term future of humanity, and that technology serves civilization rather than governs it.

Its governing order is the order already expressed throughout the FILE Odyssey: **Humans + Machines, under human leadership.**

15. The Twelve Contributions as a Human-Centered Research Program

15.1 The cumulative movement of the FILE Odyssey

The FILE Odyssey begins with a simple problem and ends with a governance demand. The problem is that AI now mediates management, leadership, education, knowledge production, judgment, and institutional life. The governance demand is that AI must remain ordered to human beings, human responsibility, human dignity, and human civilization.

The twelve contributions developed in this paper form a cumulative movement. Leadership = Intelligence + Character + Judgment provides the human-centered definition of leadership. FILE names the five intelligences required for responsible leadership in the age of artificial intelligence and robotics. FILE³, FILE⁵, FILE⁷, and the 7E Cascade explain how those intelligences mature through development, effectiveness, excellence, ecosystems, empowerment, execution, and embodiment. SQL names the character of the embodied human leader. LENS protects the conditions of human judgment. RC names the human world those intelligences must protect, cultivate, and renew. MAIL and MLT translate FILE into education. HACK translates it into a methodology of responsible knowledge production with AI. Human Irreducibility gives FILE its philosophical foundation. HTTPS translates that boundary into AI governance. BASIC and RULES extend the program into public policy. Inalienable Leadership provides the culminating civilizational and metatheoretical integration.

This movement matters because AI does not create one isolated management problem. It creates a layered transformation of work, knowledge, education, organization, judgment, power, meaning, and governance. A technical response alone cannot address this transformation. A narrow ethical checklist cannot address it either. The age of artificial intelligence and robotics requires a research program that connects leadership intelligence, human development, relational life, management education, knowledge production, judgment, philosophy, and AI governance.

FILE answers this need by placing human responsibility at the center. Its governing claim is not that AI should be rejected. Its governing claim is that AI must be governed by human judgment, human meaning, human responsibility, and laws that protect human beings. The more powerful artificial intelligence becomes, the more decisive human leadership becomes.

15.2 The internal boundaries of the twelve contributions

The strength of FILE depends on keeping its contributions distinct while showing how they reinforce one another.

A final boundary check is useful because the twelve contributions are easy to confuse. The core FILE formula remains:

$$\text{FILE} = \text{AI} + \text{EQ} + \text{CQ} + \text{PQ} + \text{AQ}$$

AI still means Augmented Intelligence (“Mind”). EQ means Emotional Intelligence (“Heart”). CQ means Cultural Intelligence (“World”). PQ means Political Intelligence (“Society”). AQ means Adaptive Intelligence (“Evolution”). None of the later contributions adds a sixth intelligence.

Within this discipline, Augmented Intelligence includes Cognitive Intelligence and Complexity Intelligence; Political Intelligence includes Purpose Intelligence, Moral Intelligence, and Sustainable Intelligence; and Adaptive Intelligence includes Decision Intelligence, Judgment Intelligence, and Learning Intelligence. These are nested dimensions, not additional primary intelligences.

The 7E Cascade explains development. RC names the human world FILE protects. MAIL and MLT translate FILE into education. HACK governs AI-assisted knowledge production under human judgment, verification, authorship, and accountability.

LENS protects the second-order governance of human judgment. Human Irreducibility explains why certain human capacities cannot be replaced by machines. HTTPS states the Five Immutable Laws of AI and Robotics.

The result is an integrated management research program. Leadership = Intelligence + Character + Judgment defines leadership; FILE defines the leadership-intelligence layer; the 7E Cascade, development; SQL, character; LENS, judgment; RC, the relational destination; MAIL/MLT, education; HACK, knowledge production; Human Irreducibility, the philosophical boundary; HTTPS, AI governance; BASIC and RULES, public policy; and Inalienable Leadership, the culminating civilizational metatheory.

15.3 The central synthesis

The central synthesis of the paper is this: AI increases the need for human leadership because it increases the consequences of human judgment.

AI technologies can automate tasks, generate knowledge outputs, support decisions, classify situations, optimize workflows, simulate care, and produce persuasive language. They can help leaders see more, compare more, generate more, and act faster. Yet this power does not remove the need for human judgment. It intensifies it. Every organization using AI systems

and robots must still decide what should be measured, what should be protected, what should be contested, what should be refused, what should be explained, and who remains answerable when consequences occur.

Leadership = Intelligence + Character + Judgment defines the human leadership architecture at the center of this responsibility. FILE: The Five Intelligences of Leadership Evolution names the intelligences needed for this responsibility. The 7E Cascade explains how they mature. SQL: The Six Qualities of Leadership defines the Character dimension of leadership through six qualities: Vision, Intuition, Taste, Courage, Discipline, and Integrity. LENS: Leadership, Ethics, Nondelegation, and Sensemaking governs the conditions under which leaders perceive, evaluate, and decide. RC: The Relational Commons names the human ecosystem that must not be depleted. MAIL: Management, AI, and Leadership and MLT: Management, Leadership, and Technology form leaders and students for this new condition. HACK: Human-AI Co-created Knowledge governs the production of knowledge when AI contributes outputs. Human Irreducibility names the philosophical line that AI must never cross. The Five Immutable Laws of AI and Robotics: HTTPS: Humanity, Truth, Transparency, Protection, and Security states the non-negotiable laws that must govern AI systems and robots so that human beings, human societies, and the future of human civilization remain protected. BASIC and RULES extend the architecture into public policy. Inalienable Leadership integrates the complete contribution system as the overarching metatheory of leadership in the age of artificial intelligence and robotics.

This synthesis gives FILE its contribution to management, leadership, and the social sciences. It refuses both technocratic overconfidence and nostalgic rejection. It does not ask leaders to choose between humans and machines. It asks leaders to put the relation in the right order: **Humans + Machines, under human leadership.**

15.4 Managerial implications: using FILE as a diagnostic lens

FILE also has direct managerial implications. Managers, leadership-development teams, executives, entrepreneurs, founders, CEOs, boards, deans, educators, researchers, public institutions, and organizations can use it as a diagnostic lens for transformation shaped by AI.

For managers, executives, and entrepreneurs, FILE asks whether AI adoption is governed as a human leadership problem rather than only as a technology project. The five intelligences diagnose whether leaders govern AI systems and robots responsibly, protect trust and psychological safety, interpret outputs across cultural and institutional contexts, understand how metrics redistribute power, and learn from consequences rather than normalizing them.

For organizations, RC provides an early warning system: an AI project can improve efficiency while depleting trust, dignity, psychological safety, meaning, voice, care, ability to challenge, or belonging.

For business schools and leadership programs, MAIL and MLT define the educational implication: AI cannot remain an elective technical topic beside management education because managers now operate where Management, AI, and Leadership meet.

For scholars, consultants, analysts, and students, HACK insists that faster synthesis is not better knowledge unless verification, source discipline, authorial voice, and final accountability remain humanly governed.

For boards, senior leaders, and governance committees, LENS asks whether leaders govern the metrics, classifications, dashboards, models, and decision architectures that shape what they see before they decide.

For institutions and societies, Human Irreducibility names the capacities that must not be replaced by machines, while HTTPS translates that boundary into AI governance through Humanity, Truth, Transparency, Protection, and Security.

Taken together, these implications make FILE usable as a diagnostic lens for examining whether organizations shaped by AI remain governed by human responsibility.

15.5 Three illustrative managerial vignettes

Three brief vignettes show how FILE can operate as a management diagnostic.

First, consider an organization that introduces an AI-supported performance management system. The system improves visibility and standardizes evaluation, but employees stop contesting ratings because they believe objections will not be heard. Managers trust the dashboard more than the conversation. FILE diagnoses more than a technical adoption problem. Emotional Intelligence asks whether fear and trust are being governed. Political Intelligence asks how the system redistributes power. LENS asks who governs the categories and metrics. RC asks whether the Relational Commons is being depleted.

Second, consider a strategy team that uses generative AI to produce a market analysis. The output is fluent, fast, and apparently comprehensive. Yet no one verifies the sources, tests the assumptions, checks omissions, or asks which categories shaped the analysis. HACK identifies the missing methodological discipline. AI contributed outputs, but human judgment did not fully govern the process. LENS adds the deeper question: did the team govern the decision frame, or did the AI system shape what counted as strategic reality before discussion began?

Third, consider a business school that adds AI literacy modules to an existing MBA curriculum. Students learn tools, prompts, analytics, and automation use cases. The initiative is useful, but incomplete. MAIL and MLT diagnose the gap: AI is not only a technical competence, but a structural condition of management and leadership. The educational question is not only whether students know how to use AI. It is whether they are formed to lead responsibly when AI changes work, knowledge, and decision-making.

These vignettes show why FILE belongs inside management theory. AI adoption is never only technical. It is emotional, cultural, political, adaptive, knowledge-related, philosophical, educational, and governance-related at the same time.

16. Research Agenda Opened by the FILE Odyssey

16.1 Construct clarification and theoretical comparison

The first research agenda concerns construct clarification. Each contribution of the FILE Odyssey requires further conceptual development, comparison, and operationalization. FILE requires deeper clarification of Augmented Intelligence, Emotional Intelligence, Cultural Intelligence, Political Intelligence, and Adaptive Intelligence, including their boundaries, overlaps, tensions, and possible indicators. The 7E Cascade requires examination as a developmental model of leadership evolution. RC requires refinement as a construct connecting trust, dignity, psychological safety, meaning, voice, legitimacy, ability to challenge, care, and belonging.

MAIL and MLT require conceptual positioning within management education, leadership education, responsible management education, and technology education. HACK requires further development as a methodology for responsible AI-assisted knowledge production. LENS requires theoretical refinement as a judgment model for the conditions that shape human judgment. Human Irreducibility requires philosophical development and future empirical exploration as a research-generating proposition. HTTPS requires further study as AI governance, including its relation to existing regulations, standards, and institutional governance regimes.

This agenda should compare FILE constructs with adjacent theories. FILE can be compared with emotional intelligence, cultural intelligence, political skill, adaptive leadership, dynamic capabilities, digital leadership, human-centered artificial intelligence, responsible leadership, sociotechnical theory, and AI governance. The purpose is not to claim that FILE replaces these literatures. The purpose is to clarify what FILE contributes to them and what it receives from them.

16.2 Leadership, organizations, and work transformed by AI

The second research agenda concerns leadership and organizations. Future research can examine whether FILE helps leaders diagnose the human consequences of AI adoption. It can study how leaders use the five intelligences when AI mediates recruitment, performance management, strategic planning, decision support, communication, coordination, customer interaction, knowledge work, and institutional governance.

This research can ask direct questions. Do leaders with stronger Augmented Intelligence govern AI systems more responsibly? Does Emotional Intelligence reduce fear, disengagement, or loss of trust during automation? Does Cultural Intelligence improve the interpretation of AI systems across different languages, institutions, and cultural contexts? Does Political Intelligence help leaders identify how data, metrics, and algorithmic systems redistribute power? Does Adaptive Intelligence help organizations learn from AI deployment rather than normalize its unintended consequences?

Research can also examine the 7E Cascade in practice. Organizations may understand the need for evolution but fail at execution. They may pursue effectiveness without excellence. They

may build AI ecosystems without empowerment. They may teach leadership language without embodiment. The 7E Cascade gives researchers a way to study where leadership development breaks down under conditions created by AI.

16.3 Relational Commons and organizational diagnosis

The third research agenda concerns the Relational Commons. RC can be developed as a diagnostic lens for studying whether AI deployment protects or depletes the human conditions of responsible leadership.

Future research can examine trust, dignity, psychological safety, meaning, voice, ability to challenge, care, legitimacy, responsibility, and belonging in organizations transformed by AI. It can study how automated evaluation affects psychological safety, how predictive ranking affects dignity, how algorithmic management affects voice, how surveillance affects trust, and how work transformed by AI affects meaning.

RC also opens a practical research question: what leadership practices protect the Relational Commons? Possible areas include mechanisms for challenging decisions, transparent communication, participatory design, protection against retaliation, spaces for human deliberation, accountable override procedures, and responsible rules for AI use. The research question is not only whether AI improves performance. It is whether AI deployment preserves the human ecosystem in which responsible leadership remains possible.

16.4 Management education and the MAIL/MLT agenda

The fourth research agenda concerns management education. MAIL and MLT open a research agenda for business schools, universities, executive education, and leadership development institutions.

Research can examine how students and leaders learn to integrate management, AI, and leadership. It can examine whether curricula built around FILE improve students' capacity to govern AI systems and robots responsibly. It can compare traditional business education with programs that integrate management, AI, leadership, ethics, judgment, cultural interpretation, political analysis, and adaptive learning.

This agenda can also develop learning outcomes for the five intelligences. Augmented Intelligence can be assessed through responsible use and governance of AI systems and robots. Emotional Intelligence can be assessed through trust, care, motivation, and psychological safety in work transformed by AI. Cultural Intelligence can be assessed through cross-cultural interpretation of AI systems. Political Intelligence can be assessed through power, legitimacy, ability to challenge, and stakeholder analysis. Adaptive Intelligence can be assessed through learning, revision, experimentation, and responsible action under uncertainty.

MAIL/MLT also invites institutional research. Business schools can examine whether AI should remain an elective topic, or whether it should become part of the core formation of managers and leaders. The FILE Odyssey argues for the second path. Management, AI, and Leadership must be taught together because organizations now operate where all three meet.

16.5 HACK and responsible knowledge production

The fifth research agenda concerns HACK. AI-assisted knowledge production is already reshaping research, teaching, writing, synthesis, translation, peer review, and professional analysis. The question is no longer whether AI will enter intellectual work. It already has. The question is how human beings should govern that work responsibly.

Research on HACK can examine the methodology across contexts. It can study how researchers use AI for brainstorming, literature mapping, drafting, critique, translation, synthesis, and revision. It can examine when AI strengthens inquiry and when it weakens reading, source discipline, interpretation, and originality. It can study the effects of AI assistance on authorial voice, scholarly responsibility, citation accuracy, argument quality, and intellectual formation.

HACK also opens a research agenda on verification. Verification is not a bureaucratic step. It is part of human judgment. Research can examine how scholars, students, leaders, and institutions verify AI-assisted outputs. It can develop protocols for source checking, claim testing, citation discipline, disclosure, and final accountability. HACK makes one principle nondelegable: the burden of judgment cannot be outsourced because accountability cannot be outsourced.

16.6 LENS, Human Irreducibility, and HTTPS

The sixth research agenda concerns judgment, philosophy, and governance. LENS, Human Irreducibility, and HTTPS together open a research agenda on how human judgment remains responsible when AI shapes the conditions of perception, decision, and governance.

LENS invites research on the governance of judgment. Researchers can study how leaders govern the systems, metrics, classifications, outputs, and decision conditions that shape what they see, measure, compare, value, and decide. They can examine algorithmic capture, capture of knowledge and judgment, dashboard dependence, metric fixation, explainability, ability to challenge, override procedures, and accountability. The central question is whether leaders still govern judgment, or whether they only approve decisions inside frames they did not govern.

Human Irreducibility invites research on the protection of the ten irreducible capacities: judgment, discernment, conscience, moral agency, dignity, meaning, responsibility, care, trust, and love. Future research can examine how leaders defend these capacities under algorithmic pressure. It can study whether organizations transformed by AI preserve responsibility, dignity, care, meaning, trust, and moral agency, or whether they normalize functional substitution. Human Irreducibility is not yet an empirically validated measurement scale. It is a philosophical boundary and research-generating proposition.

HTTPS invites research on human-centered AI governance. Researchers can study how Humanity, Truth, Transparency, Protection, and Security can guide organizational policies, public institutions, educational settings, and cross-sector governance. HTTPS does not replace existing regulations, standards, or institutional governance regimes. It gives FILE a compact

expression of AI governance that can be compared with and applied alongside existing governance instruments.

A further research direction concerns the interaction between HTTPS and existing or emerging AI governance regimes. Future work should examine how the Five Immutable Laws of AI and Robotics relate to legal, regulatory, sectoral, and institutional frameworks, including the EU AI Act (European Union 2024), professional standards, public-sector oversight, education governance, and sector-specific regulators. The purpose is not to replace these regimes, but to clarify how FILE’s human-centered AI governance can help leaders interpret, apply, and critique them through the lenses of Humanity, Truth, Transparency, Protection, and Security.

16.7 Cross-cultural, institutional, and civilizational research

The final research agenda is cross-cultural and institutional. FILE must be examined across cultures, languages, sectors, organizations, and governance traditions. AI is global, but its meanings, risks, and uses are not experienced uniformly. Leadership, dignity, care, authority, transparency, trust, education, and governance are interpreted through different cultural, institutional, legal, and historical contexts.

Future research can examine how FILE travels across regions and institutions. Does the five-intelligence model resonate in different cultural settings? How do emotional, cultural, political, and adaptive intelligence operate in different societies? How do concepts such as Human Irreducibility, RC, LENS, HACK, and HTTPS require adaptation across legal traditions, educational systems, organizational cultures, and public institutions?

This agenda is especially important because human-centered artificial intelligence must not become a culturally narrow formula. Cross-cultural research can examine, challenge, translate, and refine FILE across contexts. Its claim is universal in aspiration, but its development must remain open to plural forms of knowledge, leadership, culture, and institutional life.

16.8 Operationalizing the FILE Odyssey

The FILE Odyssey also supports operationalization. Its concepts can become researchable objects examined in organizations, classrooms, leadership programs, AI adoption projects, governance settings, and public-policy contexts.

The following table does not present final validated measurement scales. It identifies research objects, possible indicators, and settings for descriptive empirical research.

Table 2. Operationalizing the FILE Odyssey

Contribution	Research object	Possible indicators	Settings
Leadership = Intelligence + Character + Judgment	Integration of Intelligence, Character, and Judgment	FILE / SQL / LENS relation; human responsibility; decision quality	Leadership development; executive teams; management education

Contribution	Research object	Possible indicators	Settings
FILE: The Five Intelligences of Leadership Evolution	Five-intelligence integration	Human oversight; trust; cultural interpretation; power awareness; adaptive learning	Executive teams; AI adoption; MBA courses
FILE, FILE ³ , FILE ⁵ , FILE ⁷ , and the 7E Cascade	Development of leadership intelligence	Awareness; effectiveness; ecosystem thinking; empowerment; execution; embodiment	Leadership development; executive education
SQL: The Six Qualities of Leadership	Character of the embodied leader	Vision; Intuition; Taste; Courage; Discipline; Integrity	Leadership development; executive education; organizations
LENS: Leadership, Ethics, Nondelegation, and Sensemaking	Second-order governance of human judgment	Ability to challenge metrics; dashboard dependence; decision-frame review	Executive systems; risk committees; AI governance boards
RC: The Relational Commons	Protection of the Relational Commons	Trust; dignity; voice; care; meaning; belonging	HR analytics; automated evaluation; hybrid work
MAIL: Management, AI, and Leadership / MLT: Management, Leadership, and Technology	Formation for management, AI, and leadership	Curriculum integration; student judgment; responsible AI use	Business schools; MBAs; corporate universities
HACK: Human-AI Co-created Knowledge	Responsible AI-assisted knowledge production	Verification; claim testing; authorial voice; accountability	Research; consulting; student work; policy analysis
Human Irreducibility	Protection of non-substitutable human capacities	Responsibility; dignity; care; meaning; trust; moral agency	Automation-intensive organizations; education; public services

Contribution	Research object	Possible indicators	Settings
The Five Immutable Laws of AI and Robotics: HTTPS: Humanity, Truth, Transparency, Protection, and Security	Human-centered AI governance	Humanity; Truth; Transparency; Protection; Security	AI policies; governance committees; institutional oversight
BASIC / RULES	Human-centered public-policy conditions	Basic Income; AI Access; Skills; Inclusion; Care	Public policy; education; labor transition; social institutions
Inalienable Leadership	Final human authority and responsibility	Authority; responsibility; non-surrender; civilizational direction	Boards; institutions; public governance; AI and robotics governance

This operationalization pathway shows that the FILE Odyssey is not limited to conceptual synthesis. It can generate empirical questions, diagnostic instruments, teaching designs, organizational audits, governance protocols, and public-policy inquiries without making the legitimacy of its theoretical contributions dependent on a future validation program.

These operationalization ideas and research questions identify possible empirical work; the paper remains a theoretical and research-generating contribution.

FILE includes six falsifiable propositions that define explicit conditions under which its theoretical claims can be examined and bounded.

The Multi-Intelligence Claim proposes that responsible leadership in the age of artificial intelligence and robotics requires the integration of Augmented Intelligence, Emotional Intelligence, Cultural Intelligence, Political Intelligence, and Adaptive Intelligence, not reliance on one dominant capacity. The Augmented Intelligence Claim proposes that leaders work best with AI when they preserve human judgment, verification, responsibility, and critical distance. The Persistence of Human Intelligences Claim proposes that Emotional, Cultural, Political, and Adaptive Intelligence remain necessary even when AI can simulate emotionally appropriate, culturally adapted, strategically persuasive, or adaptively responsive outputs. The Interdependence Claim proposes that the five intelligences interact rather than operate as isolated traits. The Maturity and Development Claim proposes that these intelligences can develop through education, practice, feedback, and the 7E Cascade. The Organizational and Educational Scope Claim places inquiry at the level of individual leaders, teams, organizations, institutions, and management education.

The empirical agenda includes scale development for the five intelligences and RC, confirmatory factor analysis of the nested architecture, longitudinal and field studies linking

FILE profiles to outcomes such as trust, innovation, psychological safety, burnout, and legitimacy, and cross-cultural research. The research agenda also examines how these constructs operate across organizational, cultural, and institutional contexts, including their boundaries, interactions, practical applications, and possible limitations. This work can deepen understanding of FILE while contributing to broader research on leadership and decision-making in AI-mediated organizations.

Several research questions follow from this pathway. Do leadership teams that use the five FILE intelligences identify relational, political, and adaptive risks earlier than teams that use only technical or compliance-based AI adoption models? Do organizations that monitor the Relational Commons preserve trust, dignity, and voice more effectively during algorithmic management? Do MAIL- or MLT-inspired curricula improve students' ability to integrate AI use with leadership judgment? Do HACK-based protocols improve the traceability, verification, and accountability of AI-assisted knowledge work? Do LENS routines improve the ability to challenge dashboards, classifications, and decision frames? These questions illustrate how FILE can move from conceptual formulation to empirical research.

17. Conclusion: Humans + Machines, under Human Leadership

AI is changing the conditions of management, leadership, education, knowledge production, judgment, and governance. It can automate tasks, generate outputs, classify reality, support decisions, optimize processes, simulate human expression, and scale organizational action. It can expand human capability. It can also weaken the human conditions that make responsible leadership possible.

This paper has argued that the response cannot be only technical. It must be managerial, educational, knowledge-related, philosophical, and governance-oriented at the same time. FILE enters this transformation as a management theory and diagnostic framework for responsible leadership and decision-making in the age of artificial intelligence and robotics. It extends that theory into a transdisciplinary research program: the FILE Odyssey.

The twelve contributions developed in this paper define that program. Leadership = Intelligence + Character + Judgment defines human leadership through Intelligence, Character, and Judgment. FILE names the five intelligences of leadership evolution. The 7E Cascade explains their development. SQL names Character through Vision, Intuition, Taste, Courage, Discipline, and Integrity. LENS: Leadership, Ethics, Nondelegation, and Sensemaking protects human judgment. RC names the Relational Commons that leadership must protect, cultivate, and renew. MAIL and MLT redefine management education for the age of AI and robotics. HACK establishes a methodology for responsible AI-assisted knowledge production. Human Irreducibility defines the philosophical foundation and the line that AI must never cross. The Five Immutable Laws of AI and Robotics: HTTPS: Humanity, Truth, Transparency, Protection, and Security defines human-centered AI governance. BASIC and RULES provide the human-centered public-policy contribution. Inalienable Leadership is the twelfth and

culminating original contribution and the overarching metatheory of leadership in the age of artificial intelligence and robotics.

The management-theory contribution is therefore to reframe AI not as a tool added to leadership, but as a structural condition that changes how leadership, judgment, education, knowledge production, relational life, and governance must be understood.

The final claim follows from this logic. AI can assist leadership, knowledge, language, judgment, work, education, and governance. But **assistance is not equivalence**. Machines can support human activity. They cannot carry responsibility, meaning, conscience, care, trust, or love.

This creates five practical challenges. Leaders must lead people and govern AI systems and robots without surrendering judgment. Educators must form students who integrate Management, AI, and Leadership. Scholars must use AI without abandoning verification, interpretation, voice, and accountability. Organizations must protect the Relational Commons under speed, automation, and measurement. Societies must govern artificial intelligence and robotics through humanity, truth, transparency, protection, and security.

This paper presents FILE as a theoretical contribution and a research agenda. The research program includes teaching, comparison, organizational study, public-policy inquiry, and empirical research.

The age of AI will not be judged only by what machines can do. It will be judged by whether human beings still govern what machines make possible. This is the promise and demand of **the FILE School of Thought: Humans + Machines, under human leadership**.

References

- Abrams, Dominic, and Michael A. Hogg. 2004. "Metatheory: Lessons from Social Identity Research." *Personality and Social Psychology Review* 8, no. 2: 98–106.
- Acemoglu, Daron, and Pascual Restrepo. 2019. "Automation and New Tasks: How Technology Displaces and Reinstates Labor." *Journal of Economic Perspectives* 33, no. 2: 3–30.
- Acemoglu, Daron, and Pascual Restrepo. 2020. "Robots and Jobs: Evidence from US Labor Markets." *Journal of Political Economy* 128, no. 6: 2188–2244.
- Ananny, Mike, and Kate Crawford. 2018. "Seeing without Knowing: Limitations of the Transparency Ideal and Its Application to Algorithmic Accountability." *New Media & Society* 20, no. 3: 973–989.
- Ang, Soon, Linn Van Dyne, Christine Koh, K. Yee Ng, Klaus J. Templer, Cheryl Tay, and N. Anand Chandrasekar. 2007. "Cultural Intelligence: Its Measurement and Effects on Cultural Judgment and Decision Making, Cultural Adaptation and Task Performance." *Management and Organization Review* 3, no. 3: 335–371.
- Argyris, Chris, and Donald A. Schön. 1978. *Organizational Learning: A Theory of Action Perspective*. Reading, MA: Addison-Wesley.

- Autor, David H., Frank Levy, and Richard J. Murnane. 2003. "The Skill Content of Recent Technological Change: An Empirical Exploration." *Quarterly Journal of Economics* 118, no. 4: 1279–1333.
- Autor, David H. 2015. "Why Are There Still So Many Jobs? The History and Future of Workplace Automation." *Journal of Economic Perspectives* 29, no. 3: 3–30.
- Avolio, Bruce J., and William L. Gardner. 2005. "Authentic Leadership Development: Getting to the Root of Positive Forms of Leadership." *The Leadership Quarterly* 16, no. 3: 315–338.
- Bainbridge, Lisanne. 1983. "Ironies of Automation." *Automatica* 19, no. 6: 775–779.
- Bandura, Albert. 1991. "Social Cognitive Theory of Self-Regulation." *Organizational Behavior and Human Decision Processes* 50, no. 2: 248–287.
- Bass, Bernard M. 1985. *Leadership and Performance Beyond Expectations*. New York: Free Press.
- Bass, Bernard M., and Paul Steidlmeier. 1999. "Ethics, Character, and Authentic Transformational Leadership Behavior." *The Leadership Quarterly* 10, no. 2: 181–217.
- Bennis, Warren G., and James O'Toole. 2005. "How Business Schools Lost Their Way." *Harvard Business Review* 83, no. 5: 96–104, 154.
- Bouilloud, Jean-Philippe, Ghislain Deslandes, and Guillaume Mercier. 2019. "The Leader as Chief Truth Officer: The Ethical Responsibility of 'Managing the Truth' in Organizations." *Journal of Business Ethics* 157, no. 1: 1–13.
- Bovens, Mark. 2007. "Analysing and Assessing Accountability: A Conceptual Framework." *European Law Journal* 13, no. 4: 447–468.
- Brown, Michael E., Linda K. Treviño, and David A. Harrison. 2005. "Ethical Leadership: A Social Learning Perspective for Construct Development and Testing." *Organizational Behavior and Human Decision Processes* 97, no. 2: 117–134.
- Brundage, Miles, Shahar Avin, Jack Clark, Helen Toner, Peter Eckersley, Ben Garfinkel, Allan Dafoe, Paul Scharre, Thomas Zeitzoff, Bobby Filar, Hyrum Anderson, Heather Roff, Gregory C. Allen, Jacob Steinhardt, Carrick Flynn, Seán Ó hÉigeartaigh, Simon Beard, Haydn Belfield, Sebastian Farquhar, Clare Lyle, Rebecca Crotoft, Owain Evans, Michael Page, Joanna Bryson, Roman Yampolskiy, and Dario Amodei. 2018. *The Malicious Use of Artificial Intelligence: Forecasting, Prevention, and Mitigation*. arXiv preprint arXiv:1802.07228.
- Brynjolfsson, Erik, Danielle Li, and Lindsey R. Raymond. 2023. "Generative AI at Work." NBER Working Paper No. 31161. Cambridge, MA: National Bureau of Economic Research.
- Burns, James MacGregor. 1978. *Leadership*. New York: Harper & Row.
- Chesney, Robert, and Danielle Keats Citron. 2019. "Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security." *California Law Review* 107, no. 6: 1753–1820.

- Coleman, James S. 1988. "Social Capital in the Creation of Human Capital." *American Journal of Sociology* 94: S95-S120.
- Crossan, Mary M., Alyson Byrne, Gerard H. Seijts, Mark Reno, Lucas Monzani, and Jeffrey Gandz. 2017. "Toward a Framework of Leader Character in Organizations." *Journal of Management Studies* 54, no. 7: 986–1018.
- Dane, Erik, and Michael G. Pratt. 2007. "Exploring Intuition and Its Role in Managerial Decision Making." *Academy of Management Review* 32, no. 1: 33–54.
- Day, David V. 2000. "Leadership Development: A Review in Context." *The Leadership Quarterly* 11, no. 4: 581–613.
- Deming, David J. 2017. "The Growing Importance of Social Skills in the Labor Market." *Quarterly Journal of Economics* 132, no. 4: 1593–1640.
- Detert, James R., and Evan A. Bruno. 2017. "Workplace Courage: Review, Synthesis, and Future Agenda for a Complex Construct." *Academy of Management Annals* 11, no. 2: 593–639.
- Doshi-Velez, Finale, and Been Kim. 2017. "Towards a Rigorous Science of Interpretable Machine Learning." arXiv preprint arXiv:1702.08608.
- Earley, P. Christopher, and Soon Ang. 2003. *Cultural Intelligence: Individual Interactions across Cultures*. Stanford, CA: Stanford University Press.
- Edmondson, Amy C. 1999. "Psychological Safety and Learning Behavior in Work Teams." *Administrative Science Quarterly* 44, no. 2: 350–383.
- Elish, Madeleine Clare. 2019. "Moral Crumple Zones: Cautionary Tales in Human-Robot Interaction." *Engaging Science, Technology, and Society* 5: 40–60.
- Eloundou, Tyna, Sam Manning, Pamela Mishkin, and Daniel Rock. 2023. "GPTs Are GPTs: An Early Look at the Labor Market Impact Potential of Large Language Models." arXiv preprint arXiv:2303.10130.
- Engelbart, Douglas C. 1962. *Augmenting Human Intellect: A Conceptual Framework*. Summary Report AFOSR-3223. Menlo Park, CA: Stanford Research Institute.
- Eubanks, Virginia. 2018. *Automating Inequality: How High-Tech Tools Profile, Police, and Punish the Poor*. New York: St. Martin's Press.
- European Union. 2024. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence. *Official Journal of the European Union*. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>.
- Faraj, Samer, Stella Pachidi, and Karla Sayegh. 2018. "Working and Organizing in the Age of the Learning Algorithm." *Information and Organization* 28, no. 1: 62–70.
- Ferris, Gerald R., Darren C. Treadway, Robert W. Kolodinsky, Wayne A. Hochwarter, Charles J. Kacmar, Ceasar Douglas, and Dwight D. Frink. 2005. "Development and Validation of the Political Skill Inventory." *Journal of Management* 31, no. 1: 126–152.

- Floridi, Luciano, and Josh Cowls. 2019. "A Unified Framework of Five Principles for AI in Society." *Harvard Data Science Review* 1, no. 1.
- Folbre, Nancy. 2001. *The Invisible Heart: Economics and Family Values*. New York: New Press.
- Fricker, Miranda. 2007. *Epistemic Injustice: Power and the Ethics of Knowing*. Oxford: Oxford University Press.
- Ghoshal, Sumantra. 2005. "Bad Management Theories Are Destroying Good Management Practices." *Academy of Management Learning & Education* 4, no. 1: 75–91.
- Gilligan, Carol. 1982. *In a Different Voice: Psychological Theory and Women's Development*. Cambridge, MA: Harvard University Press.
- Goldman, Alvin I. 1999. *Knowledge in a Social World*. Oxford: Oxford University Press.
- Goleman, Daniel. 1995. *Emotional Intelligence: Why It Can Matter More Than IQ*. New York: Bantam Books.
- Greenleaf, Robert K. 1977. *Servant Leadership: A Journey into the Nature of Legitimate Power and Greatness*. New York: Paulist Press.
- Gronn, Peter. 2002. "Distributed Leadership as a Unit of Analysis." *The Leadership Quarterly* 13, no. 4: 423–451.
- Hansen, Hans, Arja Ropo, and Erika Sauer. 2007. "Aesthetic Leadership." *The Leadership Quarterly* 18, no. 6: 544–560.
- Haraway, Donna. 1988. "Situated Knowledges: The Science Question in Feminism and the Privilege of Partial Perspective." *Feminist Studies* 14, no. 3: 575–599.
- Hargittai, Eszter. 2002. "Second-Level Digital Divide: Differences in People's Online Skills." *First Monday* 7, no. 4.
- Heifetz, Ronald A. 1994. *Leadership Without Easy Answers*. Cambridge, MA: Belknap Press of Harvard University Press.
- Heifetz, Ronald A., Alexander Grashow, and Marty Linsky. 2009. *The Practice of Adaptive Leadership: Tools and Tactics for Changing Your Organization and the World*. Boston: Harvard Business Press.
- Honneth, Axel. 1995. *The Struggle for Recognition: The Moral Grammar of Social Conflicts*. Translated by Joel Anderson. Cambridge, MA: MIT Press.
- Hutchins, Edwin. 1995. *Cognition in the Wild*. Cambridge, MA: MIT Press.
- Iansiti, Marco, and Roy Levien. 2004. *The Keystone Advantage: What the New Dynamics of Business Ecosystems Mean for Strategy, Innovation, and Sustainability*. Boston: Harvard Business School Press.
- Jarrahi, Mohammad Hossein. 2018. "Artificial Intelligence and the Future of Work: Human-AI Symbiosis in Organizational Decision Making." *Business Horizons* 61, no. 4: 577–586.

- Jobin, Anna, Marcello Ienca, and Effy Vayena. 2019. "The Global Landscape of AI Ethics Guidelines." *Nature Machine Intelligence* 1: 389–399.
- Jonas, Hans. 1984. *The Imperative of Responsibility: In Search of an Ethics for the Technological Age*. Translated by Hans Jonas with the collaboration of David Herr. Chicago: University of Chicago Press.
- Kahneman, Daniel, and Amos Tversky. 1979. "Prospect Theory: An Analysis of Decision under Risk." *Econometrica* 47, no. 2: 263–291.
- Kant, Immanuel. 1785. *Grundlegung zur Metaphysik der Sitten*. Riga: Johann Friedrich Hartknoch.
- Katz, Robert L. 1955. "Skills of an Effective Administrator." *Harvard Business Review* 33, no. 1: 33–42.
- Kegan, Robert. 1982. *The Evolving Self: Problem and Process in Human Development*. Cambridge, MA: Harvard University Press.
- Kellogg, Katherine C., Melissa A. Valentine, and Angèle Christin. 2020. "Algorithms at Work: The New Contested Terrain of Control." *Academy of Management Annals* 14, no. 1: 366–410.
- Kolb, David A. 1984. *Experiential Learning: Experience as the Source of Learning and Development*. Englewood Cliffs, NJ: Prentice Hall.
- Licklider, J. C. R. 1960. "Man-Computer Symbiosis." *IRE Transactions on Human Factors in Electronics* HFE-1, no. 1: 4–11.
- Lévy, Pierre. 1997. *Collective Intelligence: Mankind's Emerging World in Cyberspace*. Translated by Robert Bononno. New York: Plenum Trade.
- Malone, Thomas W., and Michael S. Bernstein, eds. 2015. *Handbook of Collective Intelligence*. Cambridge, MA: MIT Press.
- Malone, Thomas W. 2018. *Superminds: The Surprising Power of People and Computers Thinking Together*. New York: Little, Brown and Company.
- Manz, Charles C. 1986. "Self-Leadership: Toward an Expanded Theory of Self-Influence Processes in Organizations." *Academy of Management Review* 11, no. 3: 585–600.
- March, James G. 1994. *A Primer on Decision Making: How Decisions Happen*. New York: Free Press.
- Matthias, Andreas. 2004. "The Responsibility Gap: Ascribing Responsibility for the Actions of Learning Automata." *Ethics and Information Technology* 6: 175–183.
- Mayer, Roger C., James H. Davis, and F. David Schoorman. 1995. "An Integrative Model of Organizational Trust." *Academy of Management Review* 20, no. 3: 709–734.
- Mayer, John D., Peter Salovey, and David R. Caruso. 2004. "Emotional Intelligence: Theory, Findings, and Implications." *Psychological Inquiry* 15, no. 3: 197–215.

- Merleau-Ponty, Maurice. 1945. *Phénoménologie de la perception*. Paris: Gallimard.
- Mezirow, Jack. 1991. *Transformative Dimensions of Adult Learning*. San Francisco: Jossey-Bass.
- Mindell, David A. 2015. *Our Robots, Ourselves: Robotics and the Myths of Autonomy*. New York: Viking.
- Mintzberg, Henry. 1973. *The Nature of Managerial Work*. New York: Harper & Row.
- Mintzberg, Henry. 2004. *Managers Not MBAs: A Hard Look at the Soft Practice of Managing and Management Development*. San Francisco: Berrett-Koehler.
- MIT Center for Collective Intelligence. n.d. "Thomas W. Malone." Massachusetts Institute of Technology. Accessed July 6, 2026. <https://cci.mit.edu/malone/>.
- Mittelstadt, Brent. 2019. "Principles Alone Cannot Guarantee Ethical AI." *Nature Machine Intelligence* 1: 501–507.
- Moore, James F. 1996. *The Death of Competition: Leadership and Strategy in the Age of Business Ecosystems*. New York: HarperBusiness.
- Mumford, Michael D., Stephen J. Zaccaro, Francis D. Harding, T. Owen Jacobs, and Edwin A. Fleishman. 2000. "Leadership Skills for a Changing World: Solving Complex Social Problems." *The Leadership Quarterly* 11, no. 1: 11–35.
- Noddings, Nel. 1984. *Caring: A Feminine Approach to Ethics and Moral Education*. Berkeley: University of California Press.
- Nonaka, Ikujiro. 1994. "A Dynamic Theory of Organizational Knowledge Creation." *Organization Science* 5, no. 1: 14–37.
- Nussbaum, Martha C. 2011. *Creating Capabilities: The Human Development Approach*. Cambridge, MA: Harvard University Press.
- Orlikowski, Wanda J. 2000. "Using Technology and Constituting Structures: A Practice Lens for Studying Technology in Organizations." *Organization Science* 11, no. 4: 404–428.
- Ostrom, Elinor. 1990. *Governing the Commons: The Evolution of Institutions for Collective Action*. Cambridge: Cambridge University Press.
- O'Neil, Cathy. 2016. *Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy*. New York: Crown.
- Palanski, Michael E., and Francis J. Yammarino. 2007. "Integrity and Leadership: Clearing the Conceptual Confusion." *European Management Journal* 25, no. 3: 171–184.
- Pasquale, Frank. 2015. *The Black Box Society: The Secret Algorithms That Control Money and Information*. Cambridge, MA: Harvard University Press.
- Pasquale, Frank. 2020. *New Laws of Robotics: Defending Human Expertise in the Age of AI*. Cambridge, MA: Belknap Press of Harvard University Press.
- Pfeffer, Jeffrey. 1992. *Managing with Power: Politics and Influence in Organizations*. Boston: Harvard Business School Press.

- Pfeffer, Jeffrey, and Christina T. Fong. 2002. "The End of Business Schools? Less Success Than Meets the Eye." *Academy of Management Learning & Education* 1, no. 1: 78–95.
- Pfeffer, Jeffrey. 2010. *Power: Why Some People Have It and Others Don't*. New York: HarperBusiness.
- Polanyi, Michael. 1966. *The Tacit Dimension*. Garden City, NY: Doubleday.
- Power, Michael. 1997. *The Audit Society: Rituals of Verification*. Oxford: Oxford University Press.
- Putnam, Robert D. 2000. *Bowling Alone: The Collapse and Revival of American Community*. New York: Simon & Schuster.
- Rahwan, Iyad, Manuel Cebrian, Nick Obradovich, Josh Bongard, Jean-François Bonnefon, Cynthia Breazeal, Jacob W. Crandall, Nicholas A. Christakis, Iain D. Couzin, Matthew O. Jackson, Nicholas R. Jennings, Ece Kamar, Isabel M. Kloumann, Hugo Larochelle, David Lazer, Richard McElreath, Alan Mislove, David C. Parkes, Alex 'Sandy' Pentland, Margaret E. Roberts, Azim Shariff, Joshua B. Tenenbaum, and Michael Wellman. 2019. "Machine Behaviour." *Nature* 568, no. 7753: 477–486.
- Raisch, Sebastian, and Sebastian Krakowski. 2021. "Artificial Intelligence and Management: The Automation-Augmentation Paradox." *Academy of Management Review* 46, no. 1: 192–210.
- Raji, Inioluwa Deborah, Andrew Smart, Rebecca N. White, Margaret Mitchell, Timnit Gebru, Ben Hutchinson, Jamila Smith-Loud, Daniel Theron, and Parker Barnes. 2020. "Closing the AI Accountability Gap: Defining an End-to-End Framework for Internal Algorithmic Auditing." *Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency* 33–44. New York: Association for Computing Machinery.
- Ricoeur, Paul. 1992. *Oneself as Another*. Translated by Kathleen Blamey. Chicago: University of Chicago Press.
- Rousseau, Denise M., Sim B. Sitkin, Ronald S. Burt, and Colin Camerer. 1998. "Not So Different After All: A Cross-Discipline View of Trust." *Academy of Management Review* 23, no. 3: 393–404.
- Russell, Stuart. 2019. *Human Compatible: Artificial Intelligence and the Problem of Control*. New York: Viking.
- Salovey, Peter, and John D. Mayer. 1990. "Emotional Intelligence." *Imagination, Cognition and Personality* 9, no. 3: 185–211.
- Santoni de Sio, Filippo, and Jeroen van den Hoven. 2018. "Meaningful Human Control over Autonomous Systems: A Philosophical Account." *Frontiers in Robotics and AI* 5: article 15.
- Schwartz, Barry, and Kenneth Sharpe. 2010. *Practical Wisdom: The Right Way to Do the Right Thing*. New York: Riverhead Books.

- Schön, Donald A. 1983. *The Reflective Practitioner: How Professionals Think in Action*. New York: Basic Books.
- Selbst, Andrew D., Danah Boyd, Sorelle A. Friedler, Suresh Venkatasubramanian, and Janet Vertesi. 2019. "Fairness and Abstraction in Sociotechnical Systems." *Proceedings of the Conference on Fairness, Accountability, and Transparency* 59–68. New York: Association for Computing Machinery.
- Selznick, Philip. 1957. *Leadership in Administration: A Sociological Interpretation*. New York: Harper & Row.
- Sen, Amartya. 1999. *Development as Freedom*. New York: Alfred A. Knopf.
- Senge, Peter M. 1990. *The Fifth Discipline: The Art and Practice of the Learning Organization*. New York: Doubleday.
- Shneiderman, Ben. 2022. *Human-Centered AI*. Oxford: Oxford University Press.
- Shotter, John, and Haridimos Tsoukas. 2014. "In Search of Phronesis: Leadership and the Art of Judgment." *Academy of Management Learning & Education* 13, no. 2: 224–243.
- Shrestha, Yash Raj, Shiko M. Ben-Menahem, and Georg von Krogh. 2019. "Organizational Decision-Making Structures in the Age of Artificial Intelligence." *California Management Review* 61, no. 4: 66–83.
- Simon, Herbert A. 1947. *Administrative Behavior: A Study of Decision-Making Processes in Administrative Organization*. New York: Macmillan.
- Stanford Institute for Human-Centered Artificial Intelligence. n.d. "About." Stanford University. Accessed July 6, 2026. <https://hai.stanford.edu/about>.
- Sternberg, Robert J. 2003. "WICS: A Model of Leadership in Organizations." *Academy of Management Learning & Education* 2, no. 4: 386–401.
- Strati, Antonio. 1992. "Aesthetic Understanding of Organizational Life." *Academy of Management Review* 17, no. 3: 568–581.
- Suchman, Mark C. 1995. "Managing Legitimacy: Strategic and Institutional Approaches." *Academy of Management Review* 20, no. 3: 571–610.
- Surowiecki, James. 2004. *The Wisdom of Crowds: Why the Many Are Smarter Than the Few and How Collective Wisdom Shapes Business, Economies, Societies, and Nations*. New York: Doubleday.
- Taylor, Charles. 1992. "The Politics of Recognition." In *Multiculturalism and 'The Politics of Recognition'*, edited by Amy Gutmann, 25–73. Princeton, NJ: Princeton University Press.
- Teece, David J., Gary Pisano, and Amy Shuen. 1997. "Dynamic Capabilities and Strategic Management." *Strategic Management Journal* 18, no. 7: 509–533.
- Teece, David J. 2007. "Explicating Dynamic Capabilities: The Nature and Microfoundations of (Sustainable) Enterprise Performance." *Strategic Management Journal* 28, no. 13: 1319–1350.

- Trist, Eric L., and Kenneth W. Bamforth. 1951. "Some Social and Psychological Consequences of the Longwall Method of Coal-Getting: An Examination of the Psychological Situation and Defences of a Work Group in Relation to the Social Structure and Technological Content of the Work System." *Human Relations* 4, no. 1: 3–38.
- Tronto, Joan C. 1993. *Moral Boundaries: A Political Argument for an Ethic of Care*. New York: Routledge.
- Tronto, Joan C. 2013. *Caring Democracy: Markets, Equality, and Justice*. New York: New York University Press.
- Tversky, Amos, and Daniel Kahneman. 1974. "Judgment under Uncertainty: Heuristics and Biases." *Science* 185, no. 4157: 1124–1131.
- Uhl-Bien, Mary, Russ Marion, and Bill McKelvey. 2007. "Complexity Leadership Theory: Shifting Leadership from the Industrial Age to the Knowledge Era." *The Leadership Quarterly* 18, no. 4: 298–318.
- van Dijk, Jan A. G. M. 2005. *The Deepening Divide: Inequality in the Information Society*. Thousand Oaks, CA: Sage.
- Van Parijs, Philippe, and Yannick Vanderborght. 2017. *Basic Income: A Radical Proposal for a Free Society and a Sane Economy*. Cambridge, MA: Harvard University Press.
- Warschauer, Mark. 2003. *Technology and Social Inclusion: Rethinking the Digital Divide*. Cambridge, MA: MIT Press.
- Weber, Max. 1978. *Economy and Society: An Outline of Interpretive Sociology*. Edited by Guenther Roth and Claus Wittich. Berkeley: University of California Press. Original work published 1922.
- Weick, Karl E. 1995. *Sensemaking in Organizations*. Thousand Oaks, CA: Sage.
- Weizenbaum, Joseph. 1976. *Computer Power and Human Reason: From Judgment to Calculation*. San Francisco: W. H. Freeman.

FILE Corpus References

- Mariani, Guillaume. 2026. *FILE vs. Major Management Frameworks*. guillaumemariani.com. <https://guillaumemariani.com/file-vs-major-management-frameworks/>.
- Mariani, Guillaume. 2026. *FILE vs. Major Leadership Theories (V2)*. guillaumemariani.com. <https://guillaumemariani.com/file-vs-major-leadership-theories-v2/>.

AI Use Disclosure

Generative AI tools were used as assistive technologies for organization, drafting support, critical review, and language refinement. The author retained full responsibility for the theory, argument, interpretation, claims, references, and final text. No AI system is listed as an author or cited as an author.

About the Author

[Guillaume Mariani](#) is a French management scholar, entrepreneur, executive, and educator, and the author and creator of Inalienable Leadership and FILE: The Five Intelligences of Leadership Evolution, a management theory and diagnostic framework for responsible leadership and decision-making in the age of artificial intelligence and robotics, and of original contributions to management, leadership, and the social sciences, including a human-centered definition of leadership (Leadership = Intelligence + Character + Judgment, with Intelligence conceptualized through FILE, Character through SQL, and Judgment through LENS within the FILE Architecture); a human-centered framework (FILE = AI + EQ + CQ + PQ + AQ, where AI means Augmented Intelligence (“Mind”), EQ means Emotional Intelligence (“Heart”), CQ means Cultural Intelligence (“World”), PQ means Political Intelligence (“Society”), and AQ means Adaptive Intelligence (“Evolution”)); a human-centered developmental theory (FILE, FILE³, FILE⁵, FILE⁷, and the 7E Cascade); a human-centered ethos (SQL: The Six Qualities of Leadership: Vision, Intuition, Taste, Courage, Discipline, and Integrity); a human-centered judgment model (LENS: Leadership, Ethics, Nondelegation, and Sensemaking); a human-centered vision (RC: The Relational Commons); a human-centered education (MAIL: Management, AI, and Leadership, a curriculum supporting proposed MLT: Management, Leadership, and Technology degrees as alternatives to the traditional BBA and MBA); a human-centered methodology (HACK: Human-AI Co-created Knowledge); a human-centered philosophy (Human Irreducibility); a human-centered AI governance codex (The Five Immutable Laws of AI and Robotics: HTTPS: Humanity, Truth, Transparency, Protection, and Security); a human-centered public policy (BASIC: Basic Income, AI Access, Skills, Inclusion, and Care, and its French equivalent, RULES : Revenu Universel, Lutte contre la pauvreté et les inégalités, Éducation et accès à l’IA et Solidarité); and a human-centered civilization (Inalienable Leadership: the overarching metatheory of leadership in the age of artificial intelligence and robotics). Humans + Machines, under human leadership.

The FILE Odyssey is an ongoing research program and will be developed further in subsequent publications.

© Guillaume Mariani, 2026.