

LEARNING AI INSTITUTE

in partnership with

MetHer By Design, LLC

Research Insight Series • LAI-RIS-004

Developing Human Interpretation

Why AI Readiness Depends on Human Judgment, Sensemaking, and Responsibility

“Artificial intelligence can generate information. Only humans can determine its meaning.”

Dr. Matt Meador

Learning AI Institute | MetHer By Design, LLC

July 2026

Publication Information

Research Insight Series	LAI-RIS-004
Title	Developing Human Interpretation
Subtitle	Why AI Readiness Depends on Human Judgment, Sensemaking, and Responsibility
Author	Dr. Matt Meador
Founder	MetHer By Design, LLC
Executive Director	Learning AI Institute
Publication Date	July 2026
Version	Public Draft 1.0
Keywords	AI readiness; human interpretation; LAIR Framework; organizational judgment; sensemaking; responsible AI; AI governance

Recommended Citation

Meador, M. (2026). Developing human interpretation: Why AI readiness depends on human judgment, sensemaking, and responsibility. Learning AI Institute Research & Implementation Series (LAI-RIS-004). Learning AI Institute.

LAIR Framework, Groundmaps, and related terminology are original intellectual property created by Dr. Matt Meador through the Learning AI Institute and MetHer By Design, LLC.

© 2026 Learning AI Institute & MetHer By Design, LLC. All Rights Reserved.

Executive Summary

Artificial intelligence is transforming how organizations generate information, analyze data, and support decision-making. While many AI initiatives focus on improving access to AI technologies or teaching users how to write better prompts, these efforts overlook a more fundamental capability: human interpretation.

Human interpretation is the ability to critically evaluate, contextualize, and responsibly act upon AI-generated outputs. It is the point at which information becomes understanding and recommendations become informed decisions. Without this capability, organizations risk substituting AI-generated confidence for human judgment.

This publication argues that interpretation represents the third stage of organizational AI capability development within the Learning AI Institute's LAIR Framework (Literacy, Application, Interpretation, and Responsibility). Building on earlier publications addressing AI readiness, organizational assessment, and implementation, this paper demonstrates that successful AI adoption depends not simply on using AI effectively but on interpreting its outputs responsibly.

Key findings include:

AI readiness extends beyond technology adoption to include organizational judgment and decision-making capabilities.

Human interpretation is a measurable organizational competency that can be developed through training, practice, governance, and reflective processes.

Organizations that intentionally develop interpretive capability are better positioned to detect errors, recognize bias, preserve human agency, and make more effective decisions.

Interpretation serves as the operational bridge between AI application and AI responsibility, ensuring that AI remains a decision-support tool rather than a decision-maker.

For executives, educators, and organizational leaders, the implication is clear: sustainable AI capability depends less on acquiring better AI tools and more on developing better human judgment. Organizations that invest in interpretive capability today will be better prepared to govern AI responsibly, adapt to future technological change, and create lasting organizational value.

Abstract

Artificial intelligence is often framed as a tool for automation, acceleration, and productivity. Yet the deeper organizational challenge is not merely whether people can access or operate AI systems, but whether they can interpret AI-supported outputs responsibly. As generative AI

becomes embedded in education, business, government, and workforce systems, human interpretation is becoming a core readiness capability.

This paper argues that AI readiness cannot be achieved through access, literacy, or application alone. Organizations must also develop the human capacity to question, contextualize, evaluate, and act on AI-generated information. Without interpretation, AI use can produce passive dependency, misplaced confidence, and weakened judgment. With interpretation, AI becomes a partner in learning, decision-making, organizational sensemaking, and responsible action.

Building from the LAIR Framework—Literacy, Application, Interpretation, and Responsibility—this paper positions interpretation as the bridge between AI use and AI accountability. It defines human interpretation as an organizational capability that can be intentionally developed, measured, and governed. The paper also outlines the interpretive skills organizations need to strengthen AI readiness, including contextual judgment, critical evaluation, domain alignment, ethical awareness, and responsible use.

By centering human interpretation, this paper reframes AI readiness as a leadership, learning, and governance challenge rather than a technical adoption problem alone.

Introduction

The next failure point in artificial intelligence adoption will not simply be technical.

It will be interpretive.

Organizations across education, business, government, healthcare, and the nonprofit sector are rapidly adopting artificial intelligence to write, summarize, analyze, classify, predict, recommend, and generate information. These systems have dramatically lowered the barriers to creating content, synthesizing knowledge, and supporting decision-making. In many respects, AI represents one of the most significant technological shifts since the emergence of the internet, fundamentally changing how individuals access information and perform knowledge work (Brynjolfsson & McAfee, 2014).

Yet the same characteristics that make AI systems valuable also introduce new organizational risks. AI systems produce responses that are increasingly fluent, persuasive, and contextually appropriate. They often appear authoritative even when they contain factual inaccuracies, unsupported claims, fabricated references, hidden biases, or incomplete reasoning (Bender et al., 2021). Consequently, speed is not synonymous with understanding. Fluency is not equivalent to expertise. A convincing AI-generated response is not inherently a responsible organizational decision. This reality creates a new readiness problem.

Much of the current discussion surrounding AI implementation focuses on whether organizations possess the technology, infrastructure, policies, or training necessary to begin using artificial intelligence. Organizations frequently ask whether their employees know how to write prompts,

select AI tools, automate workflows, or integrate AI into daily operations. These questions remain important, but they are insufficient. Basic tool proficiency does not guarantee responsible AI use. Prompt engineering does not ensure critical evaluation. Access to AI does not automatically produce sound judgment.

The more important question is whether people know how to interpret what artificial intelligence produces.

Human interpretation is the capacity to critically examine AI-generated outputs before incorporating them into decisions, communication, instruction, policy, research, or organizational action. Interpretation requires individuals to evaluate the accuracy, completeness, contextual relevance, ethical implications, underlying assumptions, and practical consequences of AI-supported information. It extends beyond technical competence by requiring domain expertise, reflective thinking, ethical reasoning, and professional judgment. Rather than accepting AI-generated outputs as definitive answers, effective interpreters recognize AI as one source of information that must be evaluated within broader organizational, disciplinary, and societal contexts.

This paper builds directly upon the developing body of scholarship established through the Learning AI Institute Research and Implementation Series (LAI-RIS). In *Fixing AI Readiness*, Meador (2026a) argued that the primary challenge facing organizations is not a lack of AI technology but a lack of organizational readiness, emphasizing that sustainable implementation begins with people, processes, and governance rather than software acquisition. *Groundmaps: Measuring Organizational Readiness Before AI Implementation* expanded this argument by introducing a framework for assessing an organization's current capabilities before implementation begins, demonstrating that organizations must understand their starting point before they can successfully plan their AI future (Meador, 2026b). Building upon these foundations, *The LAIR Framework* proposed a developmental model describing how organizations progress from AI literacy to measurable organizational capability through the sequential development of Literacy, Application, Interpretation, and Responsibility (Meador, 2026c).

LAI-RIS-004 extends this research agenda by examining the interpretive dimension of AI readiness in greater depth. While previous publications established why organizations must prepare for AI adoption and how AI capability develops over time, this paper argues that human interpretation represents the critical capability connecting AI application to responsible organizational action. Interpretation is not simply another competency within the LAIR Framework; it is the developmental turning point at which AI-generated information becomes meaningful, actionable, and accountable through human judgment.

This focus has become increasingly important because AI adoption is no longer confined to specialized technical teams or data scientists. Generative AI has become accessible to virtually every knowledge worker. Teachers use AI to develop instructional materials. Executives

summarize reports with AI before making strategic decisions. Healthcare professionals draft documentation using AI-assisted tools. Government agencies explore AI-supported policy development. Small businesses rely on AI to generate marketing content, financial summaries, and operational recommendations. As AI use becomes more widespread, the responsibility for interpreting AI-generated outputs becomes equally distributed throughout organizations.

Consequently, the definition of AI readiness must evolve. Readiness can no longer be measured solely by technology deployment, software availability, policy development, or workforce participation in AI literacy training. Instead, organizations must evaluate whether individuals possess the judgment required to question AI outputs, recognize limitations, identify uncertainty, detect bias, and determine whether AI-supported recommendations are appropriate within specific organizational contexts. Without these capabilities, organizations risk creating cultures of passive acceptance in which AI-generated information is trusted because it is efficient, persuasive, or convenient rather than because it is accurate or contextually appropriate.

These concerns are increasingly reflected within international AI governance literature. The National Institute of Standards and Technology's Artificial Intelligence Risk Management Framework emphasizes that trustworthy AI requires continuous human oversight, organizational governance, and risk management rather than technical performance alone (NIST, 2023). UNESCO's AI Competency Framework for Teachers similarly emphasizes human agency, ethical reasoning, and professional judgment as essential components of responsible AI integration within educational environments (UNESCO, 2024). Likewise, the OECD AI Principles advocate for human-centered AI systems that promote accountability, transparency, fairness, and respect for democratic values (OECD, 2019).

Although these frameworks establish important principles for trustworthy AI, they provide comparatively limited guidance regarding how organizations develop interpretive capability as an organizational competency. This paper seeks to address that gap by proposing human interpretation as a measurable component of AI readiness that can be intentionally developed through structured learning experiences, reflective practice, governance processes, and organizational measurement.

The central claim of this paper is straightforward.

Artificial intelligence can generate information.

Only humans can determine its meaning.

Organizations that intentionally develop interpretive capability will be better positioned to preserve human agency, strengthen organizational judgment, reduce overreliance on automation, improve governance, and transform AI from a productivity tool into a sustainable capability that

supports responsible organizational decision-making. In doing so, this paper extends the LAIR Framework by establishing interpretation as the indispensable bridge between AI application and AI responsibility, completing the next stage in the Learning AI Institute's evolving theory of organizational AI readiness.

Central Argument

AI readiness fails when humans become output acceptors instead of meaning makers. The presence of artificial intelligence does not diminish the need for human judgment; it amplifies it. As AI systems become increasingly capable, conversationally fluent, and embedded within everyday organizational workflows, individuals may become more willing to accept AI-generated outputs without sufficient questioning. This tendency is particularly concerning because contemporary generative AI systems often produce responses that appear coherent, authoritative, and well-reasoned even when they contain factual inaccuracies, unsupported assumptions, incomplete analyses, fabricated citations, or contextual misunderstandings (Bender et al., 2021; OpenAI, 2025).

The Learning AI Institute has previously argued that organizations frequently misunderstand AI readiness by equating technology acquisition with organizational capability (Meador, 2026a). While organizations continue investing in AI platforms, training initiatives, and automation strategies, these investments alone do not ensure that employees, educators, leaders, or decision-makers possess the judgment required to evaluate AI-generated information responsibly. As argued in *Groundmaps: Measuring Organizational Readiness Before AI Implementation*, organizations must first understand their existing human and organizational capabilities before meaningful AI transformation can occur (Meador, 2026b). Likewise, the LAIR Framework positions AI capability as a developmental process that extends beyond technical proficiency toward increasingly sophisticated forms of human judgment and organizational responsibility (Meador, 2026c).

This creates what may be described as a false sense of AI readiness. An organization may appear AI-ready because its employees have access to AI tools, complete introductory literacy training, or routinely incorporate AI into daily workflows. Yet if those same individuals cannot critically interpret AI-supported outputs, the organization remains vulnerable. It may move faster without making better decisions. It may produce more content without developing deeper understanding. It may automate routine work while simultaneously weakening the very judgment required to govern AI responsibly.

Research on technology adoption has consistently demonstrated that successful implementation depends upon far more than technological access or perceived usefulness (Davis, 1989; Venkatesh et al., 2003). Similarly, organizational learning scholars argue that sustainable organizational capability emerges through reflection, interpretation, and continuous learning rather than through technology implementation alone (Argyris & Schön, 1978; Senge, 1990). AI

readiness should therefore be understood as a human capability challenge as much as a technological one.

Human interpretation is not an optional soft skill. It is a core AI readiness capability.

Interpretation represents the cognitive and organizational process through which individuals critically examine AI-generated outputs before integrating them into communication, instruction, policy, analysis, or decision-making. It requires users to ask whether AI-generated information is accurate, contextually appropriate, supported by evidence, ethically defensible, and aligned with organizational objectives. Rather than accepting AI responses as authoritative, effective interpreters actively evaluate assumptions, identify uncertainty, recognize missing information, and determine whether AI recommendations should be accepted, revised, supplemented, or rejected.

This process closely aligns with Weick's (1995) theory of organizational sensemaking, which argues that individuals construct meaning by interpreting ambiguous information within specific organizational contexts. Artificial intelligence can generate information rapidly, but it cannot fully understand institutional history, organizational culture, stakeholder priorities, or professional responsibility. Meaning remains a fundamentally human activity. AI generates possibilities; humans determine significance.

Within the Learning AI Institute's LAIR Framework, interpretation serves as the developmental bridge between AI Application and AI Responsibility (Meador, 2026c). Literacy enables individuals to understand AI. Application enables them to use AI effectively. Interpretation enables them to evaluate AI critically. Responsibility enables them to act ethically and accountably based upon that evaluation. Without interpretation, AI capability remains incomplete because organizations possess technical proficiency without corresponding judgment.

Organizations that intentionally cultivate human interpretation will be better positioned to identify weak or incomplete AI outputs, recognize hallucinations and unsupported claims, connect AI-generated information to authentic organizational contexts, preserve human agency in decision-making, reduce automation bias (Parasuraman & Riley, 1997), strengthen governance, and foster cultures of continuous organizational learning. In this way, interpretation transforms AI from an answer-producing technology into a decision-support capability that augments rather than replaces human expertise.

The central argument of this paper is therefore straightforward:

1. Artificial intelligence can generate information, but only humans can determine its meaning.
2. AI can recommend, but only humans can exercise judgment.
3. AI can accelerate work, but only humans can assume responsibility for the consequences of decisions.

Ultimately, sustainable AI readiness is not measured by how effectively organizations deploy artificial intelligence. It is measured by how consistently individuals demonstrate the judgment, contextual awareness, and ethical responsibility necessary to transform AI-generated outputs into informed, defensible, and accountable organizational decisions. Developing this interpretive capability represents the next stage in advancing AI readiness from technical adoption toward measurable organizational capability (Meador, 2026a, 2026b, 2026c).

Conceptual Model

The conceptual model presented in this paper positions human interpretation as the critical transition between AI use and responsible organizational decision-making. While existing research has extensively examined technology adoption, digital transformation, organizational readiness, and AI governance, comparatively little attention has been given to the interpretive work humans must perform once AI generates an output. The Learning AI Institute proposes that interpretation is the organizational capability that transforms AI-generated information into informed human judgment.

Conceptual Pathway

AI Access → AI Literacy → AI Application → Human Interpretation → Responsible Action → Organizational Learning → Measurable AI Capability

This progression reflects a developmental rather than purely technical model of AI readiness. Each stage is built upon the previous stage while introducing increasingly sophisticated forms of human capability.

AI Access

AI readiness begins with access. Individuals and organizations must first have reliable opportunities to engage with AI technologies. Access includes technical infrastructure, appropriate policies, secure platforms, equitable availability, and organizational support. However, access alone creates potential rather than capability. Numerous studies of technology adoption have demonstrated that simply providing technology rarely produces meaningful organizational change without complementary investments in people, leadership, and organizational processes (Davis, 1989; Venkatesh et al., 2003).

AI Literacy

Access must be followed by AI literacy. Literacy represents foundational understanding of AI concepts, capabilities, limitations, risks, and ethical considerations. Individuals develop sufficient knowledge to recognize what AI systems can accomplish, where they should be applied, and where caution is warranted. UNESCO (2024) identifies AI literacy as a prerequisite for responsible participation in AI-supported environments, while Long and Magerko (2020) argue that AI literacy requires conceptual understanding rather than simple technical proficiency.

AI Application

Once individuals possess foundational literacy, they begin applying AI to authentic work. Applications include prompt development, workflow integration, information retrieval, content generation, data analysis, and decision support. This stage emphasizes practical experience and repeated interaction with AI systems. However, application alone does not guarantee effective decision-making. Individuals may become proficient users while remaining poor evaluators of AI-generated information. Research on sociotechnical systems consistently demonstrates that effective technology implementation depends upon the interaction between technical capability and human judgment rather than technical capability alone (Baxter & Sommerville, 2011).

Human Interpretation

Human interpretation represents the conceptual contribution of this paper. Interpretation is the process through which individuals critically examine AI-generated outputs before integrating them into decisions, communication, instruction, policy, or organizational action. Rather than accepting AI responses at face value, interpreters evaluate contextual relevance, factual accuracy, domain alignment, ethical implications, completeness, uncertainty, and organizational consequences.

Interpretation requires individuals to engage in active sensemaking process through which people construct meaning from ambiguous or incomplete information (Weick, 1995). AI systems generate information, but humans determine its meaning. This distinction shifts organizational attention away from prompt engineering alone and toward judgment engineering: developing the capacity to recognize when AI should be trusted, questioned, revised, supplemented, or rejected.

Responsible Action

Interpretation naturally leads to responsible action. Once AI outputs have been critically evaluated, individuals make informed decisions regarding whether and how AI-generated information should influence organizational activities. Responsible action incorporates accountability, transparency, ethical reasoning, disclosure, verification, and governance. This stage aligns with recommendations from the NIST Artificial Intelligence Risk Management Framework (2023), which emphasizes that trustworthy AI requires ongoing human oversight, organizational governance, and continuous evaluation throughout AI deployment.

Organizational Learning

Responsible interpretation generates organizational learning. As individuals repeatedly evaluate AI outputs, organizations accumulate institutional knowledge regarding effective prompting, appropriate use cases, governance practices, common risks, and successful implementation strategies. This learning contributes to dynamic organizational capabilities that improve over time through experience and reflection (Teece, Pisano, & Shuen, 1997).

Measurable AI Capability

The final stage is measurable AI capability. Rather than defining AI maturity by the quantity of AI tools deployed, this model defines capability as the organization's demonstrated ability to produce reliable, ethical, contextually appropriate, and defensible outcomes through effective collaboration between humans and AI. Capability therefore becomes observable through improved decision quality, governance maturity, organizational learning, workforce confidence, and responsible AI practices.

Collectively, this model extends existing research on organizational readiness, technology adoption, and AI governance by identifying human interpretation as the missing developmental capability connecting AI application to responsible organizational performance. The framework suggests that sustainable AI readiness is not achieved when organizations merely deploy AI systems. It is achieved when people consistently demonstrate the judgment necessary to transform AI-generated outputs into informed, accountable, and contextually appropriate decisions.

Defining Human Interpretation

Human interpretation is the ability to make informed meaning from AI-generated outputs before using those outputs in decisions, communication, learning, policy, or organizational action. It is the process through which individuals examine what AI produces, evaluate its usefulness and limitations, and determine whether the output should be accepted, revised, rejected, verified, disclosed, or escalated.

Within the LAIR Framework, interpretation is the capability that connects AI application to AI responsibility (Meador, 2026c). Literacy helps users understand AI. Application helps users use AI. Interpretation helps users judge AI. Responsibility helps users act ethically and accountably based on that judgment. Without interpretation, AI use remains incomplete because individuals may be able to generate output without understanding whether those outputs are accurate, appropriate, or defensible. Human interpretation includes five connected capabilities.

Contextual Judgment

Contextual judgment is the ability to determine whether an AI-generated output fits the situation, audience, purpose, and constraints in which it will be used. AI systems may generate fluent responses, but they do not automatically understand local context, institutional history, stakeholder relationships, organizational priorities, or the consequences of action. Sensemaking research emphasizes that meaning is constructed within context, not simply extracted from information (Weick, 1995). For this reason, interpretation requires users to ask whether the AI output makes sense for the specific environment in which it will be applied.

This aligns with the Groundmaps model, which argues that organizations must understand their starting point before implementing AI-supported change (Meador, 2026b). Context matters

because an AI recommendation that appears useful in one organization may be inappropriate in another.

Critical Evaluation

Critical evaluation is the ability to identify gaps, errors, unsupported claims, hallucinations, bias, and overgeneralization in AI-generated outputs. Generative AI systems can produce language that appears confident and authoritative even when the information is incomplete or inaccurate (Bender et al., 2021). This creates risk when users mistake fluency for validity.

Critical evaluation requires users to verify claims, examine evidence, question assumptions, and recognize uncertainty. It also requires awareness of automation bias, where individuals over-rely on automated systems even when contradictory evidence is available (Parasuraman & Riley, 1997). In this sense, interpretation protects organizations from treating AI outputs as conclusions rather than inputs for judgment.

Domain Alignment

Domain alignment is the ability to compare AI-generated outputs against professional knowledge, disciplinary standards, organizational expectations, and lived context. AI may produce general responses, but responsible use requires determining whether those responses align with the standards of a specific field or institution. In education, this may involve instructional design principles, assessment validity, academic integrity, or developmental appropriateness. In business, it may involve compliance, strategy, customer expectations, or operational feasibility.

This capability reflects Meador's (2026a) argument that AI readiness is not simply a technical condition but an organizational capability. A person may know how to use AI, but without domain knowledge, that person may not know whether the AI output is appropriate.

Ethical Awareness

Ethical awareness is the ability to recognize when AI-generated outputs may create harm, exclusion, privacy concerns, misrepresentation, unfair advantage, or inappropriate delegation of responsibility. AI governance frameworks emphasize that trustworthy AI requires human oversight, accountability, fairness, transparency, and risk management (NIST, 2023; OECD, 2019; UNESCO, 2024).

Interpretation therefore includes ethical reasoning. Users must consider not only whether an AI output is useful, but whether it is appropriate to use. This distinction is essential because AI-generated outputs can be efficient while still being harmful, biased, misleading, or misaligned with organizational values.

Responsible Use

Responsible use is the ability to decide what to accept, revise, reject, disclose, verify, or escalate. Interpretation becomes meaningful only when it leads to accountable action. Within the LAIR Framework, responsibility follows interpretation because users must first evaluate AI outputs before deciding how those outputs should influence action (Meador, 2026c).

Responsible use also supports organizational learning. When individuals document decisions, reflect on AI-supported work, and identify recurring risks, organizations begin developing stronger AI capability over time (Argyris & Schön, 1978; Senge, 1990). Interpretation therefore operates at both the individual and organizational levels.

Human interpretation is not anti-AI. It is the condition that makes AI useful. AI can generate information, but humans must determine meaning, relevance, risk, and responsibility. Organizations that develop human interpretation are better positioned to move beyond tool use toward sustainable AI readiness, measurable capability, and responsible organizational action (Meador, 2026a, 2026b, 2026c).

Why Interpretation Matters Now

As AI systems become increasingly intuitive, conversational, and seamlessly integrated into everyday workflows, the cognitive effort required to produce sophisticated outputs continues to decline. Tasks that once required hours of research, writing, analysis, or synthesis can now be completed within minutes. While these capabilities create substantial opportunities for productivity and innovation, they also introduce a less visible challenge: the easier AI becomes to use, the easier it becomes to accept its outputs without sufficient reflection.

Historically, difficult technologies naturally slowed users down. Individuals were required to understand software interfaces, develop technical expertise, and manually verify information throughout the decision-making process. Today's generative AI systems remove much of that friction. Although this democratizes access to advanced computational capabilities, it also creates conditions in which organizational decision-making can begin to outpace human judgment. The result is that AI readiness should no longer be viewed primarily as a software implementation challenge. It is fundamentally a human capability challenge.

Previous work within the Learning AI Institute argued that organizations frequently mistake technology adoption for organizational readiness (Meador, 2026a). Groundmaps further demonstrated that implementation should begin by understanding existing organizational capability rather than assuming readiness based on technology acquisition (Meador, 2026b). The LAIR Framework subsequently positioned AI capability as a developmental progression from literacy through responsibility (Meador, 2026c). This paper extends that progression by arguing that interpretation represents the capability that determines whether AI contributes to organizational learning or organizational dependency.

The central problem is that many people will use AI without developing the capacity to interpret what AI produces. This distinction has profound implications across sectors. Within education, students may submit polished assignments, sophisticated essays, or well-constructed analyses that they cannot explain, defend, or extend through independent reasoning. The concern is not merely academic integrity; it is the possibility that learning itself becomes displaced by dependence upon generated content. If students no longer construct meaning from information, they may successfully complete assignments while failing to develop the knowledge structures those assignments were designed to cultivate (Bransford, Brown, & Cocking, 2000).

Within business, professionals increasingly rely on AI to summarize market research, generate strategic recommendations, analyze financial information, draft reports, and support operational planning. While these applications improve efficiency, they also create opportunities for organizations to adopt recommendations without critically evaluating underlying assumptions, missing variables, or contextual limitations. Technology adoption research has consistently demonstrated that effective organizational performance depends upon human judgment operating alongside technological capability rather than being replaced by it (Davis, 1989; Venkatesh et al., 2003).

Within government and public administration, AI-supported decision systems increasingly influence policy development, resource allocation, risk assessment, and citizen services. The National Institute of Standards and Technology emphasizes that trustworthy AI requires continuous human oversight, governance, and risk management throughout the AI lifecycle (NIST, 2023). Human interpretation is therefore not simply desirable; it represents an essential safeguard against inappropriate reliance on automated recommendations.

Executive leadership faces similar challenges. Leaders receive AI-generated executive summaries, strategic analyses, competitive intelligence, and predictive insights. These outputs can improve situation awareness, but they can also create the illusion of understanding. Reading an AI-generated summary is not equivalent to understanding organizational complexity. Effective leadership continues to depend upon contextual judgment, organizational experience, stakeholder awareness, and informed decision-making that extends beyond information generation alone.

The danger, therefore, is not limited to inaccurate information, it is interpretive outsourcing.

Interpretive outsourcing occurs when individuals gradually transfer the responsibility for evaluating questioning, synthesizing, and making meaning from information to artificial intelligence systems. Rather than using AI as a decision-support technology, users begin treating AI as a substitute for judgment. Over time, this shift risks weakening critical thinking, reducing reflective practice, diminishing professional expertise, and creating increasing dependence upon AI-generated conclusions.

This concern aligns closely with Weick's (1995) theory of organizational sensemaking, which argues that organizations create knowledge by interpreting ambiguous information rather than simply collecting it. Artificial intelligence can generate enormous quantities of information, but organizations still depend upon people to determine what that information means, why it matters, and how it should influence future action. Likewise, Argyris and Schön (1978) argue that organizational learning emerges through reflection and inquiry rather than through information acquisition alone. AI can accelerate access to information, but it cannot replace the human processes through which organizations learn. Interpretation therefore represents more than an individual cognitive skill. It is an organizational capability.

Organizations that intentionally cultivate interpretive capability will be better positioned to recognize hallucinations, identify unsupported claims, evaluate contextual relevance, detect ethical concerns, preserve human agency, strengthen governance, and develop cultures of continuous organizational learning (Meador, 2026c). Conversely, organizations that neglect interpretation may experience increasing efficiency while simultaneously weakening the judgment necessary for responsible AI adoption.

Ultimately, AI readiness depends not on how quickly organizations generate information but on how effectively they transform that information into informed, ethical, and accountable decisions.

When people stop interpreting, they stop learning.

When organizations stop interpreting, they stop building judgment.

And when judgment erodes, AI readiness becomes little more than technological dependence rather than sustainable organizational capability.

Position Within the LAIR Framework

The LAIR Framework defines AI capability through four connected dimensions:

Literacy → Application → Interpretation → Responsibility

Each dimension represents a different stage of human and organizational capability. Together, they describe how individuals and organizations move from basic awareness of artificial intelligence toward responsible, measurable AI use (Meador, 2026c).

Literacy asks: What is AI, and how does it work?

At this stage, individuals develop foundational understanding of AI concepts, capabilities, limitations, risks, and appropriate uses. AI literacy is increasingly recognized as necessary for responsible participation in AI-supported environments because users must understand both what AI can do and where its limitations remain (Long & Magerko, 2020; UNESCO, 2024). However, literacy alone does not produce capability. A person may understand AI conceptually while still being unable to apply it effectively or evaluate its outputs responsibly.

Application asks: How can AI be used for a specific task or purpose?

This stage involves the practical use of AI in authentic work, including drafting, summarizing, analyzing, generating, classifying, planning, and decision support. Application is where AI becomes operational. Users begin integrating AI into workflows, learning tasks, organizational processes, and professional routines. Technology adoption research has shown that use depends partly on perceived usefulness, ease of use, social influence, and facilitating conditions (Davis, 1989; Venkatesh et al., 2003). Yet adoption does not guarantee responsible or effective use. An organization may increase AI use while still lacking the judgment needed to evaluate AI-supported outcomes.

Interpretation asks: What does the AI output mean, and should it be trusted?

Interpretation is the turning point of the LAIR Framework. It is where users move from producing AI-supported outputs to evaluating those outputs. This requires critical judgment, contextual understanding, domain knowledge, and ethical awareness. Interpretation connects directly to organizational sensemaking because individuals must determine what AI-generated information means within a particular context before acting on it (Weick, 1995). In this way, interpretation transforms AI output from generated content into evaluated knowledge.

Responsibility asks: What should we do with this output, and who is accountable?

Responsibility is the stage at which interpretation becomes action. Users decide whether to accept, revise, reject, disclose, verify, document, or escalate AI-supported outputs. This stage aligns with major AI governance frameworks emphasizing accountability, transparency, human oversight, and risk management (NIST, 2023; OECD, 2019). Responsibility cannot occur meaningfully without interpretation because users must first evaluate AI outputs before deciding how those outputs should influence decisions or actions.

Many AI training programs stop at literacy or application. They teach people what AI is, how to write prompts, how to summarize documents, how to generate content, or how to automate routine tasks. These skills are useful, but they are not sufficient. Prompting skill alone does not produce responsible AI use. An individual can create a polished AI-generated output and still fail to recognize that the output is inaccurate, biased, incomplete, inappropriate, or ethically problematic.

This is why LAI-RIS-004 focuses specifically on interpretation. Earlier work in the LAI-RIS series established the need to fix AI readiness, measure organizational starting points through

Groundmaps, and develop AI capability through the LAIR Framework (Meador, 2026a, 2026b, 2026c). This paper extends that work by arguing that interpretation is the hinge between AI application and AI responsibility. Without interpretation, AI application can become passive acceptance. With interpretation, AI application becomes disciplined inquiry. The interpretive layer is where AI use becomes human-centered.

It is the point at which users stop treating AI as an answer machine and begin treating it as a thinking partner whose outputs must be questioned, contextualized, and evaluated. It is also the point at which organizations preserve human agency. Rather than delegating judgment to AI systems, organizations develop people who can work with AI while remaining accountable for meaning, context, ethics, and action.

Within the LAIR Framework, interpretation therefore functions as both a cognitive capability and an organizational safeguard. It helps individuals determine whether AI-generated information is meaningful, appropriate, and defensible. It also helps organizations prevent overreliance, strengthen governance, and build responsible AI cultures over time.

In this sense, the LAIR Framework does not position AI readiness as a matter of tool fluency alone. It positions AI readiness as the progressive development of human capability: understanding AI, applying AI, interpreting AI, and acting responsibly with AI. Interpretation is the developmental bridge that ensures AI use remains connected to judgment, learning, and accountability.

Executive Implications

The findings of this paper suggest that AI readiness is fundamentally a leadership challenge rather than solely a technology initiative. Executives should ask not only whether their organizations have adopted AI, but whether employees consistently demonstrate the judgment required to evaluate AI-generated outputs before acting upon them (Meador, 2026a; Meador, 2026c).

Leaders can identify gaps in interpretive capability by observing several indicators:

- Employees routinely accept AI-generated recommendations without verification.
- AI-generated reports or analyses cannot be explained or defended by their authors.
- Decisions rely on AI outputs with little evidence of contextual evaluation.
- Governance focuses primarily on tool access rather than decision quality.
- Training emphasizes prompting and productivity while neglecting critical evaluation and ethical reasoning.

For CEOs, CIOs, provosts, superintendents, and executive leaders, the priority should be developing organizational judgment. This includes embedding interpretive practice into professional development, requiring transparent documentation of AI-supported decisions,

establishing governance processes that require human review for high-impact decisions, and measuring capability through evidence of contextual judgment, critical evaluation, ethical awareness, and responsible use (NIST, 2023; OECD, 2019; UNESCO, 2024).

Interpretation should be measured as an organizational capability rather than an individual opinion. Potential measures include the quality of AI-supported decisions, the frequency of verification, evidence of reflective practice, governance compliance, reduction of automation bias, and demonstrated ability to explain and defend AI-assisted work. Organizations that measure interpretation are better positioned to transform AI from a productivity tool into a sustainable organizational capability.

Implementation Model

The Learning AI Institute recommends an iterative implementation pathway that integrates the contributions of the LAI-RIS series into a continuous organizational improvement cycle (Meador, 2026a, 2026b, 2026c).

1. Assess Organizational Readiness (Groundmaps)

Conduct a baseline assessment of organizational readiness to identify strengths, gaps, risks, and implementation priorities.

2. Build AI Literacy

Develop foundational understanding of AI capabilities, limitations, ethics, governance, and organizational expectations.

3. Practice AI Application

Provide structured opportunities for employees, educators, and leaders to apply AI within authentic organizational workflows.

4. Develop Interpretive Judgment

Train users to evaluate AI-generated outputs through contextual judgment, critical evaluation, domain alignment, ethical awareness, and responsible use.

5. Establish Governance

Implement policies, oversight structures, accountability mechanisms, and human review processes consistent with trustworthy AI practices (NIST, 2023).

6. Measure Organizational Capability

Evaluate progress through measurable indicators of AI literacy, interpretive capability, governance maturity, decision quality, and organizational learning.

7. Repeat

AI readiness is not a one-time implementation project. Organizations should continuously reassess readiness, refine capability, update governance, and strengthen interpretation as AI technologies evolve. This continuous improvement cycle aligns with organizational learning theory (Argyris & Schön, 1978; Senge, 1990) and reinforces the Learning AI Institute's vision of measurable AI capability.

Conclusion

AI readiness is not complete when people gain access to tools. It is not complete when they learn prompt techniques. It is not complete when they begin applying AI to workplace, instructional, organizational, or learning tasks. These activities represent important stages of adoption, but they do not by themselves establish responsible AI capability. AI readiness becomes meaningful when people can interpret AI-supported outputs with judgment and responsibility.

This paper argued that human interpretation is the critical bridge between AI application and AI responsibility. Building on previous work in the LAI-RIS series, this paper extends the argument that AI readiness must be understood as a human and organizational capability rather than a technology deployment problem alone (Meador, 2026a, 2026b, 2026c). Within the LAIR Framework, interpretation functions as the turning point where AI use becomes subject to human judgment, contextual evaluation, ethical awareness, and accountable action.

The need for interpretation is especially urgent because AI systems are becoming easier to use, more widely available, and more deeply embedded in everyday work. As friction decreases, the risk of passive acceptance increases. Users may become more likely to treat AI-generated outputs as reliable because they are fluent, polished, and immediate. Yet as research on automation bias and human-automation interaction has shown, people can over-rely on automated systems even when those systems are incomplete or wrong (Parasuraman & Riley, 1997). This makes interpretation a necessary safeguard against overdependence, misinformation, and weakened judgment.

Developing human interpretation requires more than awareness. It requires intentional training, guided practice, reflective inquiry, governance structures, and measurement systems. Organizations must help people ask better questions of AI output: Is this accurate? Is it complete? Is it contextually appropriate? What assumptions are being made? What evidence is missing? Who is affected by this recommendation? Can this decision be defended? These questions are not peripheral to AI readiness. They are central to it.

This paper also introduced interpretation as an organizational learning capability. When individuals interpret AI outputs critically, they do more than prevent error. They strengthen the organization's ability to learn from AI-supported work. They identify patterns, risks, limitations, and opportunities that can inform future practice. In this sense, human interpretation aligns with broader theories of organizational learning and sensemaking, which emphasize that organizations

build capability by reflecting on experience, interpreting ambiguous information, and adapting over time (Argyris & Schön, 1978; Weick, 1995).

For leaders, the implication is clear: the value of AI does not come from removing humans from the loop. It comes from strengthening the human capacity to think, judge, and act in AI-supported environments. Organizations that treat AI primarily as an automation tool may gain short-term productivity while weakening the interpretive judgment needed for long-term capability. Organizations that develop interpretation as a core readiness capability will be better positioned to preserve human agency, improve decision quality, strengthen governance, and build responsible AI cultures. The future of AI readiness is not tool fluency alone.

It is human interpretation.

AI can generate information. It can summarize, recommend, classify, predict, and produce. But it cannot assume responsibility for meaning. It cannot fully understand organizational context, ethical consequence, professional standards, or human purpose. Those responsibilities remain with people.

Therefore, the central challenge for organizations is not simply to adopt AI, but to develop people who can interpret AI well. Sustainable AI readiness depends on that capability. Without interpretation, AI adoption risks becoming technological dependence. With interpretation, AI becomes a disciplined partner in learning, judgment, governance, and responsible organizational action.

References

- Argyris, C., & Schön, D. A. (1978). *Organizational learning: A theory of action perspective*. Addison-Wesley.
- Baxter, G., & Sommerville, I. (2011). Socio-technical systems: From design methods to systems engineering. *Interacting with Computers*, 23(1), 4-17.
<https://doi.org/10.1016/j.intcom.2010.07.003>
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610-623.
<https://doi.org/10.1145/3442188.3445922>
- Bransford, J. D., Brown, A. L., & Cocking, R. R. (Eds.). (2000). *How people learn: Brain, mind, experience, and school*. National Academy Press.

- Brynjolfsson, E., & McAfee, A. (2014). *The second machine age: Work, progress, and prosperity in a time of brilliant technologies*. W. W. Norton.
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319-340. <https://doi.org/10.2307/249008>
- Long, D., & Magerko, B. (2020). What is AI literacy? Competencies and design considerations. *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, 1-16. <https://doi.org/10.1145/3313831.3376727>
- Meador, M. (2026a). *Fixing AI readiness*. Learning AI Institute Research & Implementation Series (LAI-RIS-001). Learning AI Institute.
- Meador, M. (2026b). *Groundmaps: Measuring organizational readiness before AI implementation*. Learning AI Institute Research & Implementation Series (LAI-RIS-002). Learning AI Institute.
- Meador, M. (2026c). *The LAIR Framework: From AI readiness to measurable organizational capability*. Learning AI Institute Research & Implementation Series (LAI-RIS-003). Learning AI Institute.
- National Institute of Standards and Technology. (2023). *Artificial intelligence risk management framework (AI RMF 1.0) (NIST AI 100-1)*. U.S. Department of Commerce. <https://doi.org/10.6028/NIST.AI.100-1>
- OECD. (2019). *OECD principles on artificial intelligence*. Organisation for Economic Co-operation and Development. <https://oecd.ai/en/ai-principles>
- OpenAI. (2025). *Model Spec*. OpenAI.
- Parasuraman, R., & Riley, V. (1997). Humans and automation: Use, misuse, disuse, abuse. *Human Factors*, 39(2), 230-253. <https://doi.org/10.1518/001872097778543886>
- Senge, P. M. (1990). *The fifth discipline: The art and practice of the learning organization*. Doubleday.
- Teece, D. J., Pisano, G., & Shuen, A. (1997). Dynamic capabilities and strategic management. *Strategic Management Journal*, 18(7), 509-533.
- UNESCO. (2024). *AI competency framework for teachers*. United Nations Educational, Scientific and Cultural Organization.
- Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS Quarterly*, 27(3), 425-478. <https://doi.org/10.2307/30036540>
- Weick, K. E. (1995). *Sensemaking in organizations*. Sage Publications.

About the Organizations

Learning AI Institute advances research, education, and organizational capability in artificial intelligence readiness, literacy, interpretation, and responsible implementation.

MetHer By Design partners with organizations to strengthen AI readiness, governance, workforce capability, and decision intelligence through practical frameworks, assessment tools, and implementation support.