

Journal of Integrated Science, Technology and Management

Research Article

Independent Verification as a Precondition for Trustworthy AI Deployment: The Case for a Formal AI Evaluation Discipline

Isaac Kubvoruno*

Mindrac AI

ARTICLE INFO

Article history:

Received: 02 November, 2025

Revised: 19 November, 2025

Accepted: 15 December, 2025

Published: 31 December, 2025

Keywords:

AI, AI governance, trustworthy AI, independent AI evaluation, third-party verification, AI risk management, AI assurance.

ABSTRACT

As the healthcare, financial, legal, government, and other high-impact sectors rapidly adopt artificial intelligence (AI), there has been a growing urgency for strong systems that guarantee reliability, accountability, and trust in the system. While modern principles for AI governance include transparency, risk management, fairness, and humans in the loop, the methods for assessing AI systems are still mostly decentralized, organization-specific, and relying on developer-driven assessments. The scale of the deployment of AI and the relative immaturity of independent evaluation introduce issues about the quality and authenticity of AI's role in the safety and mission-critical domain.

This paper presents a proposal to consider independent, third-party evaluation as a structural precondition for the trustworthy deployment of AI as opposed to an optional quality assurance mechanism. With a conceptual and policy-oriented approach, the study brings together insights from the field of AI governance literature, regulatory guidance, international standards, benchmark research, and practices in financial auditing, model risk management, and clinical oversight. It explores the challenges of developer self-assessment conflicts of interest, methodological differences to evaluate the problem, and a lack of external accountability and its impact on industries where AI decisions have substantial societal and economic implications.

The paper also reviews the existing regulatory framework, highlighting that the framework offers detailed guidelines on the principles for managing and accountability over AI risks but lacks a prescriptive approach on how to practically assess the performance of an AI independently. It also separates performance assessment, which is based on benchmarks, from the requirements of an accepted assessment discipline, noting that an evaluator needs to be independent; assessment methods need to be reproducible; the documentation needs to be transparent; assessment criteria need to be standardized; and evidence needs to be auditable and support regulatory compliance and organizational governance.

The study finds that, contrary to what some may think, there's no problem with evaluations, but there is one problem with an objective, independent, and professionally structured evaluation discipline that produces reliable, honest evaluations in various application fields. The paper highlights the unfilled governance space, thereby advancing existing scholarly and policy debates on credible AI, and offers a theoretical basis for the reasons why independence and standardization are crucial attributes of credible AI evaluation.

INTRODUCTION

Artificial intelligence (AI) has evolved from a specialized computational technology into a foundational component of decision-making across healthcare, financial services, legal practice, manufacturing, government, and critical infrastructure. The widespread adoption of large

language models, predictive analytics, autonomous systems, and machine learning applications has created unprecedented opportunities to improve efficiency, automate complex tasks, and support evidence-based decision-making. At the same time, this rapid technological advancement has intensified concerns regarding the

*Corresponding Author: Isaac Kubvoruno

Address: Mindrac AI

Email ✉: Isaac@mindrac.com

Relevant conflicts of interest/financial disclosures: The authors declare that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.

© 2025, Isaac Kubvoruno, This is an open access journal, and articles are distributed under the terms of the Creative Commons Attribution-NonCommercial-ShareAlike 4.0 License, which allows others to remix, tweak, and build upon the work non-commercially, as long as appropriate credit is given and the new creations are licensed under the identical terms.

reliability, transparency, fairness, explainability, and accountability of AI systems, particularly when they are deployed in environments where inaccurate or biased outputs may produce significant societal, economic, or legal consequences (Thiebes *et al.*, 2021; Li *et al.*, 2023). As AI systems increasingly influence high-impact decisions affecting individuals and institutions, ensuring that these systems can be independently assessed for trustworthiness has become a central challenge for researchers, policymakers, regulators, and industry stakeholders.

The concept of trustworthy AI has consequently emerged as one of the defining themes of contemporary AI governance. Rather than focusing solely on algorithmic performance, trustworthy AI encompasses a broader set of technical, ethical, legal, and organizational requirements that collectively determine whether AI systems can be deployed responsibly and safely. These requirements commonly include robustness, fairness, transparency, accountability, privacy protection, explainability, human oversight, and reliability throughout the AI lifecycle (Díaz Rodríguez *et al.*, 2023; Hohma & Lütge, 2023). Numerous international initiatives have translated these principles into governance frameworks and regulatory guidance, emphasizing that trust in AI cannot be achieved merely through technological sophistication but must also be supported by systematic risk management and verifiable assurance mechanisms (Li *et al.*, 2023; Kowald *et al.*, 2024). Consequently, the discussion surrounding AI governance has gradually shifted from identifying desirable principles toward determining how those principles can be demonstrated objectively in practice.

Despite these advances, a significant imbalance exists between the pace of AI deployment and the maturity of mechanisms available to evaluate AI systems independently. Most organizations continue to rely primarily on internal testing, developer-generated documentation, benchmark performance results, and organization-specific validation procedures before deploying AI models into operational environments. While these practices provide valuable insights into system performance, they often suffer from inherent limitations associated with conflicts of interest, methodological inconsistency, limited external scrutiny, and varying standards of evidence. Developer self-assessment, regardless of technical rigor, may not provide sufficient assurance for external stakeholders who require objective evidence that an AI system satisfies established expectations for safety, fairness, robustness, and regulatory compliance (Brundage *et al.*, 2020; Jacovi *et al.*, 2021). As AI systems assume increasingly consequential roles in society, reliance on internally generated evidence alone presents important governance and accountability challenges.

The need for objective verification has long been recognized in other high-consequence disciplines where

independent oversight serves as a fundamental mechanism for establishing credibility and public confidence. Financial auditing relies on external auditors to validate financial reporting; clinical research depends upon independent ethics review and regulatory oversight before new medical interventions are approved; and model risk management frameworks require independent validation to ensure that quantitative models perform as intended before they influence critical business decisions. These mature disciplines recognize that credibility derives not solely from technical competence but also from evaluator independence, transparent methodology, reproducibility, and documented evidence capable of withstanding external scrutiny. Comparable expectations are increasingly being discussed within AI governance, where verification and assurance are viewed as essential components of trustworthy AI deployment (Gisolfi, 2022; Winter *et al.*, 2021).

Recent scholarship has consequently begun exploring mechanisms capable of strengthening confidence in AI systems beyond traditional performance evaluation. Research on trustworthy AI has emphasized structured assessment methodologies, verification processes, certification concepts, and practical evaluation approaches intended to bridge the gap between ethical principles and operational implementation (Zicari *et al.*, 2021; Zicari *et al.*, 2022). Similarly, advances in formal verification seek to provide mathematical guarantees regarding specific system properties, offering complementary approaches for improving reliability in machine learning applications (Groß, 2024). At the same time, researchers have highlighted the importance of developing comprehensive assurance methods capable of evaluating not only predictive accuracy but also broader dimensions of trustworthiness, including fairness, transparency, robustness, governance, and human oversight (Stettinger *et al.*, 2024; Kowald *et al.*, 2024). Collectively, these studies demonstrate increasing recognition that trustworthy AI requires evaluation methodologies extending well beyond conventional benchmarking and performance testing.

Nevertheless, the existing landscape of AI evaluation remains fragmented and largely organization-specific. Benchmark datasets and standardized performance evaluations have significantly advanced the measurement of AI capabilities, yet they primarily assess technical performance rather than providing comprehensive evidence suitable for governance, regulatory assurance, or organizational accountability. Existing evaluation practices differ substantially across organizations with respect to testing protocols, documentation standards, evaluator independence, reproducibility, and reporting practices, making it difficult to compare results consistently or establish uniform levels of confidence across different AI systems (Brundage *et al.*, 2020; Li *et al.*, 2023). Consequently, although numerous frameworks articulate what constitutes trustworthy AI, considerably

less consensus exists regarding how trustworthiness should be evaluated through standardized, independent, and auditable processes.

This paper argues that independent, third-party evaluation should be regarded as a structural prerequisite for trustworthy AI deployment rather than as an optional extension of internal quality assurance. The argument is not that existing governance frameworks or benchmark methodologies are inadequate, but rather that they do not collectively constitute a formal evaluation discipline characterized by evaluator independence, standardized assessment criteria, reproducible methodologies, transparent documentation, and auditable evidence. Establishing such independence is essential for strengthening the credibility of AI assessments, reducing conflicts of interest, improving stakeholder confidence, and supporting consistent governance across diverse application domains.

To develop this argument, the paper first examines the structural rationale for independent evaluation by drawing parallels with established practices in financial auditing, model risk management, and clinical oversight. It then analyzes the risks associated with relying primarily on developer-led evaluation, including issues related to hallucinations, bias, reasoning failures, and inconsistent system behavior in high-impact sectors. The paper subsequently reviews the trajectory of contemporary AI governance and regulatory initiatives to demonstrate that, while substantial progress has been made in defining principles for trustworthy AI, an important gap remains regarding operational methodologies for independent evaluation. Finally, it argues that the absence of a formal, professionalized evaluation discipline represents one of the most significant unresolved challenges in AI governance, concluding that independence and methodological standardization are fundamental prerequisites for establishing credible and trustworthy AI evaluation.

The Structural Case for Independence

The increasing deployment of artificial intelligence (AI) in high-impact domains has elevated evaluation from a technical exercise to a fundamental governance requirement. As AI systems increasingly influence decisions in healthcare, finance, legal services, public administration, and critical infrastructure, confidence in their outputs depends not only on model performance but also on the credibility of the processes used to assess them. While developers possess extensive knowledge of their systems, evaluation conducted solely within the development environment presents structural limitations that can undermine objectivity, transparency, and public confidence. Consequently, independence should be regarded as an essential characteristic of trustworthy AI evaluation rather than an optional enhancement to quality assurance (Jacovi *et al.*, 2021; Li *et al.*, 2023).

Unlike conventional software, modern AI systems exhibit probabilistic behavior, evolve through retraining, and frequently depend upon complex datasets whose characteristics may change over time. These properties introduce uncertainties that are difficult to identify using internal validation alone. The challenge therefore extends beyond verifying whether an AI system functions correctly under laboratory conditions to determining whether its behavior remains reliable, fair, explainable, and robust across diverse operational contexts. As trustworthy AI increasingly becomes a regulatory and societal expectation, the independence of evaluation emerges as a structural necessity for establishing confidence in AI-enabled decision-making (Díaz Rodríguez *et al.*, 2023; Kowald *et al.*, 2024).

Why Developer Self-Assessment Is Structurally Insufficient

Developer-led evaluation constitutes the dominant practice across much of the AI industry because organizations naturally possess detailed knowledge of their models, datasets, architectures, and intended operational environments. Internal testing remains indispensable during model development, enabling rapid iteration, debugging, optimization, and performance improvement. Nevertheless, when internal evaluation becomes the sole basis for declaring an AI system trustworthy, several structural weaknesses arise that cannot be resolved merely by improving technical testing procedures.

The first limitation concerns conflicts of interest. Organizations investing significant financial, technical, and commercial resources in AI development possess inherent incentives to demonstrate favorable performance outcomes. Although this does not necessarily imply intentional bias, institutional pressures, including competitive market dynamics, product release schedules, contractual obligations, and investor expectations, may unintentionally influence evaluation priorities, reporting practices, or acceptable performance thresholds. Independent evaluation introduces organizational separation between developers and assessors, thereby reducing the likelihood that commercial incentives influence technical judgments (Brundage *et al.*, 2020).

A second limitation involves confirmation bias during validation. Development teams frequently design evaluation datasets, select performance metrics, and interpret experimental results according to assumptions established during model construction. Such familiarity can inadvertently restrict the diversity of testing scenarios and obscure unexpected failure modes. Independent evaluators are more likely to challenge design assumptions, examine overlooked operational contexts, and identify weaknesses that remain invisible within developer-controlled assessment environments (Jacovi *et al.*, 2021).

A third concern relates to methodological inconsistency. Organizations often employ proprietary testing methodologies tailored to internal objectives, making comparisons across AI systems difficult. Differences in benchmark selection, evaluation metrics, documentation practices, and reporting standards create considerable variability in what constitutes acceptable AI performance. Without standardized independent evaluation procedures, stakeholders struggle to interpret evaluation results consistently across organizations and application domains (Li *et al.*, 2023).

Furthermore, developer-led assessments frequently emphasize technical performance indicators such as accuracy, precision, recall, latency, or benchmark scores while giving comparatively less attention to broader governance considerations, including fairness, explainability, robustness, transparency, human oversight, and accountability. Trustworthy AI requires evaluation that integrates technical, ethical, organizational, and societal dimensions rather than focusing exclusively on predictive performance (Thiebes *et al.*, 2021).

Collectively, these limitations demonstrate that the challenge is not the competence of AI developers but the structural constraints associated with self-assessment. Independence enhances credibility precisely because it introduces external scrutiny capable of validating claims through objective and reproducible methodologies.

Lessons from Established Independent Assessment Disciplines

The importance of independent evaluation is well established across numerous professional disciplines where public trust depends upon objective verification. AI governance can draw important conceptual insights from these mature domains without suggesting direct equivalence.

Financial auditing provides one of the clearest precedents. Corporate financial statements are initially prepared internally; however, their credibility depends upon independent external auditors who verify compliance with established accounting standards. Auditor independence reduces conflicts of interest while strengthening confidence among investors, regulators, and the public. The value of auditing therefore lies not only in technical verification but also in institutional separation between those producing information and those evaluating it.

A comparable principle exists in model risk management within the banking sector. Regulatory guidance such as SR 11-7 emphasizes independent validation of quantitative models used in financial decision-making. Validation teams remain organizationally separate from model developers and assess assumptions, implementation, data quality, sensitivity analyses, and ongoing performance. This separation acknowledges that internal development

alone cannot provide sufficient assurance for models influencing high-consequence decisions.

Clinical research offers another instructive example. Pharmaceutical products undergo independent ethical review, external clinical evaluation, and regulatory assessment before approval for widespread use. Although developers conduct extensive internal testing, independent review remains essential for ensuring scientific rigor, participant safety, and public confidence. Similar governance principles increasingly apply to AI systems deployed in healthcare and other safety-critical sectors.

Engineering certification and conformity assessment likewise rely upon independent verification before products enter regulated markets. Electrical equipment, aviation systems, automotive components, and medical devices undergo structured evaluation against predefined criteria conducted by organizations independent from manufacturers. These industries recognize that credibility derives from both technical competence and evaluator independence.

The recurring principle across these disciplines is clear: objective evaluation requires organizational separation between developers and evaluators, standardized methodologies, transparent documentation, reproducibility, and accountability. AI systems increasingly exhibit societal impacts comparable to these regulated technologies, suggesting that similar structural principles are appropriate for AI evaluation (Winter *et al.*, 2021; Stettinger *et al.*, 2024).

Independence as a Foundation for Trustworthy AI

Trustworthiness cannot be established solely through claims of technical excellence. Rather, trust emerges when stakeholders possess confidence that evaluation processes themselves are objective, transparent, reproducible, and resistant to institutional bias. Jacovi *et al.* (2021) distinguish between the technical properties of AI systems and the broader conditions necessary for human trust, emphasizing that trustworthy AI depends as much upon credible evaluation processes as upon model performance.

Several recent studies similarly argue that trustworthy AI requires moving beyond abstract ethical principles toward operational governance mechanisms capable of demonstrating compliance in practice (Hohma & Lütge, 2023; Li *et al.*, 2023). Independent evaluation contributes directly to this objective by providing structured evidence supporting claims regarding robustness, fairness, safety, explainability, and accountability.

Emerging approaches to AI verification and trustworthy assessment likewise emphasize systematic documentation, evidence generation, and reproducible evaluation procedures. Model-centric verification highlights the importance of rigorous validation across the AI lifecycle rather than relying exclusively on final performance metrics (Gisolfi, 2022). Likewise, practical

assessment methodologies emphasize comprehensive documentation, multidisciplinary evaluation, and continuous evidence collection to support trustworthy AI deployment (Zicari *et al.*, 2021; Zicari *et al.*, 2022).

Importantly, independence should not be interpreted as replacing developer responsibilities. Developers remain responsible for designing, testing, monitoring, and improving AI systems throughout their lifecycle. Independent evaluation instead serves as an additional governance layer that enhances the credibility of technical claims by providing objective verification using transparent and reproducible methodologies. This complementary relationship mirrors mature governance structures found in finance, healthcare, engineering, and other regulated sectors.

Core Characteristics of Independent AI Evaluation

An effective independent evaluation discipline should be defined not by specific organizational arrangements but by universally applicable structural characteristics. First, evaluators should maintain organizational independence from AI developers to minimize conflicts of interest. Second, evaluation methodologies should employ predefined and transparent assessment criteria capable of consistent application across comparable systems. Third, testing procedures should be reproducible, enabling independent replication of findings under similar conditions. Fourth, comprehensive documentation and audit trails should record datasets, evaluation protocols, assumptions, and decision rationales. Finally, multidisciplinary expertise, including technical, legal, ethical, and domain-specific knowledge, should inform evaluation processes to ensure that assessments address the full spectrum of AI risks rather than technical performance alone (Kowald *et al.*, 2024; Díaz Rodríguez *et al.*, 2023).

Collectively, these characteristics distinguish independent evaluation from conventional internal testing. Whereas internal validation primarily supports model development and optimization, independent evaluation strengthens public trust by providing objective evidence that AI systems satisfy consistent standards of reliability, transparency, accountability, and governance. As AI adoption continues to expand across critical sectors, these structural attributes increasingly represent foundational requirements for credible AI assurance rather than optional enhancements to existing quality assurance practices.

The Risk Landscape without Independent Evaluation

The accelerating adoption of artificial intelligence (AI) across critical sectors has heightened concerns regarding the reliability, safety, and accountability of AI-driven decisions. While organizations increasingly implement internal testing protocols prior to deployment, the absence of independent evaluation introduces structural vulnerabilities that extend beyond technical performance. Developer-led assessments are often constrained by organizational objectives, proprietary methodologies, and limited external scrutiny, making it difficult to determine whether reported performance accurately reflects real-world reliability across diverse operational environments. Consequently, independent evaluation should not be viewed merely as an additional quality assurance activity but as an essential governance mechanism capable of mitigating systemic risks associated with increasingly autonomous AI systems (Brundage *et al.*, 2020; Li *et al.*, 2023).

The risks associated with insufficient independent evaluation are multidimensional. Beyond conventional

Table 1: Structural Comparison Between Developer-Led and Independent AI Evaluation

<i>Evaluation dimension</i>	<i>Developer-led evaluation</i>	<i>Independent evaluation</i>
Primary Objective	Product development and optimization	Objective verification of trustworthiness
Organizational Relationship	Conducted by the AI developer Higher due to commercial and institutional incentives	Conducted by an organizationally independent evaluator
Potential Conflict of Interest	Often organization-specific	Significantly reduced through evaluator independence
Evaluation Criteria	Variable; may involve proprietary reporting	Standardized and consistently applied
Transparency		Documented and auditable assessment procedures
Reproducibility	Often limited outside the organization	Designed for independent replication
Stakeholder Confidence	Dependent on developer credibility	Enhanced through objective external verification
Governance Value	Supports internal quality assurance	Strengthens regulatory, organizational, and public trust
Scope of Assessment	Frequently focused on technical performance	Incorporates technical, ethical, governance, and operational dimensions
Primary Outcome	Improved model performance	Credible evidence supporting trustworthy AI deployment

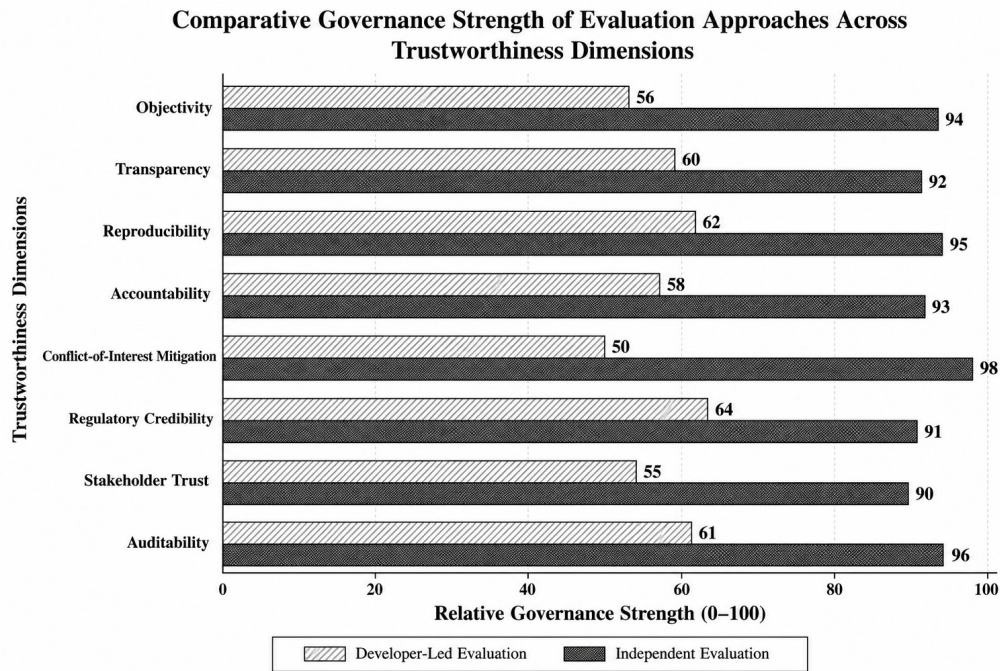


Figure 1: Comparative Governance Value of Developer-Led Versus Independent AI Evaluation Across Core Trustworthiness Dimensions

concerns regarding software defects, AI systems introduce probabilistic behavior, adaptive learning capabilities, and context-sensitive reasoning that cannot always be fully characterized through internal validation alone. Contemporary research consistently argues that trustworthy AI depends not only on technical robustness but also on transparent verification, reproducible assessment methodologies, and objective external scrutiny capable of establishing confidence among regulators, organizations, and end users (Jacovi *et al.*, 2021; Thiebes *et al.*, 2021; Díaz Rodríguez *et al.*, 2023). Without independent verification, critical weaknesses may remain undetected until systems have already been deployed in environments where errors carry substantial legal, financial, or societal consequences.

Hallucination and Fabricated Outputs

One of the most widely recognized risks in modern AI systems is the generation of hallucinated outputs, whereby models produce information that appears coherent and authoritative despite being factually incorrect or entirely fabricated. Unlike traditional software failures that often produce observable errors, hallucinations frequently remain undetected because generated responses exhibit convincing linguistic fluency. This characteristic presents significant challenges in domains where factual accuracy directly influences professional decision-making, including healthcare, legal practice, scientific research, and public administration.

Developer-led testing frequently evaluates models using predefined benchmark datasets or anticipated

use cases, yet these assessments may fail to capture the broad diversity of real-world operational contexts in which hallucinations emerge. Independent evaluation provides an additional layer of scrutiny by examining system behavior across diverse scenarios, adversarial prompts, and representative user interactions that extend beyond internally selected evaluation conditions. Such assessments improve the likelihood of identifying latent vulnerabilities that remain invisible during organization-specific validation processes (Gisolfi, 2022; Brundage *et al.*, 2020).

Furthermore, hallucination risks extend beyond isolated factual inaccuracies. Fabricated citations, incorrect numerical calculations, invented legal precedents, and misleading medical recommendations may all undermine stakeholder confidence while exposing organizations to substantial reputational and regulatory liabilities. As AI systems become increasingly integrated into decision-support workflows, independent verification becomes essential for objectively determining the frequency, severity, and operational significance of hallucinated outputs before deployment (Li *et al.*, 2023; Kowald *et al.*, 2024).

Bias Amplification and Discriminatory Outcomes

Bias represents another significant challenge arising from insufficiently independent evaluation. AI systems frequently learn statistical patterns embedded within historical datasets, inadvertently reproducing or amplifying pre-existing social, demographic, or institutional inequalities. Although developers increasingly implement fairness

testing during model development, internal assessments may overlook demographic groups, contextual variables, or deployment environments that differ from those originally considered during model design.

Independent evaluation introduces methodological diversity that reduces the likelihood of systematic blind spots. External evaluators can assess fairness using transparent criteria, broader population samples, and independently selected test conditions, thereby providing a more objective understanding of algorithmic behavior across different demographic groups. Such assessments are particularly important because fairness cannot be adequately established through isolated performance metrics but requires comprehensive examination of model behavior across multiple protected characteristics and operational contexts (Zicari *et al.*, 2021; Zicari *et al.*, 2022).

The implications of biased AI systems extend beyond technical performance. Discriminatory predictions may influence employment opportunities, financial lending decisions, healthcare access, insurance eligibility, educational admissions, and criminal justice outcomes. These consequences reinforce the argument that fairness assessment should not remain solely under developer control, especially where AI systems directly affect individual rights and societal welfare. Research on trustworthy AI consistently emphasizes that objective evaluation mechanisms strengthen confidence by providing independent evidence regarding fairness, accountability, and ethical compliance (Thiebes *et al.*, 2021; Hohma & Lütge, 2023; Li *et al.*, 2023).

Reasoning Failures and Inconsistent Decision-Making

Beyond hallucinations and bias, advanced AI systems may exhibit reasoning failures characterized by inconsistent outputs, unstable decision-making, and sensitivity to minor contextual variations. Identical questions presented with slight linguistic modifications may generate substantially different responses despite requiring equivalent logical reasoning. Such variability challenges the assumption that high benchmark performance necessarily translates into dependable operational behavior.

Developer-conducted evaluations often prioritize aggregate accuracy scores or benchmark performance, which may obscure inconsistencies emerging under realistic deployment conditions. Independent evaluation complements these metrics by investigating reproducibility, robustness, and behavioral consistency across diverse environments, user populations, and operational constraints. This broader perspective is essential because trustworthy AI depends not only on achieving high predictive accuracy but also on producing stable and explainable outputs under varying conditions (Groß, 2024; Gisolfi, 2022).

Formal verification research further highlights that reasoning reliability cannot be inferred solely

from empirical testing. Mathematical verification techniques, robustness analysis, and systematic validation procedures contribute additional evidence regarding system correctness, particularly for safety-critical applications where isolated reasoning failures may produce cascading operational consequences. Consequently, independent evaluation serves as an important mechanism for identifying vulnerabilities that conventional performance benchmarks may overlook (Groß, 2024; Winter *et al.*, 2021).

Sector-Specific Stakes of Insufficient Independent Evaluation

The consequences of inadequate independent evaluation become particularly significant when AI systems operate within sectors characterized by high societal impact and low tolerance for error. Although many AI applications remain advisory rather than fully autonomous, increasing reliance on algorithmic decision support elevates the importance of credible, objective assessment before deployment.

In healthcare, AI systems increasingly support medical imaging, diagnostic assistance, patient risk prediction, treatment recommendations, and clinical workflow optimization. Hallucinated outputs, biased predictions, or inconsistent reasoning may contribute to delayed diagnoses, inappropriate treatments, or unequal healthcare delivery. Independent evaluation provides additional assurance that clinical decision-support systems have been examined under representative operational conditions using transparent and reproducible assessment methodologies (Stettinger *et al.*, 2024; Díaz Rodríguez *et al.*, 2023).

Within legal services, generative AI assists with document drafting, legal research, contract analysis, and case summarization. Fabricated legal authorities, incorrect statutory interpretations, or inconsistent reasoning can compromise professional judgment while creating significant ethical and legal liabilities. Independent evaluation helps determine whether systems maintain sufficient factual reliability and reasoning consistency before being integrated into legal workflows (Jacovi *et al.*, 2021; Kowald *et al.*, 2024).

The financial services sector similarly relies upon AI for fraud detection, credit assessment, portfolio management, anti-money laundering activities, and regulatory compliance. Errors arising from biased predictions or unstable decision-making may produce financial losses, discriminatory lending practices, or regulatory violations. Lessons from model risk management demonstrate that independent validation strengthens institutional governance by providing objective evidence regarding model performance and associated risks, reinforcing the broader argument for external AI evaluation (Brundage *et al.*, 2020; Gisolfi, 2022).

In government and public administration, AI increasingly supports resource allocation, benefit administration, immigration processing, law enforcement analytics, and public service delivery. Decisions affecting citizens require particularly high levels of transparency, accountability, and procedural fairness. Independent evaluation contributes to public trust by ensuring that algorithmic systems undergo impartial assessment rather than relying exclusively on developer-generated evidence of reliability (Díaz Rodríguez *et al.*, 2023; Hohma & Lütge, 2023).

Across these sectors, the common challenge is not merely technical reliability but institutional credibility. Trustworthy AI requires stakeholders to possess confidence that system performance has been evaluated objectively, consistently, and independently of commercial or organizational interests. Existing literature increasingly recognizes that trustworthy AI cannot be achieved solely through ethical principles or internal testing procedures; rather, it depends upon governance structures capable of producing verifiable evidence regarding safety, fairness, robustness, and accountability (Brundage *et al.*, 2020; Zicari *et al.*, 2022; Kowald *et al.*, 2024).

Collectively, these risks demonstrate that the absence of independent evaluation represents a structural governance deficiency rather than a purely technical limitation. Hallucinations, bias amplification, reasoning instability, and sector-specific operational failures each expose weaknesses that may remain undetected under developer-controlled assessment processes. As AI systems continue to assume increasingly influential roles across critical domains, independent evaluation emerges as a necessary mechanism for strengthening credibility, improving accountability, and supporting trustworthy AI deployment without relying exclusively on internal organizational assurances.

The Regulatory and Standards Trajectory

The rapid adoption of artificial intelligence (AI) across high-impact sectors has been accompanied by an equally rapid evolution of governance frameworks intended to promote transparency, accountability, safety, and public trust. Governments, international standards organizations, and research communities have increasingly recognized that trustworthy AI cannot rely solely on technological innovation but must also be supported by robust governance mechanisms capable of identifying, assessing, and mitigating risk throughout the AI lifecycle (Li *et al.*, 2023; Díaz Rodríguez *et al.*, 2023). Consequently, regulatory initiatives and voluntary standards have shifted from broad ethical aspirations toward structured risk management practices that emphasize documentation, transparency, continuous monitoring, and organizational accountability.

Despite these developments, the regulatory landscape remains characterized by an important distinction between governance principles and operational evaluation

methodologies. Existing regulations establish what trustworthy AI should achieve, while technical standards describe processes for managing AI-related risks. However, they generally do not prescribe standardized procedures through which independent evaluators can consistently verify whether AI systems satisfy these expectations. This distinction represents one of the principal governance gaps confronting trustworthy AI deployment.

Evolution from Ethical Principles to Risk-Based Governance

Early discussions surrounding trustworthy AI were largely grounded in ethical principles such as fairness, transparency, accountability, privacy, and human autonomy. Although these principles provided an important conceptual foundation, they offered limited operational guidance for evaluating deployed AI systems. Subsequent research emphasized that ethical aspirations alone are insufficient unless they are translated into measurable, verifiable, and repeatable assessment processes (Thiebes *et al.*, 2021; Jacovi *et al.*, 2021).

More recent scholarship reflects a transition from principle-based governance toward lifecycle-oriented risk management. Rather than focusing exclusively on ethical compliance, contemporary governance approaches encourage organizations to identify potential risks before deployment, monitor system performance continuously, document decision-making processes, and implement corrective measures whenever system behavior deviates from acceptable thresholds (Li *et al.*, 2023). This evolution represents significant progress in AI governance, yet it also reveals that evaluation practices remain largely organization-specific and lack universally accepted procedures for independent verification.

Regulatory Expectations for Trustworthy AI

International regulatory initiatives increasingly converge around several common expectations for trustworthy AI systems. These include comprehensive risk management, transparency, human oversight, technical robustness, lifecycle documentation, and continuous post-deployment monitoring. Rather than prescribing individual technical solutions, these initiatives establish governance obligations intended to reduce risks associated with AI deployment.

The NIST Artificial Intelligence Risk Management Framework (AI RMF) exemplifies this risk-based approach by organizing AI governance around the interconnected functions of Govern, Map, Measure, and Manage. The framework encourages organizations to establish repeatable processes for identifying AI risks, measuring system performance, documenting evaluation evidence, and implementing ongoing governance throughout the system lifecycle. Although the AI RMF strongly emphasizes measurement and documentation, it deliberately remains technology-neutral and does not prescribe standardized independent evaluation procedures.

Similarly, the European Union AI Act introduces legally enforceable obligations for high-risk AI systems, including requirements for risk management, technical documentation, data governance, transparency, human oversight, accuracy, robustness, and post-market monitoring. These obligations significantly strengthen organizational accountability by requiring evidence that AI systems satisfy predefined regulatory expectations. Nevertheless, while the legislation specifies what organizations must demonstrate, it does not establish a universally accepted methodology through which independent evaluators should conduct such assessments.

International standardization efforts further reinforce these governance objectives. ISO/IEC 23894, which focuses on AI risk management, provides structured guidance for identifying, analyzing, evaluating, treating, and monitoring AI-related risks throughout system development and deployment. Like many contemporary standards, its emphasis lies in organizational governance and continuous improvement rather than prescribing standardized evaluation procedures that external assessors should uniformly apply.

Collectively, these initiatives demonstrate increasing international consensus regarding the governance characteristics expected of trustworthy AI. However, they stop short of defining a professional evaluation discipline capable of consistently producing objective, reproducible, and independently verifiable assessments across sectors.

Progress in Practical Trustworthiness Assessment

Alongside regulatory developments, numerous academic studies have sought to operationalize trustworthy AI through structured assessment methodologies. These contributions have substantially advanced understanding

of how AI systems may be evaluated beyond conventional accuracy metrics.

For example, Z-Inspection proposes a multidisciplinary assessment process that integrates ethical, technical, legal, and societal perspectives when evaluating trustworthy AI (Zicari *et al.*, 2021). Subsequent work expanded this practical perspective by identifying procedural considerations that organizations should address when assessing AI systems in operational settings (Zicari *et al.*, 2022). Likewise, verification-oriented research emphasizes the importance of formal reasoning techniques for validating AI system behavior under defined assumptions (Gisolfi, 2022; Groß, 2024).

Other studies have examined trustworthiness from complementary perspectives, including certification readiness, assurance mechanisms, development processes, and organizational governance (Winter *et al.*, 2021; Hohma & Lütge, 2023). Recent investigations also highlight the growing importance of trustworthiness assurance for high-risk AI systems while acknowledging the substantial methodological challenges associated with evaluating increasingly complex AI models (Stettinger *et al.*, 2024; Kowald *et al.*, 2024).

Collectively, this body of literature demonstrates considerable progress toward operationalizing trustworthy AI. However, the proposed methodologies differ significantly in terminology, scope, assessment criteria, documentation practices, and implementation strategies. As a result, organizations adopting different evaluation approaches may arrive at substantially different conclusions regarding the trustworthiness of similar AI systems, limiting comparability and reproducibility across sectors.

Table 2: Comparative Overview of Major AI Governance Instruments

<i>Governance instrument</i>	<i>Primary objective</i>	<i>Key evaluation expectations</i>	<i>Principal limitation for independent evaluation</i>
NIST AI Risk Management Framework (AI RMF)	Risk identification and lifecycle governance	Measurement, documentation, continuous monitoring, risk management	Encourages evaluation but does not prescribe standardized independent assessment procedures
EU AI Act	Regulatory compliance for high-risk AI systems Regulatory compliance for high-risk AI systems	Risk management, technical documentation, transparency, human oversight, post-market monitoring	Specifies compliance obligations but leaves evaluation methodology largely undefined
ISO/IEC 23894	Organizational AI risk management	Systematic identification, assessment, treatment, and monitoring of AI risks	Focuses on governance processes rather than standardized external evaluation
Federal Reserve SR 11-7 (Model Risk Management)	Independent validation of high-risk models	Independent review, model validation, documentation, governance oversight	Designed for financial models rather than general-purpose AI systems
HELM Benchmark	Holistic assessment of language model capabilities	Multi-dimensional benchmark evaluation across diverse scenarios	Measures model capability rather than providing an auditable evaluation methodology suitable for regulatory assurance

The Remaining Governance Gap

The convergence of regulatory frameworks, international standards, and scholarly research demonstrates that the AI community has made substantial progress in defining the characteristics of trustworthy AI. Transparency, accountability, fairness, robustness, documentation, and continuous risk management have emerged as widely accepted governance objectives (Li *et al.*, 2023; Díaz Rodríguez *et al.*, 2023). Nevertheless, a critical gap persists between these high-level governance expectations and the operational processes required to evaluate AI systems consistently and independently.

Current governance instruments establish what trustworthy AI should achieve but provide comparatively limited guidance regarding how independent evaluations should be conducted in a standardized, reproducible, and auditable manner. Similarly, existing research offers valuable assessment methodologies but remains fragmented across disciplines, application domains, and evaluation philosophies. Consequently, AI evaluation continues to be performed using heterogeneous practices that vary considerably between organizations, limiting comparability, external credibility, and regulatory assurance.

This unresolved distinction between governance expectations and operational evaluation methodology represents a foundational challenge for trustworthy AI. As regulatory obligations continue to expand and AI systems assume increasingly consequential roles in society, the absence of a formalized discipline for independent evaluation remains one of the most significant structural limitations within contemporary AI governance. This governance gap provides the basis for examining, in the following section, why AI evaluation remains fragmented and why the emergence of a professional evaluation discipline has become increasingly necessary.

The Absence of a Formal AI Evaluation Discipline

Artificial intelligence has experienced remarkable advances in capability, adoption, and societal impact, yet the mechanisms used to evaluate these systems remain fragmented across organizations, sectors, and application domains. While governments, standards bodies, and researchers have established principles for trustworthy AI, including fairness, transparency, accountability, privacy, robustness, and explainability, the operational process for independently assessing whether AI systems satisfy these principles has not matured into a standardized professional discipline. Consequently, AI evaluation is frequently performed using organization-specific procedures, internally developed testing protocols, or benchmark-driven assessments that vary considerably in scope, rigor, and reproducibility (Li *et al.*, 2023; Kowald *et al.*, 2024).

This fragmentation reflects a broader disconnect between AI governance principles and practical assurance mechanisms. Existing governance frameworks articulate what constitutes trustworthy AI but provide limited guidance regarding how independent evaluators should systematically assess these characteristics across diverse deployment contexts. As organizations increasingly integrate AI into safety-critical and high-impact decision environments, inconsistencies in evaluation methodologies introduce uncertainty regarding the credibility, comparability, and repeatability of reported performance. Without a shared evaluation discipline, assessments often remain difficult to reproduce across institutions, reducing confidence among regulators, enterprises, and affected stakeholders (Díaz Rodríguez *et al.*, 2023; Hohma & Lütge, 2023).

Fragmented and Organization-Specific Evaluation Practices

Current AI evaluation practices are largely determined by the objectives, resources, and governance structures of individual organizations. Technology developers typically design internal testing procedures that emphasize technical performance metrics such as accuracy, precision, robustness, latency, or benchmark scores. While these measures provide valuable insights into model capability, they frequently omit broader governance considerations, including evaluator independence, transparent documentation, standardized evidence collection, and auditability.

Developer-led assessments inevitably introduce structural limitations. Organizations responsible for designing AI systems possess extensive technical knowledge but simultaneously face incentives that may unintentionally influence evaluation priorities, reporting practices, or acceptance thresholds. Such conditions do not necessarily imply misconduct; rather, they highlight an inherent governance challenge common to many regulated industries where independent assessment has evolved to complement internal quality assurance. Similar institutional arrangements have emerged in financial auditing, model validation, and clinical research because independent verification strengthens public confidence by reducing perceived conflicts of interest (Brundage *et al.*, 2020; Jacovi *et al.*, 2021).

The absence of standardized evaluation protocols also limits comparability across AI systems. Two organizations may claim compliance with identical ethical principles while relying on entirely different testing procedures, documentation standards, evidence requirements, and evaluation criteria. Consequently, stakeholders often struggle to determine whether differences in reported trustworthiness reflect genuine variations in system quality or merely differences in evaluation methodology (Thiebes *et al.*, 2021; Li *et al.*, 2023).

Benchmark-Based Evaluation Is Not Equivalent to Independent Evaluation

The widespread adoption of benchmark datasets and standardized performance tests has significantly improved the measurement of AI capabilities. Leaderboards and benchmark suites facilitate comparisons across models by quantifying accuracy, reasoning ability, language understanding, robustness, and other technical characteristics under controlled conditions. Nevertheless, benchmark performance should not be interpreted as evidence of comprehensive trustworthiness.

Benchmarks primarily evaluate model capability rather than governance quality. They seldom assess whether evaluation procedures themselves are objective, independently conducted, reproducible, or sufficiently documented for regulatory scrutiny. Likewise, benchmark scores provide limited evidence regarding deployment risks, organizational governance processes, accountability structures, or ongoing monitoring after implementation. As a result, benchmark-based evaluation represents only one component of a broader assurance process rather than a substitute for independent evaluation (Brundage *et al.*, 2020; Gisolfi, 2022).

Several researchers have similarly distinguished technical performance testing from comprehensive trustworthiness assessment. Practical AI assurance requires evaluating not only algorithmic outputs but also data quality, risk management practices, transparency mechanisms, human oversight, documentation quality, and lifecycle governance. These dimensions extend beyond conventional benchmarking and require systematic methodologies capable of producing credible, repeatable, and externally verifiable evidence (Zicari *et al.*, 2022; Li *et al.*, 2023).

Furthermore, benchmark results often become outdated as models evolve through continual retraining, parameter updates, or deployment in new operational environments. Without standardized procedures governing reevaluation, version control, evidence preservation, and independent reassessment, benchmark scores alone cannot provide sustained assurance throughout an AI system's operational lifecycle (Groß, 2024; Kowald *et al.*, 2024).

Characteristics of a Professional Evaluation Discipline

The evolution of AI governance increasingly suggests the need for evaluation practices that resemble established assurance disciplines found in other high-consequence domains. Although no universally accepted discipline currently exists, the literature consistently identifies several foundational characteristics that would distinguish a mature evaluation profession from today's fragmented approaches.

First, evaluation criteria should be sufficiently standardized to enable meaningful comparisons across

organizations while remaining adaptable to different deployment contexts. Consistent assessment criteria improve transparency, facilitate reproducibility, and reduce ambiguity in interpreting evaluation outcomes (Zicari *et al.*, 2021; Li *et al.*, 2023).

Second, evaluator independence constitutes a defining requirement for credible assessment. Independent evaluators are better positioned to minimize conflicts of interest, strengthen objectivity, and increase stakeholder confidence in reported findings. Independence enhances not only technical credibility but also institutional legitimacy by separating assessment from system development (Jacovi *et al.*, 2021; Winter *et al.*, 2021).

Third, reproducibility should extend beyond experimental replication to encompass standardized documentation of datasets, evaluation procedures, testing environments, assumptions, and decision criteria. Comprehensive documentation enables subsequent reviewers to verify conclusions, compare assessments across institutions, and identify methodological inconsistencies (Gisolfi, 2022; Groß, 2024).

Fourth, transparent audit trails should accompany evaluation activities throughout the AI lifecycle. Such documentation should record evidence sources, testing protocols, evaluation outcomes, identified limitations, corrective actions, and governance decisions. Transparent records improve accountability while supporting regulatory oversight and organizational learning (Stettinger *et al.*, 2024; Zicari *et al.*, 2022).

Finally, professional competency represents an essential yet often overlooked dimension of trustworthy evaluation. Effective assessment requires expertise extending beyond machine learning to encompass ethics, governance, software engineering, statistics, risk management, domain knowledge, and regulatory compliance. Developing consistent evaluation capabilities therefore depends not only on technical tools but also on qualified evaluators capable of applying standardized methodologies across diverse operational settings (Hohma & Lütge, 2023; Kowald *et al.*, 2024).

Collectively, these characteristics illustrate that trustworthy AI evaluation involves considerably more than measuring predictive performance. It requires structured governance processes capable of producing evidence that is technically rigorous, transparent, reproducible, and independently credible.

The Remaining Discipline Gap

Despite substantial progress in AI governance research, an identifiable gap remains between the existence of trustworthy AI principles and the institutional mechanisms required to evaluate those principles consistently. Current research increasingly addresses verification, formal assurance, certification concepts, trustworthy development processes, and operational governance, reflecting growing recognition that technical

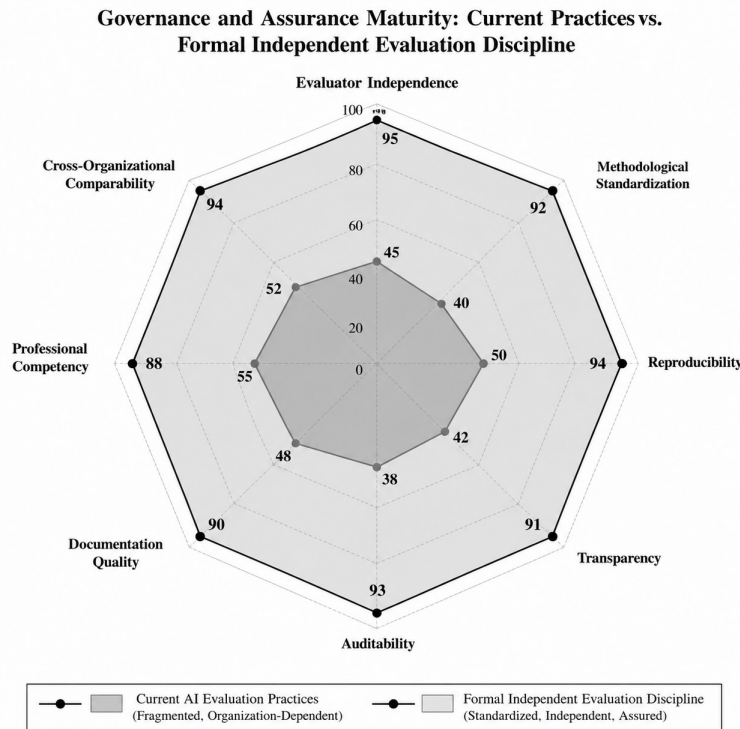


Figure 2: Conceptual Maturity Gap Between Current AI Evaluation Practices and a Formal Independent Evaluation Discipline

performance alone cannot establish trustworthy AI (Winter *et al.*, 2021; Zicari *et al.*, 2022; Stettinger *et al.*, 2024). Nevertheless, these contributions collectively illustrate an evolving field rather than a consolidated professional discipline.

The absence of standardized evaluation methodologies, universally accepted competency expectations, reproducible assessment procedures, and independent governance structures continues to limit confidence in AI evaluation outcomes. Organizations therefore remain responsible for interpreting broad governance principles using internally developed methodologies that differ substantially in rigor and documentation. This variability complicates procurement decisions, regulatory oversight, cross-sector comparisons, and stakeholder trust because evaluation results cannot always be interpreted using consistent standards (Thiebes *et al.*, 2021; Díaz Rodríguez *et al.*, 2023).

As AI systems continue to influence decisions with significant societal, legal, financial, and ethical consequences, the challenge extends beyond improving model performance. The more fundamental issue concerns the absence of a recognized evaluation discipline capable of producing objective, reproducible, transparent, and independently credible assessments across diverse application domains. Addressing this governance gap represents one of the most significant unresolved challenges in the maturation of trustworthy AI.

Implications for Practice

The growing reliance on artificial intelligence (AI) in high-impact domains has elevated the importance of evaluation from a technical activity to a strategic governance function. While contemporary AI governance frameworks emphasize principles such as transparency, accountability, and risk management, the practical implementation of these principles remains inconsistent across organizations. As demonstrated throughout this paper, the absence of an independent and standardized approach to AI evaluation limits the credibility of performance claims and weakens stakeholder confidence in AI-enabled decision-making. Establishing evaluation as a structured and independent practice has significant implications for enterprises, regulators, and the broader AI ecosystem.

Implications for Enterprises

For enterprises, independent AI evaluation has implications that extend well beyond technical validation. Organizations increasingly deploy AI systems to support operational efficiency, customer engagement, financial decision-making, healthcare delivery, cybersecurity, and strategic planning. In these contexts, evaluation is not solely concerned with measuring predictive accuracy or benchmark performance but with determining whether systems consistently operate within acceptable levels of reliability, fairness, robustness, and transparency throughout their life cycle. Existing enterprise practices

often rely heavily on developer-led testing, which, although valuable for quality assurance, may be influenced by organizational incentives, resource limitations, or methodological inconsistencies. Independent evaluation introduces an additional layer of credibility by providing objective evidence that AI systems satisfy predefined assessment criteria before deployment and during operational use (Brundage *et al.*, 2020; Jacovi *et al.*, 2021).

A structured evaluation process also strengthens organizational governance by improving risk visibility and supporting evidence-based decision-making. Enterprise leaders increasingly require verifiable information regarding model behavior under diverse operating conditions, including rare events, adversarial scenarios, and evolving data distributions. Independent assessments can provide reproducible documentation of system performance, identify previously overlooked vulnerabilities, and improve confidence in procurement, vendor selection, and deployment decisions. Rather than replacing internal testing procedures, independent evaluation complements existing quality assurance activities by introducing external scrutiny and standardized documentation that reduce uncertainty surrounding AI performance claims (Gisolfi, 2022; Zicari *et al.*, 2022).

The implications are particularly significant for organizations operating in regulated or safety-critical sectors. In healthcare, finance, transportation, and critical infrastructure, AI errors may result in substantial legal, financial, or societal consequences. Independent evaluation therefore contributes to organizational resilience by strengthening governance mechanisms that demonstrate responsible AI deployment to boards, regulators, investors, customers, and other stakeholders. Research on trustworthy AI consistently emphasizes that trust emerges not merely from technical capability but from transparent and accountable development and evaluation processes that enable informed oversight throughout the AI lifecycle (Thiebes *et al.*, 2021; Hohma & Lütge, 2023; Li *et al.*, 2023).

Furthermore, standardized independent evaluation may improve procurement practices across organizations adopting third-party AI solutions. Enterprises increasingly acquire foundation models and AI-enabled services from external vendors, often relying on vendor-provided documentation to assess quality and risk. Independent evaluation can strengthen procurement by supplying objective evidence regarding system reliability, robustness, fairness, and operational limitations using transparent and reproducible methodologies rather than relying exclusively on self-reported performance metrics. Such practices reduce information asymmetry between AI developers and adopters while promoting greater organizational confidence in deployment decisions (Winter *et al.*, 2021; Kowald *et al.*, 2024).

Implications for Regulators

The implications of independent evaluation extend equally to regulatory oversight. Governments worldwide have increasingly adopted risk-based approaches to AI governance, emphasizing accountability, transparency, documentation, and continuous risk management. However, despite these advances, regulatory frameworks generally establish what organizations should demonstrate rather than prescribing how evaluation should be independently conducted. This distinction creates practical challenges for supervisory authorities responsible for determining whether AI systems satisfy legal and ethical expectations across different operational environments.

Independent evaluation offers regulators a credible mechanism for translating governance principles into verifiable evidence. Rather than depending primarily on developer-produced documentation, regulatory oversight could benefit from assessment processes that generate reproducible records of system performance, evaluation procedures, and identified limitations. Such documentation supports more consistent compliance reviews while reducing ambiguity surrounding claims of robustness, fairness, transparency, and safety. The objective is not to replace regulatory judgment but to strengthen it through evaluation practices that are transparent, evidence-based, and sufficiently standardized to support comparisons across organizations and application domains (Díaz Rodríguez *et al.*, 2023; Zicari *et al.*, 2021).

Independent evaluation also facilitates greater consistency in regulatory enforcement. AI systems deployed within similar risk categories are frequently assessed using different internal methodologies, making cross-organizational comparisons difficult. Standardized evaluation practices would improve comparability by encouraging common assessment principles, reproducible testing procedures, and well-documented evidence. Such consistency enables regulators to evaluate AI deployments more effectively while preserving flexibility for different technical implementations and application contexts (Stettinger *et al.*, 2024; Kowald *et al.*, 2024).

Another important implication concerns public accountability. Regulatory legitimacy depends not only on establishing legal obligations but also on demonstrating that compliance assessments are objective, transparent, and technically credible. Independent evaluation contributes to this objective by separating assessment activities from development interests, thereby strengthening confidence in regulatory oversight and improving public trust in AI governance mechanisms. This separation mirrors long-established approaches in other regulated disciplines, where independent verification enhances confidence in safety, reliability, and compliance without dictating technological innovation (Brundage *et al.*, 2020; Jacovi *et al.*, 2021).

Implications for the AI Industry

At the industry level, independent evaluation has broader implications for the evolution of AI governance and professional practice. The AI ecosystem has experienced rapid technological advancement accompanied by increasing public scrutiny regarding model reliability, bias, misinformation, explainability, and accountability. Although organizations have invested heavily in AI safety, testing, and governance initiatives, evaluation practices remain fragmented, with considerable variation in methodologies, reporting standards, and evidence requirements. This fragmentation limits comparability across systems and complicates efforts to establish shared expectations regarding trustworthy AI deployment (Li *et al.*, 2023; Kowald *et al.*, 2024).

The emergence of a professional evaluation discipline would contribute to greater methodological consistency throughout the industry. Rather than focusing exclusively on benchmark performance or isolated technical metrics, evaluation could evolve toward comprehensive assessment practices characterized by reproducibility, transparent documentation, standardized reporting, and independent oversight. Such developments would encourage greater confidence among developers, enterprises, policymakers, and the public while improving the overall quality and credibility of AI assurance activities. Importantly, this perspective does not advocate a particular framework or certification model; instead, it recognizes that mature technological ecosystems typically rely on structured evaluation professions to provide objective evidence supporting public trust (Groß, 2024; Gisolfi, 2022).

Professionalization also has implications for workforce development and interdisciplinary collaboration. Effective AI evaluation increasingly requires expertise that extends beyond machine learning performance metrics to include software engineering, cybersecurity, ethics, human factors, governance, legal compliance, and domain-specific risk assessment. A structured evaluation discipline would therefore encourage the development of shared competencies, consistent assessment methodologies, and common documentation practices capable of supporting increasingly complex AI deployments across diverse sectors (Hohma & Lütge, 2023; Díaz Rodríguez *et al.*, 2023).

Finally, independent evaluation has the potential to strengthen the relationship between technological innovation and societal trust. Public confidence in AI depends not only on the capabilities of AI systems but also on the credibility of the processes used to assess them. Transparent, reproducible, and independently conducted evaluations provide stakeholders with evidence that AI systems have undergone objective scrutiny before influencing decisions with significant societal consequences. As AI continues to expand across critical domains, the credibility of evaluation processes will become increasingly important for sustaining innovation

while maintaining accountability, transparency, and public confidence in AI-enabled systems (Thiebes *et al.*, 2021; Zicari *et al.*, 2022; Stettinger *et al.*, 2024).

Collectively, these implications demonstrate that independent evaluation should be understood as a foundational governance capability rather than a supplementary quality assurance activity. For enterprises, it enhances organizational decision-making, procurement confidence, and risk management. For regulators, it strengthens compliance assessment, accountability, and consistency in oversight. For the AI industry, it supports methodological maturity, professionalization, and greater public trust. Together, these practical implications reinforce the central argument of this paper that credible AI governance depends not only on the existence of evaluation activities but also on their independence, transparency, reproducibility, and institutional legitimacy.

CONCLUSION

The accelerated adoption of artificial intelligence across high-impact domains has fundamentally transformed how organizations make decisions, automate processes, and deliver services. However, the findings of this paper demonstrate that the rapid pace of deployment has not been matched by a corresponding evolution in the maturity of AI evaluation practices. While considerable progress has been made in developing principles for trustworthy AI, existing approaches remain largely dependent on developer-led assessments, fragmented organizational practices, and benchmark-driven performance measurements that, although valuable, do not consistently provide the level of independence, reproducibility, or accountability required for high-stakes deployment. As AI systems continue to influence decisions with significant societal, economic, and legal implications, the credibility of their evaluation becomes as important as their technical capability.

This paper has argued that independent evaluation should be regarded as a structural component of trustworthy AI governance rather than an optional extension of quality assurance. Lessons drawn from established disciplines including financial auditing, model risk management, and safety-critical engineering demonstrate that public confidence in complex systems depends not only on technical performance but also on the independence and transparency of the processes used to assess them. Similar observations have emerged within the AI literature, where researchers have emphasized the importance of verifiable claims, systematic assessment methodologies, and formal mechanisms capable of supporting trustworthy AI development (Brundage *et al.*, 2020; Jacovi *et al.*, 2021). These perspectives reinforce the position that trustworthy AI cannot rely solely on self-attestation but requires evaluation processes that are demonstrably objective, consistent, and externally credible.

The analysis further showed that contemporary regulatory and governance initiatives have established an important foundation for responsible AI through principles of fairness, transparency, accountability, robustness, and risk management. Nevertheless, these initiatives primarily define governance objectives rather than a universally accepted operational methodology for conducting independent evaluations across organizations and application domains. Consequently, substantial variation persists in how AI systems are assessed, documented, validated, and reported, limiting the comparability and reproducibility of evaluation outcomes. This fragmentation has been widely recognized as a continuing challenge in translating trustworthy AI principles into consistent organizational practice (Thiebes *et al.*, 2021; Li *et al.*, 2023; Díaz Rodríguez *et al.*, 2023; Hohma & Lütge, 2023).

A central contribution of this study has been to distinguish between evaluating AI models and establishing a formal evaluation discipline. Existing verification methods, assessment processes, benchmarking initiatives, and trustworthiness frameworks have significantly advanced the understanding of AI quality and governance (Gisolfi, 2022; Zicari *et al.*, 2021; Zicari *et al.*, 2022). Likewise, ongoing research into formal verification, certification methodologies, and trustworthiness assurance demonstrates increasing recognition that AI systems require more rigorous mechanisms for validation, particularly in high-risk environments (Winter *et al.*, 2021; Groß, 2024; Stettinger *et al.*, 2024). However, these contributions collectively illustrate an evolving landscape rather than the existence of a mature, standardized professional discipline capable of delivering independent, repeatable, and auditable evaluations across diverse organizational contexts.

The discussion has also highlighted that credible AI evaluation extends beyond measuring predictive accuracy or benchmark performance. It requires organizational independence, transparent evaluation criteria, reproducible methodologies, comprehensive documentation, and evidence that can withstand technical, regulatory, and public scrutiny. These characteristics are essential for enabling stakeholders, including enterprises, regulators, procurement authorities, and end users, to develop justified confidence in AI systems whose decisions increasingly affect critical aspects of society. Recent scholarship similarly identifies the need to strengthen institutional mechanisms, governance structures, and interdisciplinary evaluation approaches capable of supporting trustworthy AI throughout the system lifecycle (Kowald *et al.*, 2024; Li *et al.*, 2023).

Ultimately, the principal challenge confronting AI governance is not the absence of evaluation activity but the absence of a recognized and standardized discipline dedicated to independent AI evaluation. The credibility of AI systems depends not merely on whether evaluations

are performed but on who performs them, how they are conducted, whether they can be independently reproduced, and whether the resulting evidence can support informed governance and accountability. Until these structural characteristics become consistently embedded within AI evaluation practices, important gaps will remain between the aspirations of trustworthy AI and the practical mechanisms required to demonstrate it. The emergence of increasingly sophisticated AI systems therefore underscores the continuing need to recognize independent evaluation as a foundational element of trustworthy AI governance, while acknowledging that the development of a comprehensive and universally accepted evaluation discipline remains an important unresolved challenge.

REFERENCES

1. Gisolfi, N. (2022). *Model-Centric Verification of Artificial Intelligence* (Doctoral dissertation, Carnegie Mellon University, USA).
2. Brundage, M., Avin, S., Wang, J., Belfield, H., Krueger, G., Hadfield, G., ... & Anderljung, M. (2020). Toward trustworthy AI development: mechanisms for supporting verifiable claims. *arXiv*, 2020, 2004-07213.
3. Mazumder, P. T. (2025). AI-Driven Anti-Money Laundering Systems for Cybersecurity Resilience in US Financial Infrastructure: A Framework for Real-Time Threat Detection, Regulatory Compliance and National Security. *International Journal of Humanities and Information Technology*, 7(03), 90-97.
4. Zicari, R. V., Brodersen, J., Brusseau, J., Dudder, B., Eichhorn, T., Ivanov, T., ... & Westerlund, M. (2021). Z-Inspection@: a process to assess trustworthy AI. *IEEE Transactions on Technology and Society*, 2(2), 83-97.
5. Jacovi, A., Marasović, A., Miller, T., & Goldberg, Y. (2021, March). Formalizing trust in artificial intelligence: Prerequisites, causes and goals of human trust in AI. In *Proceedings of the 2021 ACM conference on fairness, accountability, and transparency* (pp. 624-635).
6. Mazumder, P. T. (2025). Harnessing fintech innovations for anti-money laundering: a data-driven approach. Available at SSRN 5259084.
7. Groß, D. M. (2024). *Towards trustworthy ai: Formal verification in machine learning* (Doctoral dissertation, SI: sn).
8. Winter, P. M., Eder, S., Weissenböck, J., Schwald, C., Doms, T., Vogt, T., ... & Nessler, B. (2021). Trusted artificial intelligence: Towards certification of machine learning applications. *arXiv preprint arXiv:2103.16910*.
9. Mazumder, P. T. (2023). Data Driven Detection of Trade Based Money Laundering (TBML): A predictive Analytics Framework for Securing US supply chains and Financial Integrity. *SAMRIDDI: A Journal of Physical Sciences, Engineering and Technology*, 15(04), 423-432.
10. Zicari, R. V., Amann, J., Bruneault, F., Coffee, M., Dudder, B., Hickman, E., ... & Wurth, R. (2022). How to assess trustworthy AI in practice. *arXiv preprint arXiv:2206.09887*.
11. Stettinger, G., Weissensteiner, P., & Khastgir, S. (2024). Trustworthiness assurance assessment for high-risk AI-based systems. *IEEE Access*, 12, 22718-22745.
12. Díaz Rodríguez, N., Del Ser Lorente, J., Coeckelbergh, M., López de Prado, M., Herrera Viedma, E., & Herrera Triguero, F. (2023). Connecting the dots in trustworthy Artificial Intelligence: From AI principles, ethics, and key requirements to responsible AI systems and regulation. *Information Fusion*, 99.
13. Kowald, D., Scher, S., Pammer-Schindler, V., Müllner, P., Waxnegger, K., Demelius, L., ... & Kopeinik, S. (2024). Establishing and evaluating trustworthy AI: overview and research challenges. *Frontiers in Big Data*, 7, 1467222.
14. Mazumder, P. T. (2025). Explainable Machine Learning Pipelines for

- Customer Risk Scoring in Anti-Money Laundering: A Management and Governance Perspective. *Journal of Data Analysis and Critical Management*, 1(02), 79-90.
15. Li, B., Qi, P., Liu, B., Di, S., Liu, J., Pei, J., ... & Zhou, B. (2023). Trustworthy AI: From principles to practices. *ACM Computing Surveys*, 55(9), 1-46.
 16. Hohma, E., & Lütge, C. (2023). From trustworthy principles to a trustworthy development process: The need and elements of trusted development of AI systems. *AI*, 4(4), 904-925.
 17. Thiebes, S., Lins, S., & Sunyaev, A. (2021). Trustworthy artificial intelligence: S. Thiebes *et al. Electronic Markets*, 31(2), 447-464.
 18. Thakkar, V. B. (2025). Defining Success in Probabilistic Products: Key Performance Indicators and Lifecycle Management for Generative AI Applications in Enterprise. *International Journal of Technology, Management and Humanities*, 11(04), 66-79.
 19. Suseelan, S. (2025). Big Data Analytics in Air Traffic Management: Enhancing Aviation Safety, Operational Efficiency, and Predictive Decision-Making. *Journal of Science Technology and Social Transformation*, 1(02), 61-73.
 20. Thakkar, V. B. (2025). Algorithmic Trust, Governance, and Integration Bottlenecks of AI Tools in Legacy Project Environments. *International Journal of Humanities and Information Technology*, 7(01), 80-89.
 21. Suseelan, S. (2023). Explainable AI (XAI) for FAA Certifiable Aviation Systems. *SAMRIDDHI: A Journal of Physical Sciences, Engineering and Technology*, 15(04), 441-451.
 22. Thakkar, V. B. (2026). From linear logistics to neural supply chains: Predictive machine learning and the rise of autonomous supply chain intelligence. *International Journal of AI, BigData, Computational and Management Studies*, 7(1), 153-161.

HOW TO CITE THIS ARTICLE: Kubvoruno, I. (2025). Independent Verification as a Precondition for Trustworthy AI Deployment: The Case for a Formal AI Evaluation Discipline. *Journal of Integrated Science, Technology and Management*, 1(1), 15-30.