



FUNDAMENTALS OF ARTIFICIAL INTELLIGENCE RISK MANAGEMENT FOR MANAGERS

Haris Hamidović 

Microcredit Foundation EKI and Microcredit Company EKI, Sarajevo, Bosnia and Herzegovina
<https://orcid.org/0000-0002-1296-5008>

ARTICLE INFO



 Open access

JEL Category:
K22, M15

Keywords:

artificial intelligence
AI risk
AI governance
risk management
personal data protection
compliance
accountability

ABSTRACT

Artificial Intelligence (AI) is increasingly becoming an integral part of business processes, products, and services, providing opportunities for automation, data analysis, prediction, content generation, and decision-making support. At the same time, AI introduces specific risks related to data quality and bias, reliability, explainability, privacy, security, accountability, human oversight, vendor management, and changes throughout the AI lifecycle. This paper examines AI risk management from a management perspective, drawing on established approaches to AI risk management and governance. Particular attention is given to integrating AI risks into existing organizational risk management, security, privacy, vendor management, and compliance processes, as well as to the regulatory context relevant to organizations in Bosnia and Herzegovina. The main contribution of the paper is a management-oriented AI risk management framework integrating AI-specific risks with existing governance and risk management processes across seven dimensions: business value, risk, data, accountability, vendors, compliance, and continuous monitoring. The framework is intended to support management in assessing and governing AI use cases throughout their lifecycle.

1 INTRODUCTION

Artificial intelligence is no longer exclusively a subject for technology and research teams. Generative AI, large language models, predictive analytics systems, and increasingly autonomous AI systems are already becoming part of business processes, customer services, document management, software development, security

operations, financial decision-making, and other areas of business.

For management, the question of AI can therefore no longer be reduced to selecting the appropriate technology. The key question is how to realize the business value of AI while keeping the associated risks within the boundaries that the organization is willing to accept.

Address of the author:
Haris Hamidović
 haris.hamidovic@eki.ba

Received: 09.08.2026
Revised: 20.08.2026
Accepted: 21.08.2026
Available online: 21.08.2026



This approach is particularly important because AI does not only introduce entirely new risks. To a significant extent, it changes how existing risks manifest. In AI systems, the traditional risk of data quality can translate into the risk of inaccurate predictions or discriminatory outcomes. Vendor risk may involve not only service availability and security, but also how a vendor uses data, trains and modifies models, engages subcontractors, and allocates responsibilities. Change risk may include changes in the behavior of an AI model following a new version, even when the formal business functionality of the system has not changed.

For this reason, AI should not be introduced as a separate process outside the organization's existing governance and management systems. Existing processes for risk management, information security, privacy, vendor management, change management, and business continuity should provide the foundation upon which AI-specific controls are built.

The objective, therefore, should not be to create a completely separate AI risk management system, but rather to appropriately extend the existing governance and risk management framework wherever AI technologies introduce new or different risks.

Before approving the implementation of an AI system, management should therefore ask three fundamental questions:

- What business problem are we trying to solve?
- What business value do we expect to achieve?
- What level of risk are we willing to accept to realize that value?

These three questions form the foundation of a responsible approach to AI.

1.1 Research objective and contribution

The objective of this paper is to develop a management-oriented conceptual framework for AI risk management that can be integrated into existing organizational governance and enterprise risk management processes. Rather than proposing a technically oriented AI model, the paper focuses on the managerial mechanisms

required to govern AI-related risks throughout the AI system lifecycle. The contribution of the paper is an integrated seven-dimensional management framework comprising business value, risk, data, accountability, vendors, compliance, and continuous monitoring. The framework synthesizes concepts from established AI risk management and governance approaches and connects them with organizational responsibilities and regulatory considerations relevant to organizations operating in Bosnia and Herzegovina. The proposed framework is intended to provide managers with a practical structure for assessing whether an AI use case is sufficiently justified, governed, controlled, and monitored, while avoiding the creation of a separate AI risk management system outside existing organizational governance structures.

1.2 Literature Review and Conceptual Basis

The growing adoption of artificial intelligence has resulted in an increasing body of research addressing AI-related risks, governance, security, fairness, accountability, and risk measurement. Existing approaches differ in scope and purpose, ranging from general AI risk management frameworks to more specific approaches addressing particular categories of AI risk.

Research on AI risk measurement emphasizes the need to assess AI-related risks systematically rather than treating AI as a purely technological issue. Giudici et al. (2024), for example, examine approaches to artificial intelligence risk measurement, supporting the need to integrate AI risk assessment into broader organizational risk management processes. Similarly, the FAIR-AIR approach applies established risk analysis principles to AI adoption and provides a link between AI-specific risks and enterprise risk management (FAIR Institute, 2024).

Other sources focus on identifying and structuring specific AI risks. The MIT AI Risk Repository provides a structured overview of AI-related risks, while MITRE ATLAS and OWASP resources address security threats associated with AI and machine learning systems, including adversarial attacks and risks specific to large language model applications (MIT AI Risk Initiative, 2024; OWASP, 2025; MITRE, 2023). These approaches

complement broader AI risk management frameworks by providing more detailed perspectives on technical and security-related risks.

AI governance frameworks and professional guidance further emphasize accountability, transparency, human oversight, privacy, fairness, and continuous monitoring as essential elements of responsible AI management. The NIST AI Risk Management Framework and the ISACA Advanced in AI Risk (AAIR) approach provide structured guidance for incorporating these considerations into organizational governance and risk management (U.S. Department of Commerce, 2023; ISACA, 2025a). The EU AI Act further reinforces a risk-based approach by linking the level of regulatory and governance requirements to the potential impact of AI systems. (European Parliament & Council of the European Union, 2024)

As highlighted in the ISACA Advanced in AI Security Management (AAISM) guidance, the rapidly evolving nature of artificial intelligence and emerging technologies requires focused and tailored frameworks to support effective management. Regulatory and legislative bodies, standards organizations, professional associations, and research institutions are increasingly developing frameworks and guidance to keep pace with the rapidly evolving field of AI. These include the European Parliament and the Council of the European Union, the FAIR Institute, the Institute of Electrical and Electronics Engineers (IEEE), the Organisation for Economic Co-operation and Development (OECD), the Digital Ethics Lab of the Oxford Internet Institute (OII), ISO/IEC, UNESCO, and NIST (ISACA, 2025b; Oxford Internet Institute, 2017; ISO/IEC, 2023a; ISO/IEC, 2023b, UNESCO, 2021). As these AI-related frameworks continue to evolve, organizations need to review them regularly to remain aligned with current guidance, standards, and regulatory requirements. The distinction between regulatory and voluntary frameworks is particularly relevant in this context. Regulatory instruments, such as the EU AI Act, establish legally binding obligations, whereas voluntary frameworks, such as the NIST AI Risk Management Framework, provide structured

guidance and good practices for identifying and mitigating AI-related risks.

Although these frameworks and approaches provide important and complementary elements of AI risk management, they differ in scope, purpose, and level of abstraction. Some primarily support risk measurement and analysis, others focus on technical and security threats, while regulatory and governance frameworks emphasize compliance, accountability, and organizational oversight. Consequently, organizations may need to draw on multiple sources when establishing an integrated approach to AI risk management and decision-making. Building on these complementary perspectives, this paper proposes an integrative management-oriented framework that connects AI-specific risks with existing organizational governance and risk management processes across seven dimensions: business value, risk, data, accountability, vendors, compliance, and continuous monitoring.

2 FUNDAMENTAL CONCEPTS OF ARTIFICIAL INTELLIGENCE

Artificial intelligence is an umbrella term encompassing computer systems capable of performing tasks associated with human capabilities, such as analysis, pattern recognition, prediction, reasoning, classification, content generation, and decision-making support.

AI is not a single technology or model. It encompasses a range of approaches, from systems based on rules and algorithms to complex systems that learn from data and can operate with varying degrees of autonomy.

One important category is Machine Learning (ML), in which a system uses data to identify patterns and develop models that can subsequently be used for prediction, classification, generation, or other tasks. Consequently, the quality, relevance, completeness, and integrity of data represent key considerations in AI risk management.

Deep Learning (DL) is a subfield of machine learning based on multilayer neural networks and has become a fundamental technological basis for many contemporary AI applications. Generative AI (GenAI), in turn, refers to AI systems capable of generating new content, including text, images, software code, audio, and other types of content.

Contemporary GenAI systems are predominantly based on deep learning techniques, but GenAI should not be understood as a subsequent hierarchical level of AI beyond deep learning. Rather, it represents a class of AI applications defined primarily by their ability to generate content.

A further development is Agentic AI, referring to AI systems or architectures capable of operating with a greater degree of autonomy, including planning and executing tasks, using external tools, maintaining context, and interacting with other systems. Agentic AI therefore represents primarily a dimension of system autonomy and operation rather than a separate hierarchical level within AI.

For risk management purposes, the relationship among these concepts can therefore be viewed as overlapping rather than strictly hierarchical: AI represents the broader field; machine learning and deep learning represent important approaches for developing AI systems; GenAI represents a class of AI systems focused on content generation, often enabled by deep learning; while Agentic AI describes systems that employ AI capabilities with an increased degree of autonomy.

For management, however, understanding several fundamental questions is more important than the classification itself:

- Who controls the AI system?
- Who is accountable for its outputs?
- What data is the system based on?
- Can its outputs be verified?
- Can the organization reconstruct what happened?
- What happens when the system makes a mistake?

The question of accountability goes beyond just the technical correctness of the system; it requires setting up ethical standards that ensure human oversight remains meaningful. As Scarpino (2025) points out, striving for responsible AI use is inseparable from clearly defining who ultimately bears responsibility for decisions made based on AI model outputs.

3 AI AS A DISTINCT RISK

AI can generate significant business value through automation, the processing of large volumes of

data, pattern recognition, predictive analytics, anomaly detection, personalization, content generation, and decision-making support.

However, the same capabilities that create value can also create risk.

Automation can increase efficiency, but it can also automate a poorly designed or incorrectly defined process. Predictive models can improve decision-making, but their predictions can be wrong. Processing large volumes of data enables analysis that would be difficult for humans to perform at the same scale, but the results may be unreliable if the underlying data is inaccurate, incomplete, or biased.

Generative AI can significantly increase productivity, but it can also produce inaccurate or inappropriate content. Autonomous task execution can further increase efficiency while simultaneously increasing the risk of losing control over the process.

Therefore, for AI systems, it is not sufficient to demonstrate that the technology works. It is necessary to determine whether the output is sufficiently reliable for the specific business purpose and the level of risk the organization is willing to accept.

The level of accuracy that may be acceptable for generating marketing content recommendations may not be acceptable for an AI system that supports credit decisions, candidate selection, or another process that may significantly affect an individual.

Among the most important AI-specific risks are the following:

- **Data Quality.** If data is inaccurate, incomplete, outdated, biased, or inadequately labeled, these issues can directly affect the outputs of the model.
- **Bias and Fairness.** AI can reproduce or amplify existing patterns of bias, potentially resulting in different treatment of individuals or groups.
- **Explainability.** With more complex models, it can be difficult to explain why a system produced a particular output. This becomes particularly important when an output affects an individual or a significant business decision.

- Reliability and Accuracy. An AI output can be convincing but inaccurate. Therefore, a convincing output should not be mistaken for a validated output.
- Model Drift. Model performance can change as a result of changes in data, user behavior, market conditions, or the broader business environment.
- Security. AI systems can introduce additional attack surfaces and create risks related to misuse, data manipulation, and unauthorized access.
- Privacy. The processing of personal data through AI can create additional risks to the rights and freedoms of individuals.
- Intellectual Property. Organizations need to consider the rights associated with data entered into AI systems, how such data is used, and the rights associated with AI-generated outputs.
- Accountability. When an AI system provides a recommendation or performs an activity, it must be clear who is accountable for the final decision and outcome.

The greater the potential impact of an AI system, the higher the level of governance and control that should be applied.

4 AI GOVERNANCE

AI governance frameworks provide organizations with a structured approach to managing risks arising from the development, procurement, and use of AI systems.

Their purpose is not to create another parallel governance system, but rather to address key questions such as:

- what needs to be controlled;
- who is accountable;
- when risk should be assessed;
- how the AI system should be tested;
- what should be monitored after implementation;
- how incidents should be managed;
- how compliance can be demonstrated.

Across different AI governance approaches, similar principles recur, including accountability, transparency, explainability, fairness, privacy, security, reliability, human oversight, data governance, and continuous monitoring. (U.S.

Department of Commerce, 2023; European Parliament & Council of the European Union, 2024)

These principles form the foundation of what is increasingly referred to as AI trustworthiness.

For management, it is therefore more important to understand the purpose of these principles than to adopt a particular framework as an objective in itself. A framework should help the organization strengthen its existing governance and management systems wherever the specific characteristics of AI technologies introduce new or different risks.

5 ACCOUNTABILITY FOR AI DECISIONS AND OUTCOMES

One of the most important questions in AI governance is not simply: “Who developed the AI?”

but rather: “Who is accountable for the decision or business outcome based on the AI output?”

This question becomes particularly important when an organization uses an AI system provided by an external vendor.

The fact that a vendor developed the model does not mean that the vendor has automatically assumed responsibility for how the organization uses that system.

The organization should clearly define:

- the AI use case owner;
- the business process owner;
- the owner of the associated risk;
- the person or function responsible for approving the use of the system;
- responsibility for monitoring;
- responsibility for incident management;
- the authority to restrict or shut down the system.

The AAIR framework specifically emphasizes AI ownership, oversight, and accountability as key elements of AI risk governance.

Therefore, a vendor may be accountable for contracted services, functionalities, or specific failures within the scope of the contractual relationship, but the organization cannot simply transfer its own accountability for the decision to use a particular AI system, the manner in which it

is used, and the management of the risks arising from its use.

When an organization uses AI as a service, traditional vendor management needs to be expanded.

In addition to availability, service-level agreements (SLAs), security, and business continuity, the organization needs to understand:

- what data the vendor processes;
- where the data is processed and stored;
- whether the data is used to train or improve models;
- how long the data is retained;
- who can access it;
- which subcontractors are involved in delivering the service;
- how model changes are managed;
- how new versions are tested and validated;
- how AI incidents are reported and handled;
- how bias and model drift are managed;
- how can the organization conduct audits or obtain appropriate evidence.

A particularly important issue is exiting management.

The organization must understand what happens if it decides to discontinue the use of an AI solution. Can it retrieve its data? Can the data be deleted? Can the organization migrate to another solution? How will business continuity be maintained?

For AI systems, model change management is also particularly important. A new model version can change the behavior of the system even when its core functionality has not formally changed.

Therefore, where applicable, contracts with AI vendors should address the obligation to notify the organization of significant changes, the organization's ability to test and revalidate new versions, the availability of relevant documentation, and, where technically feasible, the ability to roll back to a previous version or implement an alternative solution.

Vendor management therefore becomes an integral part of AI risk management and AI system lifecycle management.

6 PEOPLE AND ORGANIZATIONAL CAPABILITIES

AI does not only change technology. It also changes the requirements placed on people.

AI risk management should therefore address:

- AI awareness;
- employee training;
- development of new competencies;
- changes in job roles;
- human oversight;
- employee accountability;
- the risk of overreliance on AI.

A particular risk is the loss of critical human judgment.

If an employee begins to automatically accept an AI output simply because it was produced by an "intelligent" system, human oversight becomes merely a formality.

7 REGULATORY AND PRIVACY CONSIDERATIONS FOR AI IN BOSNIA AND HERZEGOVINA

For organizations in Bosnia and Herzegovina, it is particularly important to understand the relationship between the domestic regulatory environment and the development of the European regulatory framework for artificial intelligence.

The European Union has adopted Regulation (EU) 2024/1689 – Artificial Intelligence Act (EU AI Act), which establishes harmonized rules for artificial intelligence and introduces a risk-based regulatory approach. For organizations in Bosnia and Herzegovina, however, it is important to distinguish between the relevance of the EU AI Act and its direct applicability. The fact that Bosnia and Herzegovina is engaged in the European integration process does not mean that EU regulations automatically and directly apply as domestic legislation. Nevertheless, the development of the European regulatory framework represents an important reference point for organizations in Bosnia and Herzegovina, particularly those operating in the European market, providing services to EU-based organizations, or relying on European business partners.

The EU AI Act may also be directly relevant to certain organizations from Bosnia and Herzegovina depending on their role, the markets in which they operate, and how their AI systems are developed, placed on the market, deployed, or used. Organizations should therefore not assume that the EU AI Act is automatically irrelevant simply because they are established outside the European Union. Instead, they should assess their role in the AI value chain, the markets they serve, the location of AI system deployment and use, and whether the outputs of their AI systems are used within the European Union.

From a management perspective, one of the most important features of the EU AI Act is its risk-based approach. Not all AI systems present the same level of risk, and the level of governance and control should correspond to the potential impact of the system. For certain high-risk AI systems, requirements include, among other areas, risk management, data governance, technical documentation, record-keeping, transparency, human oversight, accuracy, robustness, and cybersecurity.

This risk-based principle is also highly relevant beyond the direct scope of the EU AI Act. Organizations should assess AI use cases according to their potential impact rather than treating all AI systems as equally significant. An AI system used to generate internal marketing content does not necessarily require the same level of governance as a system that supports credit decisions, employee selection, customer profiling, fraud detection, or other decisions that may significantly affect individuals.

Personal data protection represents another critical regulatory consideration for organizations using AI in Bosnia and Herzegovina. Of particular importance is the Law on the Protection of Personal Data of Bosnia and Herzegovina, which is largely aligned with the principles and requirements of the EU General Data Protection Regulation (GDPR). (Parliamentary Assembly of Bosnia and Herzegovina, 2025; European Parliament & Council of the European Union, 2016)

Its relevance to AI arises from the fact that many AI applications involve the processing of personal data, profiling, automated decision-making, or the

analysis of individual behavior. Consequently, AI risk assessment and privacy risk assessment should not be treated as completely separate processes. Where an AI system processes personal data, privacy and data protection considerations should be incorporated into the AI risk assessment from the earliest stages of the system lifecycle.

Depending on the specific use case, the organization should consider, among other factors:

- the legal basis for processing;
- the purpose and proportionality of the processing;
- the quality and source of the data;
- transparency;
- the rights of data subjects;
- the possibility of human intervention;
- the responsibilities of controllers and processors;
- the risks associated with automated decision-making;
- the need for a Data Protection Impact Assessment (DPIA), where applicable.

Particular attention should be given to AI systems that support decisions relating to credit assessment, employment, customers, profiling, fraud detection, or other processes that may significantly affect individuals. In such cases, demonstrating the technical accuracy of the model is not sufficient. The organization must also consider the legal, ethical, privacy, and organizational context in which the AI system is used.

This also illustrates an important principle of AI risk management: the risk associated with an AI system cannot be determined solely by examining the technology itself. The same model may present very different levels of risk depending on the data it processes, the purpose for which it is used, the people affected by its outputs, and the consequences of an incorrect result.

For management, regulatory and privacy considerations should therefore be addressed before the implementation of an AI system rather than treated as a subsequent compliance review. Changes to the model, data sources, processing purposes, or business processes may also change the regulatory and privacy risk profile.

Regulatory compliance should consequently be integrated into AI governance, change management, and continuous risk monitoring throughout the AI system lifecycle.

The key management question is therefore not simply: “Is this AI technology compliant?” but rather: “Is the way our organization develops, procures, deploys, and uses this AI system compliant with the applicable legal, regulatory, privacy, and organizational requirements.

8 MANAGEMENT-ORIENTED FRAMEWORK FOR AI RISK MANAGEMENT

Based on the synthesis of the AI risks, governance principles, organizational responsibilities, and regulatory considerations discussed in the preceding sections, this paper proposes a management-oriented AI risk management framework consisting of seven interconnected dimensions: business value, risk, data, accountability, vendors, compliance, and continuous monitoring.

The framework is designed as an extension of existing organizational governance and enterprise risk management processes rather than as a separate AI-specific management system. Its purpose is to provide management with a

structured approach for evaluating AI use cases before implementation, defining responsibilities and controls during deployment, and continuously monitoring risks throughout the AI system lifecycle.

The seven dimensions should not be understood as independent control areas. They are interconnected and should be considered together when an organization evaluates an AI use case. For example, the expected business value influences the level of risk that may be acceptable, while the risk profile determines the required level of governance, human oversight, security, privacy protection, and monitoring. Similarly, the use of external AI providers introduces additional vendors, data, accountability, and compliance considerations.

The proposed framework can therefore be applied as a management decision-making structure consisting of three basic questions: whether the AI use case provides sufficient business value, whether the associated risks are understood and acceptable, and whether the organization has established sufficient ownership, controls, and monitoring mechanisms to manage those risks.

Table 1 presents the seven dimensions of the proposed framework and summarizes their primary management focus.

Table 1. *Seven Dimensions of the Proposed Management-Oriented AI Risk Management Framework*

Dimension	Management focus	Source
Business Value	Why is AI needed and what value is expected?	(U.S. Department of Commerce, 2023); (International Organization for Standardization, 2023)
Risk	What can go wrong and is the risk acceptable?	(Palo Alto Networks, 2026); (Villanueva, 2025)
Data	Is the data appropriate, reliable, lawful and sufficiently governed?	(Peters, 2025a); (Future of Life Institute, 2026); (U.S. Department of Commerce, 2023)
Accountability	Who owns the use case, risk and final decision?	(OECD, 2024); (Stanley & Sheh, 2026); (U.S. Department of Commerce, 2023)
Vendors	How are third-party AI providers and model changes controlled?	(Peters, 2025b); (Hemmersbach, 2026); (ISACA, 2025b)
Compliance	Which legal, regulatory, privacy and organizational requirements apply?	(Bidault, 2026); (Osano, 2025); (European Parliament and Council of the European Union, 2016)
Continuous Monitoring	How are performance, drift, bias, incidents and significant changes detected?	(Data Privacy Office, 2025); (NIST, 2026)

Nevertheless, it should be recognized that the sources presented in Table 1 represent a supporting rather than exhaustive basis for the proposed dimensions. Given the rapid development of AI technologies, standards, regulatory requirements, and emerging risk practices, relevant sources should be continuously monitored and reviewed, with the framework updated as necessary.

The seven dimensions provide the management-level structure of the framework, while individual AI risk scenarios provide the operational level at which the framework can be implemented. Accordingly, each identified AI risk should be linked to an appropriate management control, a clearly defined responsible function, and the relevant regulatory or governance requirement.

This mapping is important because identifying an AI risk alone does not provide sufficient assurance that the risk is being managed. Management

should be able to determine what control addresses the risk, who is responsible for implementing and monitoring that control, and which regulatory, privacy, security, or governance requirements are relevant.

The proposed mapping therefore creates a direct connection between risk identification and management action. It also supports the integration of AI risk management into existing organizational processes such as enterprise risk management, information security management, privacy management, vendor management, compliance, internal audit, and business continuity.

Table 2 illustrates this approach by mapping selected AI risk scenarios to corresponding management controls, responsible organizational functions, and regulatory or governance relevance.

Table 2. Mapping of AI Risk Scenarios to Management Controls, Responsibilities, and Regulatory Requirements

AI risk	Management control	Responsible function	Regulatory / governance relevance
Data quality	Data governance, validation and testing	Data Owner / Business Owner	GDPR / AI Act
Bias and unfair outcomes	Fairness assessment and monitoring	Business/Risk Owner	AI Act/responsible AI
Model drift	Continuous performance monitoring and revalidation	Model/AI Owner	AI governance
Lack of explainability	Documentation, explainability assessment and human oversight	AI Owner / Business Owner	AI Act / accountability
Vendor risk	Due diligence, contractual controls, change notification and exit planning	Procurement / Risk / AI Owner	Third-party risk
Privacy risk	Privacy assessment, data minimization and DPIA where applicable	DPO / Privacy Function	BiH data protection law / GDPR
Security risk	Security testing, access control, monitoring and incident response	CISO / Security Function	AI Act / information security
Accountability risk	Defined ownership, approval authority and escalation	Management / Business Owner	AI governance
Intellectual property risk	Data and output rights assessment, contractual controls	Legal / Procurement / Business Owner	IP / contractual compliance
Human overreliance	Human oversight, training and intervention mechanisms	Business Owner / HR / Management	AI governance

The proposed framework provides a structured management layer through which AI-related risks can be incorporated into existing organizational processes. Its primary purpose is not to prescribe a fixed set of technical controls, but to establish a clear relationship between business objectives, AI risks, organizational responsibilities, and applicable governance requirements.

The framework also supports a risk-proportionate approach. AI use cases with limited potential impact may require relatively simple governance and monitoring mechanisms, while use cases involving personal data, significant decisions, financial consequences, or autonomous actions may require enhanced assessment, human oversight, security controls, documentation, and continuous monitoring.

From a management perspective, the framework can therefore be used throughout the AI lifecycle: during the initial evaluation of a proposed use case, during procurement or development, before deployment, during operational use, and when significant changes occur. This lifecycle perspective is important because the risk profile of an AI system may change as its data, model, users, purpose, vendor, or operating environment changes.

The framework should consequently be viewed as a continuous management process rather than a one-time assessment. Its effectiveness depends on clearly assigned accountability, integration with existing governance processes, and the organization's ability to detect and respond to changes in AI-related risk.

9 CONCLUSIONS

Artificial intelligence offers significant potential for automation, data analysis, prediction, content generation, and business decision-making support. However, its value does not arise solely from its technological capabilities, but from an organization's ability to use these capabilities within clearly defined objectives, responsibilities, and controls.

AI risk management should therefore not be viewed as an entirely new and separate system. A more effective approach is to treat it as an extension of the existing enterprise risk management and governance framework,

complemented by additional controls that address the specific characteristics and risks of AI technologies.

This approach is also consistent with the fundamental philosophy of the ISACA Advanced in AI Risk (AAIR) program, which views AI risk through the lens of established risk management principles and their extension to the AI environment. Of particular importance are the integration of AI risks into enterprise risk management, lifecycle management, ownership and accountability, regulatory compliance, supply chain risk management, monitoring, and continuous risk management.

For organizations in Bosnia and Herzegovina, this approach is particularly relevant given the development of the domestic regulatory framework and the country's ongoing alignment with the European legal acquis. The new Law on the Protection of Personal Data of Bosnia and Herzegovina provides the applicable domestic framework for the processing of personal data, while the EU AI Act represents an important regulatory reference and, in certain circumstances, may also have direct practical implications for organizations in Bosnia and Herzegovina that operate in or provide services to the European Union market.

Ultimately, effective AI risk management is not about preventing organizations from using AI. It is about enabling them to realize its benefits while ensuring that AI systems are used in a manner that is controlled, accountable, secure, compliant, and consistent with the organization's risk appetite.

The main contribution of this paper is the proposed management-oriented framework that translates these principles into seven interconnected dimensions: business value, risk, data, accountability, vendors, compliance, and continuous monitoring. The framework is intended as an integrative management layer that can be embedded into existing enterprise risk management and governance processes rather than established as a separate AI management system.

Its particular value for organizations in Bosnia and Herzegovina lies in connecting AI risk management with the domestic personal data protection framework and the evolving European AI regulatory environment.

REFERENCES

- Bidault, P.-E. (2026). *AI Act timeline: Key compliance deadlines under the European artificial intelligence regulation*. Dastra. <https://www.dastra.eu/en/blog/ai-act-timeline-key-compliance-deadlines-under-the-european-artifici/60278>
- Data Privacy Office. (2025). *Navigating the AI landscape: Understanding AI risk management frameworks*. Data Privacy Office. <https://data-privacy-office.eu/navigating-the-ai-landscape-understanding-ai-risk-management-frameworks/>
- European Parliament and Council of the European Union. (2016). *Regulation (EU) 2016/679 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data (General Data Protection Regulation)*. European Union. <https://eur-lex.europa.eu/eli/reg/2016/679/oj>
- European Parliament & Council of the European Union. (2024). *Regulation (EU) 2024/1689 – Artificial Intelligence Act*. *Official Journal of the European Union*.
- FAIR Institute. (2024). *A FAIR Artificial Intelligence (AI) Cyber Risk Playbook*. FAIR Institute.
- Future of Life Institute. (2026). *Article 10: Data and data governance*. EU Artificial Intelligence Act. <https://artificialintelligenceact.eu/article/10/>
- Giudici, P., Centurelli, M., & Turchetta, S. (2024). Artificial intelligence risk measurement. *Expert Systems with Applications*, 235, 121220. <https://doi.org/10.1016/j.eswa.2023.121220>
- Hemmersbach, S. (2026). *The NIST AI RMF and third-party risk: An implementation guide for TPRM programs*. Mitrtech. <https://mitrtech.com/resource-hub/blog/nist-ai-risk-management-framework-rmf/>
- International Organization for Standardization. (2023). *ISO 42001 explained*. <https://www.iso.org/home/insights-news/resources/iso-42001-explained-what-it-is.html>
- ISACA. (2025a). *Advanced in AI Risk (AAIR) Official Review Manual*.
- ISACA. (2025b). *Advanced in AI Security Management (AAISM) Official Review Manual*.
- International Organization for Standardization (ISO) & International Electrotechnical Commission (IEC). (2023a). *ISO/IEC 42001:2023 – Information technology — Artificial intelligence — Management system*. ISO.
- International Organization for Standardization (ISO) & International Electrotechnical Commission (IEC). (2023b). *ISO/IEC 23894:2023 – Information technology — Artificial intelligence — Guidance on risk management*. ISO.
- MIT AI Risk Initiative. (2024). *The MIT AI Risk Repository*. MIT FutureTech.
- MITRE. (2023). *MITRE ATLAS: Adversarial Threat Landscape for Artificial-Intelligence Systems*. MITRE Corporation.
- NIST. (2026). *AI risk management framework*. <https://www.nist.gov/itl/ai-risk-management-framework>
- OECD. (2024). *AI principles*. <https://www.oecd.org/en/topics/ai-principles.html>
- Osano. (2025). *AI risk management frameworks to manage risks in artificial intelligence systems*. JDSupra. <https://www.jdsupra.com/legalnews/ai-risk-management-frameworks-to-manage-5746651/>
- Oxford Internet Institute. (2017). *Oxford Internet Institute launches Digital Ethics Lab: Tackles ethical challenges posed by digital innovation*. University of Oxford.

- OWASP. (2025). *OWASP Top 10 for LLM Applications 2025*. Open Worldwide Application Security Project.
- Palo Alto Networks. (2026, July 21). *AI risk management frameworks: Everything you need to know*. Palo Alto Networks. <https://www.paloaltonetworks.com/cyberpedia/ai-risk-management-framework>
- Parliamentary Assembly of Bosnia and Herzegovina. (2025). *Personal Data Protection Law*. *Official Gazette of Bosnia and Herzegovina*, No. 12/25.
- Peters, S. (2025a). *ISO 42001 Annex A Control A.8 explained*. ISMS Online. <https://www.isms.online/iso-42001/annex-a-controls/a-8-information-for-interested-parties-of-ai-systems/>
- Peters, S. (2025b). *ISO 42001 Annex A Control A.10 explained*. ISMS Online. <https://www.isms.online/iso-42001/annex-a-controls/a-10-third-party-and-customer-relationships/>
- Scarpino, J. (2025). Ethics, accountability, and the pursuit of responsible AI. *ISACA Journal*, 3. <https://www.isaca.org/resources/isaca-journal/issues/2025/volume-3/ethics-accountability-and-the-pursuit-of-responsible-ai>
- Stanley, M., & Sheh, R. (2026). *Concept note: AI RMF profile on trustworthy AI in critical infrastructure*. NIST. <https://www.nist.gov/programs-projects/concept-note-ai-rmf-profile-trustworthy-ai-critical-infrastructure>
- U.S. Department of Commerce. (2023). *NIST AI 100-1: Artificial Intelligence risk management framework (AI RMF 1.0)*. <https://doi.org/10.6028/NIST.AI.100-1>
- UNESCO. (2021). *Recommendation on the ethics of artificial intelligence*. UNESCO.
- Villanueva, E. (2025, March 8). *Safeguard the future of AI: The core functions of the NIST AI RMF*. Optro. <https://optro.ai/blog/nist-ai-rmf>