

Documentary Heritage in the era of Artificial Intelligence: A Systematic Analysis of Risks and Challenges

Azam Najafgholinejad¹ 

1. PhD, knowledge and information science, Assistant Professor, Data Science, Information, and Artificial Intelligence Department, National library and archives of I.R. Iran, Tehran, Iran; a-najafgholinejad@nlai.ir

Abstract

Purpose: Documentary heritage constitutes a significant component of the historical memory, cultural capital, and informational resources of any nation. It reflects the collective memory of society and plays a pivotal role in linking the past to the present while shaping future perspectives. Artificial intelligence, with its capabilities in data processing and pattern recognition, offers novel tools for the preservation and restoration of written heritage. Applications such as text recognition through deep learning, digitization and reconstruction of damaged materials, large-scale analysis of historical data, and intelligent cataloging enable more accurate extraction and organization of information. Despite the broad potential of AI, the technology possesses a dual nature, and its risks cannot be overlooked. Consequently, engaging with AI requires a dialectical approach that simultaneously examines its opportunities and threats. The integration of AI into libraries and archives introduces new possibilities while also presenting fresh challenges and risks. The present study aims to identify and analyze the challenges and risks associated with the use of artificial intelligence in preserving documentary heritage within libraries and archives.

Method: This research is applied in terms of purpose, qualitative in terms of approach, and exploratory interview-based in terms of data collection. To identify the risks and challenges associated with the implementation of artificial intelligence (AI) in libraries and archives, 20 experts in Knowledge and Information Science, archival science, information technology, and artificial intelligence were selected through purposive and snowball sampling. Subsequently, semi-structured exploratory interviews were conducted with these experts. Expert selection was based on three key criteria: professional experience and active engagement in reputable national institutions (such as the National Library and Archives of Iran, the Iranian Research Institute for Information Science and Technology, and the Computer Research Center of Islamic Sciences) and leading AI companies, with emphasis on backgrounds in information science, librarianship, archival studies, IT, and AI applications; possession of at least a master's degree in relevant fields and specialized knowledge in AI as well as the preservation of written and documentary heritage; and research and practical experience in national projects or similar initiatives demonstrating applied expertise. Data were coded and categorized using inductive qualitative content analysis, resulting in 167 final codes (4 main categories and 15 subcategories). Data analysis was performed using MAXQDA software, and Guba and Lincoln's criteria were applied to ensure validity and reliability.

Findings: Data revealed that the implementation of AI technologies in libraries and archives—particularly in the preservation of documentary heritage—faces a range of multidimensional challenges and risks. These risks span technological, legal and security, managerial, and socio-cultural dimensions. Managerial, infrastructural, and cultural risks exhibited the highest frequency.

Conclusion: The successful implementation of artificial intelligence in organizations such as the National Library and Archives of Iran, as the country's central library, requires a comprehensive and systematic approach. Beyond purely technological considerations, such an approach must simultaneously encompass governance, strategy, data and infrastructure management, security, ethics, as well as human-capital development and organizational culture. Neglecting these interconnected components significantly heightens the risk of project failure, resource wastage, and erosion of public trust. Therefore, sustainable success in AI-driven initiatives depends on effective data and technology management, the establishment of clear legal and security policies, provision of necessary financial and technical infrastructures, enhancement of training and organizational culture, and the implementation of continuous monitoring and evaluation mechanisms. This study represents one of the first attempts to provide a comprehensive picture of the risks associated with artificial intelligence in organizations responsible for preserving the nation's documentary heritage. It also offers a foundation for developing risk mitigation strategies and promoting the responsible adoption of artificial intelligence in this field.

Keywords: Artificial Intelligence, Risk Management, Libraries and Archives, Documentary and Written Heritage, Data Governance, Algorithm



Publisher: *National Library and Archives of I.R. of Iran*

© The Author(s).

Doi: 10.30484/NASTINFO.2021.2971.2077

Extended Abstract

Introduction

Documentary heritage constitutes a significant component of the historical memory, cultural capital, and informational resources of any nation. It reflects the collective memory of society and plays a pivotal role in linking the past to the present while shaping future perspectives. Artificial intelligence, with its capabilities in data processing and pattern recognition, offers novel tools for the preservation and restoration of written heritage. Applications such as text recognition through deep learning, digitization and reconstruction of damaged materials, large-scale analysis of historical data, and intelligent cataloging enable more accurate extraction and organization of information. Despite the broad potential of AI, the technology possesses a dual nature, and its risks cannot be overlooked. Consequently, engaging with AI requires a dialectical approach that simultaneously examines its opportunities and threats. The integration of AI into libraries and archives introduces new possibilities while also presenting fresh challenges and risks.

Purpose

The present study aims to identify and analyze the challenges and risks associated with the use of artificial intelligence in preserving documentary heritage within libraries and archives.

Method

This study employed a qualitative research approach. To identify the risks and challenges of implementing AI in libraries and archives, targeted exploratory interviews were conducted with 20 experts in information science, archival studies, information technology, and artificial intelligence. Expert selection was based on three key criteria: professional experience and active engagement in reputable national institutions (such as the National Library and Archives of Iran, the Iranian Research Institute for Information Science and Technology, and the Computer Research Center of Islamic Sciences) and leading AI companies, with emphasis on backgrounds in information science, librarianship, archival studies, IT, and AI applications; possession of at least a master's degree in relevant fields and specialized knowledge in AI as well as the preservation of written and documentary heritage; and research and practical experience in national projects or similar initiatives demonstrating applied expertise. Data were coded and categorized using inductive qualitative content analysis, resulting in 167 final codes. Data analysis was performed using MAXQDA software, and Guba and Lincoln's criteria were applied to ensure validity and reliability.

Findings

Insights derived from the expert interviews revealed that the implementation of AI technologies in libraries and archives—particularly in the preservation of documentary heritage—faces a range of multidimensional challenges and risks. These risks span technological, legal and security, managerial, and socio-cultural dimensions. Managerial, infrastructural, and cultural risks exhibited the highest frequency. Findings derived from the interviews, which reflect a wide spectrum of professional concerns, are comprehensively summarized in Table 1.

Table 1. Challenges and Risks of Artificial Intelligence in the Preservation of Documentary Heritage

Main Category	Subcategory	Definition / Brief Explanation	Frequency
Technological Risks	Data Limitations, Quality, and Compatibility	Errors in learning and retrieval caused by noisy or incomplete datasets; need for accuracy and task-relevant data; requirement for preprocessing, cleaning, and precise annotation; refining, aligning, and integrating data to ensure consistency; adequacy and comprehensiveness of training data.	15
	Linguistic Limitations	Limited compatibility of large language models and text-processing tools with Persian and other non-English languages.	4
	Data Bias and Inaccurate Model Evaluation	Restricted learning of models from only certain aspects of a phenomenon; biased or incomplete outputs due to skewed data; imbalanced training data affecting model accuracy; inability to detect rare cases; evaluating models using inadequate metrics; producing misleading assessments of model performance.	8
	Hallucination	AI hallucination and incorrect generalization due to uneven or insufficient data.	5

	System Reliability and Dependence	Reduced performance accuracy and algorithmic errors; increased organizational dependence on automated technologies; interruptions or breakdowns in intelligent system operations.	14
	Inefficiency in Transferring Domain Knowledge	Failure to accurately transfer expert knowledge into AI systems.	1
Legal and Security Risks	Legal, Confidentiality, and Access Issues	Risks of copyright infringement and disclosure of restricted or confidential data; privacy risks.	25
	Security Risks and Access Restrictions	Cyberattacks on servers; unauthorized access; network outages or disruptions; execution of malicious commands.	12
Managerial and Governance Risks	Financial and Infrastructural Challenges	Need for powerful computing hardware (advanced GPUs, stable electricity, cooling systems); software infrastructure and data governance (local model installation and training, data-sharing restrictions due to security); dependence on foreign technologies (operational and sanctions-related risks, difficulty procuring up-to-date hardware/models); operational costs and risks of AI implementation; maintenance, updating, and continuous monitoring of AI models.	28
	Access to AI Tools	Availability and usability of AI tools.	2
	Human Resource Skills and Expertise Risks	Shortage of technical and specialized skills; lack of skilled personnel; difficulty in recruiting and retaining experts.	7
	Managerial, Strategic, and IT Governance Risks	Lack of clear prioritization for AI deployment; absence of systematic needs assessment; weak sustainability and project management; insufficient organizational culture; undervaluation of emerging technologies.	19
	Unrealistic Expectations	Exaggerated or unrealistic expectations of AI capabilities.	4
Socio-Cultural Risks	Cultural and Organizational Challenges	Organizational resistance and slow adoption of new technologies; skill gaps and fear of learning complex technologies; weak interdepartmental collaboration.	20
	Authenticity and Credibility of AI-Generated Texts	Concerns regarding accuracy and reliability of AI-generated information; ambiguity about the contribution of AI versus humans in content creation.	3
Total	—	—	167

Conclusion

The successful implementation of artificial intelligence in organizations such as the National Library and Archives of Iran, as the country's central library, requires a comprehensive and systematic approach. Beyond purely technological considerations, such an approach must simultaneously encompass governance, strategy, data and infrastructure management, security, ethics, as well as human-capital development and organizational culture. Neglecting these interconnected components significantly heightens the risk of project failure, resource wastage, and erosion of public trust. Therefore, sustainable success in AI-driven initiatives depends on effective data and technology management, the establishment of clear legal and security policies, provision of necessary financial and technical infrastructures, enhancement of training and organizational culture, and the implementation of continuous monitoring and evaluation mechanisms.

Acknowledgments

The author gratefully acknowledges all experts who participated in the interviews and contributed their professional knowledge and experience to this study.

Ethical Considerations

Participation in the study was voluntary, and the participants were informed about the purpose and procedures of the research. The confidentiality and anonymity of participants and their responses were maintained throughout the research and reporting processes.

Use of Artificial Intelligence

Artificial intelligence tools were not used to conduct the interviews, analyze the research data, generate the findings, or draw the conclusions of this study. Any use of AI-assisted tools, if applicable, was limited to language editing and improving the clarity of the manuscript.

Conflict of Interest

The authors declare that they have no conflict of interest regarding the publication of this article.

میراث مستند در عصر هوش مصنوعی: تحلیل نظام‌مند ریسک‌ها و چالش‌ها^۱

اعظم نجفقلی‌نژاد^۱ 

۱. دکترای علم اطلاعات و دانش‌شناسی، استادیار گروه پژوهشی علوم داده، اطلاعات و هوش مصنوعی سازمان اسناد و کتابخانه ملی

ایران، تهران، ایران؛ a-najafgholinejad@nlai.ir

چکیده

هدف: میراث مستند بخش مهمی از حافظه تاریخی، سرمایه فرهنگی و منابع اطلاعاتی هر کشور را تشکیل می‌دهد. این میراث بازتاب‌دهنده حافظه جمعی جامعه است و در ایجاد پیوند میان گذشته و حال و همچنین در شکل‌دهی چشم‌انداز آینده نقشی اساسی ایفا می‌کند. هوش مصنوعی با توانمندی‌های خود در پردازش داده‌ها و تشخیص الگوها، ابزارهایی نوین برای حفاظت و بازسازی میراث مکتوب فراهم می‌کند. کاربردهایی همچون شناسایی متون با روش‌های یادگیری عمیق، دیجیتال‌سازی و بازسازی بخش‌های آسیب‌دیده، تحلیل کلان داده‌های تاریخی و فهرست‌نویسی هوشمند منابع، امکان استخراج و سازمان‌دهی دقیق‌تر اطلاعات را فراهم می‌سازد. باوجود ظرفیت‌های گسترده هوش مصنوعی، این فناوری ماهیتی دوگانه دارد و تهدیدهای آن قابل چشم‌پوشی نیست. بنابراین، مواجهه با آن مستلزم رویکردی متوازن و مبتنی بر سنجش هم‌زمان فرصت‌ها و مخاطرات است. ورود هوش مصنوعی به کتابخانه‌ها و آرشیوها، در کنار ایجاد امکانات نوین، چالش‌ها و ریسک‌های تازه‌ای را نیز به همراه دارد. پژوهش حاضر با هدف شناسایی و تحلیل چالش‌ها و ریسک‌های هوش مصنوعی در حفاظت از میراث مستند در کتابخانه‌ها و آرشیوها انجام شده است.

روش: پژوهش حاضر از نظر هدف، کاربردی، از نظر رویکرد، کیفی و از نظر نحوه گردآوری داده‌ها، مصاحبه اکتشافی است. به‌منظور شناسایی ریسک‌ها و چالش‌های پیاده‌سازی هوش مصنوعی در کتابخانه‌ها و آرشیوها، ۲۰ نفر از خبرگان حوزه‌های علم اطلاعات و دانش‌شناسی، آرشیو، فناوری اطلاعات و هوش مصنوعی از طریق نمونه‌گیری هدفمند و گلوله‌برفی انتخاب و مصاحبه‌های اکتشافی نیمه‌ساختاریافته با آنان انجام گرفت. در فرآیند گزینش خبرگان، سه معیار کلیدی مدنظر قرار گرفت: تجربه حرفه‌ای و تمرکز بر متخصصان فعال در نهادهای معتبر ملی (مانند سازمان اسناد و کتابخانه ملی ایران، پژوهشگاه علوم و فناوری اطلاعات ایران، مرکز تحقیقات کامپیوتری علوم اسلامی) و شرکت‌های پیشرو در حوزه هوش مصنوعی، با تأکید بر سوابق در علم اطلاعات، کتابداری، آرشیو، فناوری اطلاعات و کاربرد هوش مصنوعی؛ دارا بودن حداقل مدرک کارشناسی ارشد در رشته‌های مرتبط و کسب دانش تخصصی در زمینه‌های هوش مصنوعی و حفاظت از میراث مکتوب و مستند؛ و سوابق پژوهشی و اجرایی در طرح‌های ملی و پروژه‌های مشابه که گواه توانمندی و تجربه عملی خبرگان در حوزه‌های مرتبط باشد. داده‌ها با روش تحلیل محتوای کیفی استقرایی کدگذاری و مقوله‌بندی شد و در نهایت ۱۶۷ کد نهایی (شامل ۴ مقوله اصلی و ۱۵ زیرمقوله) استخراج شد. تحلیل داده‌ها با استفاده از نرم‌افزار مکس کیودی‌ای انجام شد و برای تضمین اعتبار و پایایی از معیارهای گویا و لینکلن بهره گرفته شد.

یافته‌ها: یافته‌ها نشان داد پیاده‌سازی فناوری‌های هوش مصنوعی در کتابخانه‌ها و آرشیوها، به‌ویژه در حوزه حفاظت از میراث مستند، با مجموعه‌ای از چالش‌ها و ریسک‌های چندبعدی مواجه است. این ریسک‌ها شامل ابعاد فناوری، حقوقی و امنیتی، مدیریتی و فرهنگی - اجتماعی هستند. بیشترین فراوانی به ریسک‌های مدیریتی، زیرساختی و فرهنگی اختصاص دارد.

نتیجه‌گیری: پیاده‌سازی موفقیت‌آمیز هوش مصنوعی در سازمان‌هایی نظیر سازمان اسناد و کتابخانه ملی ایران، به‌عنوان کتابخانه مادر کشور، مستلزم اتخاذ رویکردی جامع و نظام‌مند است. این رویکرد باید فراتر از ملاحظات صرفاً فناورانه، ابعاد حکمرانی، راهبرد، مدیریت داده و زیرساخت، امنیت، اخلاق و همچنین توسعه سرمایه انسانی و فرهنگ سازمانی را به‌طور هم‌زمان در بر گیرد. عدم توجه به این مؤلفه‌های چندگانه، ریسک شکست پروژه‌ها، اتلاف منابع و تضعیف اعتماد عمومی را به‌طور قابل‌ملاحظه‌ای افزایش می‌دهد. بنابراین موفقیت در اجرای پایدار پروژه‌های هوش مصنوعی، در گرو مدیریت مؤثر داده و فناوری، تدوین سیاست‌های امنیتی و حقوقی شفاف، تأمین زیرساخت‌های مالی و فنی لازم، ارتقاء آموزش و فرهنگ سازمانی و استقرار سازوکارهای نظارت و پایش مستمر است. این پژوهش، یکی از نخستین تلاش‌ها برای ارائه تصویری جامع از مخاطرات این فناوری در سازمان‌های

^۱ مقاله پیش‌رو مستخرج از طرح پژوهشی موظف با عنوان «شناسایی و تحلیل طرح‌های بین‌المللی هوش مصنوعی در حفظ میراث مستند با تمرکز بر مدیریت ریسک AI در سازمان اسناد و کتابخانه ملی ایران» به کارفرمایی معاونت پژوهش و منابع دیجیتال سازمان اسناد و کتابخانه ملی ایران است.

حافظ میراث مستند کشور به شمار می‌آید و می‌تواند مبنایی برای طراحی راهبردهای کاهش ریسک و بهره‌برداری مسئولانه از هوش مصنوعی در این حوزه باشد.

کلیدواژه‌ها: هوش مصنوعی، مدیریت ریسک، کتابخانه‌ها و مراکز آرشیوی، میراث مستند و مکتوب، حکمرانی داده، الگوریتم

نوع مقاله: پژوهشی / مروری / نقد و بررسی

تاریخ دریافت: ۱۴۰۱/۰۱/۰۱؛ **دریافت آخرین اصلاحات:** روز/ماه/سال؛ **پذیرش:** روز/ماه/سال

استناد:

نجفقلی نژاد، اعظم (۱۴۰۵). میراث مستند در عصر هوش مصنوعی: تحلیل نظام‌مند ریسک‌ها و چالش‌ها. *مطالعات کتابداری و سازماندهی*

اطلاعات، (۰)، -، Doi:



© نویسندگان

ناشر: سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران

Doi: 10.30484/NASTINFO.2021.2971.2077

مقدمه

میراث مستند بخش وسیعی از حافظه، میراث و منابع اطلاعاتی هر کشور شمرده می‌شود. این میراث در حقیقت معرف حافظه جمعی مردم است و در برقراری ارتباط میان گذشته و حال، و ترسیم آینده از اهمیت حیاتی برخوردار است (کریمی، ۱۳۹۱). طبق تعریف یونسکو^۱ (۲۰۲۵) میراث مستند به منابع و اسناد منفرد یا مجموعه‌هایی از اسناد گفته می‌شود که دارای ارزش پایدار برای جامعه یا بشریت هستند و از دست رفتن آن‌ها موجب آسیب جبران‌ناپذیر به حافظه جمعی می‌شود. این اسناد می‌توانند در قالب‌های مختلفی از جمله مکتوب، چاپی، دیداری، شنیداری یا دیجیتال باشند و به صورت مجموعه‌هایی منسجم براساس موضوع، منشأ یا زمینه تاریخی گردآوری شده باشند.

تکامل مداوم دیجیتالی شدن و تحول دیجیتال، به‌ویژه از زمان بروز انقلاب صنعتی چهارم، منجر به تولید حجم وسیعی از داده‌ها، اطلاعات و میراث مستند شده است. روش‌های سنتی حفاظت و مدیریت منابع و اسناد دیگر پاسخگوی نیازهای روزافزون نیستند و خطر از دست رفتن یا آسیب دیدن منابع و اسناد ارزشمند وجود دارد. هوش مصنوعی پتانسیل زیادی برای حل این مشکل دارد. ابزارهای هوش مصنوعی به‌عنوان راه‌حلی برای تحلیل حجم عظیم داده‌ها ظهور کرده و پتانسیل تأثیرگذاری بر هر جنبه‌ای از زندگی را دارند (Badawy, 2023).

هوش مصنوعی با قابلیت‌های خود در پردازش داده‌ها و شناسایی الگوها، می‌تواند به پژوهشگران و متخصصان کمک کند تا روش‌های نوینی برای حفظ و نگهداری میراث مکتوب و مستند توسعه دهند. تشخیص و شناسایی متون با استفاده از الگوریتم‌های یادگیری عمیق^۲ و پردازش زبان طبیعی، دیجیتال‌سازی و تولید نسخه‌های دیجیتال با کیفیت بالا با استفاده از فنون پردازش تصویر و بازسازی بخش‌های آسیب‌دیده یا گم‌شده با استفاده از الگوریتم‌های یادگیری عمیق، تحلیل کلان داده‌های مرتبط با میراث مکتوب شامل شناسایی الگوهای تاریخی، بررسی روندهای تاریخی و اجتماعی و استخراج اطلاعات جدید از متون قدیمی و فهرست‌نویسی هوشمند و خودکار منابع به‌منظور تسهیل دسترسی به آن‌ها با استفاده از الگوریتم‌های یادگیری ماشین از آن جمله است (Yu et al., 2019).

باوجود تمام پتانسیل‌های مذکور، فناوری شمشیری دو لبه است و ریسک‌های آن نمی‌توانند نادیده گرفته شوند. ورود هوش مصنوعی به حوزه کتابخانه‌ها و آرشیوها، در کنار فرصت‌های فراوان، همراه با ریسک‌ها و چالش‌های جدیدی است (Wang & Liu, 2023). ریسک در حوزه هوش مصنوعی به معنای احتمال وقوع رویداد نامطلوب یا ناخواسته‌ای است که می‌تواند به داده‌ها، سامانه‌ها، فرایندها، یا شهرت سازمان آسیب رسانده و یا به اهداف راهبردی آن خدشه وارد کند. مناطقی همچون سوگیری الگوریتمی، نگرانی‌های حریم خصوصی (Okoroma, 2024)، وابستگی به زیرساخت‌های پرهزینه، فقدان شفافیت تصمیم‌گیری (ISO 31000, 2018)، تهدیدهای امنیت سایبری، کاهش اعتماد کاربران (NIST, 2023) و نیز مقاومت‌های فرهنگی و سازمانی، تنها بخشی از ریسک‌هایی هستند که می‌توانند پیاده‌سازی هوش مصنوعی را با ناکارآمدی یا حتی شکست مواجه سازند. ازاین‌رو، ضرورت دارد هرگونه حرکت به‌سمت بهره‌گیری از هوش مصنوعی در حفاظت از میراث مستند، با مدیریت ریسک نظام‌مند همراه باشد.

^۱UNESCO^۲deep learning^۳National Institute of Standards and Technology

استانداردهای متعددی در خصوص مدیریت ریسک در سازمان‌ها تدوین و منتشر شده است که از جمله آن‌ها می‌توان به استاندارد ایزو/آی.ای.سی. ۲۳۸۹۴^۱ (۲۰۲۳) و ایزو ۳۱۰۰۰ (۲۰۱۸) اشاره کرد. درحالی‌که این استانداردها و روش‌ها به سازمان‌ها کمک می‌کنند تا خطرات سامانه‌های نرم‌افزاری یا اطلاعاتی سنتی را کاهش دهند، اما خطرات ناشی از سامانه‌های هوش مصنوعی منحصر به فرد هستند. به عنوان مثال، سامانه‌های هوش مصنوعی ممکن است براساس داده‌هایی آموزش ببینند که ممکن است با گذشت زمان تغییر کنند و این قضیه می‌تواند بر عملکرد و قابلیت اعتماد سامانه تأثیرگذار باشد (Ames & Lewis, 2020; NIST, 2023; Yazdi et al., 2024). همچنین سامانه‌های هوش مصنوعی ذاتاً اجتماعی-فنی هستند، به این معنی که تحت تأثیر دینامیک‌های اجتماعی و رفتار انسانی قرار دارند (Ames & Lewis, 2020; NIST, 2023). با این حال، با کنترل‌های مناسب، می‌توان بسیاری از پیامدهای نامطلوب ناشی از هوش مصنوعی، از جمله سوگیری در پردازش اطلاعات، نابرابری در دسترسی کاربران به خدمات و خطاهای تصمیم‌گیری مبتنی بر داده را کاهش داد و مدیریت کرد (NIST, 2023).

طبق استانداردها و چهارچوب‌های موجود، مدیریت ریسک باید در طول چرخه عمر سامانه هوش مصنوعی مشتمل بر مراحل قبل از پیاده‌سازی، مرحله طراحی و توسعه و مرحله پیاده‌سازی و اجرا انجام شود تا اطمینان حاصل شود که مدیریت ریسک به صورت مداوم و به موقع صورت می‌گیرد (NIST, 2023; ISO 31000, 2018). شناسایی و تحلیل ریسک‌ها و چالش‌ها در هر مرحله می‌تواند به شناخت وضعیت کنونی این فناوری‌ها کمک کند. علی‌رغم گسترش روزافزون پژوهش‌ها در حوزه کاربردهای هوش مصنوعی در کتابخانه‌ها و آرشیوهای جهان، تاکنون پژوهشی که به صورت نظام‌مند و با اتکا به نظر خبرگان ایرانی، به شناسایی و اولویت‌بندی ریسک‌های این فناوری در بافت میراث مستند کشور پرداخته باشد، انجام نشده است. از منظر علمی، شناخت نظام‌مند ریسک‌های هوش مصنوعی در بافت خاص میراث مستند، می‌تواند به توسعه ادبیات مدیریت ریسک فناوری در حوزه علم اطلاعات و دانش‌شناسی کمک کرده و الگویی برای پژوهش‌های آتی در سایر نهادهای فرهنگی فراهم آورد. در سطح عملیاتی، سازمان اسناد و کتابخانه ملی ایران به عنوان اصلی‌ترین نهاد متولی حفاظت از میراث مستند کشور، در آستانه تصمیم‌گیری برای سرمایه‌گذاری در پروژه‌های هوش مصنوعی قرار دارد. نبود تصویر شفافی از ریسک‌های پیش‌رو، می‌تواند منجر به اتلاف منابع، شکست پروژه‌ها و حتی آسیب به اعتماد عمومی شود. بنابراین، شناسایی دقیق ریسک‌ها از دیدگاه خبرگان، نیاز فوری و راهبردی برای این سازمان محسوب می‌شود.

سازمان اسناد و کتابخانه ملی ایران نه تنها وظیفه گردآوری، سازماندهی و حفاظت از میراث مکتوب و اسناد تاریخی کشور را بر عهده دارد، بلکه به دلیل نقش الگوساز خود برای دیگر کتابخانه‌ها و مراکز آرشیوی، هرگونه تجربه موفق یا ناموفق در پیاده‌سازی هوش مصنوعی در این سازمان، می‌تواند مسیر فناوری در کل شبکه نهادهای حافظه‌ای کشور را تحت تأثیر قرار دهد. با این اوصاف، پژوهش حاضر در پی پاسخ به این پرسش است که چالش‌ها و ریسک‌های مرتبط با هوش مصنوعی در حفاظت از میراث مستند در کتابخانه‌ها و آرشیوها از جمله سازمان اسناد و کتابخانه ملی ایران به عنوان کتابخانه مادر کدام‌اند؟

پیشینه پژوهش

به منظور آگاهی از وضعیت پژوهش‌های انجام‌شده در حوزه ریسک‌های هوش مصنوعی در حفاظت از میراث مستند، مرور نظام‌مندی بر منابع منتشرشده در دهه اخیر انجام شد. هدف اصلی این مرور، شناسایی خلأهای پژوهشی موجود، به‌ویژه در زمینه مطالعات بومی و سازمان‌محور و نیز دستیابی به الگوهای نظری و روش‌شناختی مورد استفاده در پژوهش‌های پیشین

^۱ISO/IEC 23894

است. ساختار این بخش به گونه‌ای طراحی شده است که ابتدا به پژوهش‌های داخلی، سپس به پژوهش‌های خارجی پرداخته و در پایان با جمع‌بندی و نتیجه‌گیری، شکاف دانش موجود و جایگاه پژوهش حاضر را مشخص می‌کند.

مقربى منظرى (۱۳۹۸) با روش پیمایشی-تحلیلی به ارزیابی ریسک‌های منابع انسانی کتابخانه‌های دیجیتالی دانشگاه‌های دولتی شهر تهران پرداخت. این ریسک‌ها در چهارگروه عملیاتی، سرمایه‌های انسانی، مهارت‌های فردی و ادراکی و مهارت‌های تخصصی شناسایی و براساس سه معیار احتمال وقوع، میزان تأثیر و عدم توانایی سازمان در واکنش به ریسک، میان اعضای رتبه‌بندی شد. نتایج حاصل نشان داد به‌ترتیب ریسک‌های مهارت‌های تخصصی، عملیاتی، سرمایه‌های انسانی و در نهایت ریسک‌های مهارت‌های فردی، اولویت‌های یک تا چهار را به خود اختصاص داده‌اند.

محمدصالحی و همکاران (۱۴۰۱) به ارائه الگوی پیشنهادی برای مدیریت ریسک تولید محتوا در ۲۸ کتابخانه دیجیتالی دولتی شهر تهران پرداختند. در این پژوهش نیز، سیاهه‌ای از ریسک‌های تولید محتوا براساس سه معیار شدت اثر، احتمال وقوع و توانایی سازمان در واکنش به ریسک رتبه‌بندی شدند. براساس رتبه‌بندی ریسک‌ها به‌ترتیب ریسک‌های حفاظت و نگهداری، نیروی انسانی، ریسک‌های فنی، امنیت اطلاعات و ریسک‌های زیرساختی، ارزیابی محتوای منابع و پدیدآورندگان، حقوق پدیدآورندگان، یکپارچه‌سازی و ریسک‌های محیطی رتبه ۱ تا ۹ را از نظر اهمیت ریسک‌ها کسب نمودند.

صمیعی و همکاران (۱۴۰۱) با روش پیمایشی تحلیلی با فنون آنتروپی شانون برای وزن‌دهی ریسک‌ها و تاپسیس فازی برای رتبه‌بندی ریسک‌ها، به شناسایی و رتبه‌بندی ریسک‌های زنجیره تأمین در کتابخانه‌های دیجیتالی دانشگاه‌های دولتی شهر تهران براساس استاندارد ایزو ۳۱۰۰۰ پرداختند. براساس رتبه‌بندی ریسک‌ها به‌ترتیب «ریسک‌های مربوط به تأمین‌کنندگان»، «ریسک‌های مشترک بین اجزای زنجیره تأمین»، «ریسک‌های مربوط به خدمات به کاربران»، «ریسک‌های مربوط به تولید و مدیریت منابع دیجیتالی» و «ریسک‌های مربوط به توزیع منابع دیجیتالی» رتبه ۱ تا ۵ را از نظر اهمیت ریسک‌ها کسب نموده‌اند. در خارج از کشور امز و لويس^۱ (۲۰۲۰) به بررسی چالش‌های ارائه‌شده توسط کلان داده میراث فرهنگی در زمینه توسعه خدمات دانش‌اندوزی دیجیتال کتابخانه ملی اسکاتلند پرداختند که هدف آن تشویق و فعال کردن استفاده از روش‌های محاسباتی در ارتباط با مجموعه بود. نتایج نشان داد علاوه بر فرصت‌هایی که در حال حاضر برای پشتیبانی از تجزیه و تحلیل در مقیاس بزرگ از مجموعه‌های میراث فرهنگی ایجاد شده است، چالش‌های جدیدی در مورد مدیریت داده‌ها، ذخیره‌سازی، حقوق، قالب‌ها، مهارت‌ها و دسترسی پیش آمده است.

کورکا^۲ و همکاران (۲۰۲۲) نیز موضوع حفاظت از میراث طبیعی و فرهنگی، موزه‌ها و آرشیوها در برابر ریسک‌ها را در پیوند با هوش مصنوعی، ارزیابی ریسک و مسئله ذی‌نفعان بررسی نمود. هدف پژوهش ایشان، ارائه نقش کمیته یونانی آرم آبی^۳ در یونان از یک سو و روش کاربردی جدیدی از سوی دیگر بود که شامل فنون ارزیابی ریسک و فناوری‌های برتر مانند هوش مصنوعی و اینترنت اشیا بود. این پژوهش، روش جدید جامع و قابل انطباق به نام INBO را معرفی نمود. INBO، نحوه واکنش و مدیریت پاسخ‌دهندگان اولیه و مدیران بناها، حتی بازدیدکنندگان را در موقعیت‌های خطرناک یا اضطراری بهبود می‌بخشد و پیش‌بینی صحیح را همراه با تصمیم‌های هوشمندانه ممکن می‌سازد که به جلوگیری از آسیب احتمالی به میراث فرهنگی و طبیعی بدون تلفات انسانی کمک می‌کند.

^۱Ames & Lewis

^۲Korka

^۳Blue Shield: Blue Shield به حفاظت از میراث فرهنگی و طبیعی در سطح جهانی پرداخته و بر اهمیت حفظ این منابع در شرایط بحرانی تأکید دارد.

وانگ و لیو^۱ (۲۰۲۳) پژوهشی در ارتباط با سامانه پیشگیری و کنترل ریسک هوش مصنوعی برای کتابخانه‌های هوشمند انجام دادند. آن‌ها دریافتند هوش مصنوعی در حال بازسازی مدل مدیریت و مدل خدمات کتابخانه‌های هوشمند است و کتابخانه‌ها باید در حین استقبال از تغییر مدل، به پیشگیری و کنترل ریسک توجه کنند.

باداوی^۲ (۲۰۲۳) تحلیل جامعی از چهارچوب مبتنی بر داده (DDF)^۳ برای ارزیابی ریسک دیجیتالی شدن و هوش مصنوعی ارائه داد. این چهارچوب ارزیابی را در مراحل مختلف چرخه حیات داده‌ها دسته‌بندی و از تحلیل پستل به عنوان مثال توصیفی برای درک ریسک‌های قانونی، اخلاقی و زمینه‌ای مرتبط با هر مرحله استفاده می‌کند. با اتخاذ این رویکرد جامع، پروژه بر تعامل پیچیده بین دیجیتالی شدن، هوش مصنوعی و عوامل محیطی مختلف پرتو می‌افکند و تصمیم‌گیری آگاهانه و اجرای مسئولانه برای رشد پایدار را تضمین می‌کند.

پانسونی^۴ و همکاران (۲۰۲۳) نیز با طراحی و ارزیابی چهارچوبی اخلاقی، به بررسی کاربردها، چالش‌ها و فرصت‌های اخلاقی مرتبط با استفاده از هوش مصنوعی در حوزه میراث فرهنگی پرداختند و اصولی برای تضمین پیاده‌سازی مسئولانه و قابل اعتماد این فناوری ارائه شده است.

اوکوروما^۵ (۲۰۲۴) نیز با استفاده از رویکردی کیفی و بررسی نظام‌مند متون به بررسی تأثیر دگرگون‌کننده هوش مصنوعی بر کتابخانه‌ها پرداخت و نقش هوش مصنوعی در تحلیل رفتار کاربر، پیش‌بینی روندها و بهبود مدیریت منابع و پیامدهای اخلاقی پیاده‌سازی هوش مصنوعی، به‌ویژه در مورد حریم خصوصی و رضایت در تحلیل یادگیری را برجسته نمود که در نهایت منجر به بهبود تجربه کاربری می‌شود.

علی^۶ و همکاران (۲۰۲۴) با رویکرد قوم‌نگاری، تحلیلی مختصر از نقاط قوت، ضعف، فرصت‌ها و تهدیدات (SWOT) مربوط به کاربرد هوش مصنوعی در کتابخانه‌های دانشگاهی پاکستان ارائه کردند. یافته‌ها نشان داد هوش مصنوعی به آرامی در حال معرفی شدن به کتابخانه‌های دانشگاهی پاکستان است و نیازمند سرمایه‌گذاری در زمینه تأمین مالی، زمان و نیروی انسانی است. تیل^۸ (۲۰۲۴) نیز با مرور متون به بررسی نقش تحول‌آفرین هوش مصنوعی در حفاظت دیجیتالی آرشیوهای تاریخی پرداخت و چالش‌های پیش‌روی کتابخانه‌ها، آرشیوها و موزه‌ها در نگهداری مجموعه‌های تاریخی را مورد بررسی قرار داد. یافته‌ها بیان می‌کند هوش مصنوعی می‌تواند به عنوان کاتالیزوری برای تغییر عمل کند و به‌طور بالقوه با افزایش کارایی و کاهش تلاش دستی، فرآیندهای نگهداری را دوباره تعریف کند. یافته‌ها همچنین نشان می‌دهد درحالی‌که هوش مصنوعی فرصت‌های قابل توجهی را ارائه می‌دهد، توجه دقیق و تحقیقات بیشتر برای هدایت خطرات و نگرانی‌های اخلاقی مرتبط با آن ضروری است.

مانهایمر^۹ و همکاران (۲۰۲۴ب) نیز در شماره ویژه، نمونه‌هایی از چگونگی تعامل اخلاقی و مسئولانه متخصصان با ابزارها و سامانه‌های هوش مصنوعی را ارائه نمودند. هفت مسئله اخلاقی کلیدی در این مطالعات موردی روشن گردید: حریم خصوصی،

^۱Wang & Liu

^۲Badawy

^۳Data-Driven Framework

^۴PESTEL

^۵Pansoni

^۶Okoroma

^۷Ali

^۸Teel

^۹Mannheimer

رضایت (استفاده از داده‌ها بدون رضایت صریح)، دقت (چالش‌های اخلاقی ناشی از دقت داده‌ها)، ملاحظات نیروی کار (مسائل اخلاقی مرتبط با کار)، شکاف دیجیتال، سوگیری و شفافیت. مرور پیشینه‌های پژوهشی اهمیت و ضرورت مدیریت ریسک در کتابخانه‌ها، آرشیوها و مراکز فرهنگی را نشان می‌دهد. شناسایی ریسک‌ها و اولویت‌بندی آن‌ها در بافت کتابخانه‌های دیجیتال (مقرببی منطقی، ۱۳۹۸؛ صمیعی و همکاران، ۱۴۰۱) و ارائه الگوی مدیریت ریسک (محمدصالحی و همکاران، ۱۴۰۱) در پژوهش‌های داخل کشور بیشتر مدنظر قرار گرفته است. پیشینه‌های بررسی شده در خارج از کشور نیز مبین اهمیت و چالش‌های مرتبط با فناوری‌های نوین، به‌ویژه هوش مصنوعی، در حفاظت و مدیریت میراث فرهنگی و مستند هستند. این پژوهش‌ها مبین آن است که فناوری‌های نوین می‌توانند به بهبود فرآیندهای مدیریتی کمک کنند و نیاز به رویکردهای جدید در حفاظت از میراث فرهنگی را نمایان می‌سازند (Ames & Lewis, 2020). توجه به مدیریت ریسک در زمینه هوش مصنوعی و حفاظت از میراث مستند (Badawy, 2023; Korka et al., 2022) Korka et al., 2024; Wang & Liu, 2023; Ali et al., 2024) از ابعاد اصلی پژوهش‌های خارجی است. به‌طور کلی، پیشینه‌های اخیر (Korka et al., 2024; Wang & Liu, 2023; Badawy, 2023; et al., 2022) بر ضرورت توجه به مدیریت ریسک، اخلاقیات و چالش‌های مرتبط با استفاده از هوش مصنوعی در حفاظت از میراث مستند تأکید دارند.

با توجه به مرور انجام‌شده، تاکنون پژوهشی که به‌طور جامع و با روش کیفی و مبتنی بر نظر خبرگان، به شناسایی و طبقه‌بندی ریسک‌های هوش مصنوعی در حفاظت از میراث مستند در یک سازمان ملی (همچون سازمان اسناد و کتابخانه ملی ایران) پرداخته، انجام نشده است. همچنین، از منظر بومی، پژوهشی که این ریسک‌ها را در بافت فرهنگی، حقوقی و مدیریتی ایران واکاوی کرده باشد، در پیشینه‌ها مشاهده نمی‌شود. بنابراین، پژوهش حاضر درصدد است تا این خلأ را با تمرکز بر سازمان اسناد و کتابخانه ملی ایران (به‌عنوان کتابخانه مادر و الگو) پر کند.

روش پژوهش

پژوهش حاضر از نظر هدف، کاربردی، از نظر رویکرد کیفی و از نظر نحوه گردآوری داده‌ها، از نوع مصاحبه نیمه‌ساختاریافته اکتشافی است. در این پژوهش تحلیل محتوای کیفی مبتنی بر رویکرد گرانهم و لاندمن^۱ (۲۰۰۴) استفاده شد. در این روش، داده‌های متنی حاصل از مصاحبه‌ها بدون پیش‌فرض نظری اولیه، کدگذاری شده و کدها در فرآیندی استقرایی در قالب مقوله‌ها و درون‌مایه‌ها سازماندهی می‌شوند. پیش از انجام مصاحبه‌ها، پژوهشگر با مرور مستندات و استانداردهای موجود مدیریت ریسک (مانند ISO/IEC 38507, ISO 31000) و چهارچوب‌های مدیریت ریسک (NIST, 2023) با مفاهیم عمومی مدیریت ریسک آشنا شد، اما این مستندات صرفاً برای طراحی راهنمای مصاحبه (مانند تعیین ابعاد کلی ریسک برای پرسش از خبرگان) به کار رفتند و هیچ نقشی در فرایند کدگذاری و مقوله‌بندی استقرایی نداشتند؛ بدین معنا که کدها و مقوله‌ها کاملاً از دل متن مصاحبه‌ها و بدون تحمیل مقوله‌های ازپیش‌تعیین‌شده استخراج شدند. جامعه پژوهش متشکل از ۲۰ خبره از کارشناسان علم اطلاعات و دانش‌شناسی و آرشیویست‌ها، متخصصان فناوری اطلاعات بود که تجربه و دانش کافی در زمینه‌های مرتبط با هوش مصنوعی و حفاظت از میراث مکتوب و مستند را داشتند. انتخاب نمونه به‌صورت هدفمند انجام شد تا اطمینان حاصل شود که شرکت‌کنندگان تجربه و دانش کافی در زمینه‌های مرتبط با هوش مصنوعی و حفاظت از میراث مکتوب و مستند را دارند. در انتخاب نمونه، رسیدن به اشباع، ملاک عمل بود.

^۱Graneheim & Lundman

در فرایند انتخاب خبرگان، سه معیار اصلی مدنظر قرار گرفت. نخست، تجربه حرفه‌ای افراد در حوزه‌های علم اطلاعات، کتابداری، آرشیو، فناوری اطلاعات و هوش مصنوعی، به‌ویژه متخصصان فعال در نهادهای معتبر مانند سازمان اسناد و کتابخانه ملی ایران، پژوهشگاه علوم و فناوری اطلاعات ایران، مرکز تحقیقات کامپیوتری علوم اسلامی و شرکت‌های فعال در حوزه هوش مصنوعی؛ دوم، سطح تحصیلات آنان، به‌ویژه دارا بودن حداقل مدرک کارشناسی ارشد در رشته‌های مرتبط و برخورداری از دانش تخصصی لازم در زمینه‌های هوش مصنوعی و حفاظت از میراث مکتوب و مستند؛ سوم، سوابق مشارکت در طرح‌های ملی و پروژه‌های مشابه که توانایی و تجربه عملی آنان را در این حوزه‌ها تأیید می‌کند. در جدول ۱، اطلاعات دقیق جامعه پژوهش ارائه می‌شود.

جدول ۱. توزیع فراوانی مشخصات خبرگان

ویژگی	دسته‌بندی	فراوانی	درصد فراوانی
مدرک تحصیلی	دکتری	۱۲	۶۰
	کارشناسی ارشد	۷	۳۵
	کارشناسی (هم‌زمان سطح ۴ حوزه)	۱	۵
رشته تحصیلی	علم اطلاعات و دانش‌شناسی	۷	۳۵
	فناوری اطلاعات / کامپیوتر	۵	۲۵
	هوش مصنوعی	۳	۱۵
	سایر (بیوانفورماتیک، صنایع، برق و...)	۵	۲۵
محل کار	سازمان اسناد و کتابخانه ملی ایران	۶	۳۰
	پژوهشگاه علوم و فناوری اطلاعات ایران (ایرانداک)	۵	۲۵
	دانشگاه‌ها و مراکز آموزش عالی	۳	۱۵
	شرکت‌های دانش‌بنیان هوش مصنوعی	۳	۱۵
	سایر مراکز (مرکز تحقیقات کامپیوتری علوم اسلامی و مؤسسه نجم، مرکز اسناد انقلاب اسلامی)	۳	۱۵
سابقه کار مرتبط با تحلیل داده و هوش مصنوعی	۱-۳ سال	۶	۳۰
	۴-۶ سال	۷	۳۵
	۷-۹ سال	۲	۱۰
	۱۰ سال و بالاتر	۵	۲۵
تخصص اصلی	توسعه فناوری‌های نوین و هوش مصنوعی	۷	۳۵
	عضو هیئت‌علمی / پژوهشگر حوزه فناوری‌های نوین و هوش مصنوعی	۶	۳۰
	پردازش زبان طبیعی و داده‌کاوی	۵	۲۵
	سایر تخصص‌ها (متخصص و مدرس هوش مصنوعی، سیاستگذاری علم و فناوری، پردازش صوت و تصویر)	۴	۲۰

مصاحبه‌ها به‌صورت نیمه‌ساختاریافته و به مدت ۴۵ تا ۷۵ دقیقه اجرا شدند و با اجازه کتبی مشارکت‌کنندگان ضبط صوتی شدند. سپس تمامی مصاحبه‌ها کلمه‌به‌کلمه پیاده‌سازی (رونویسی) شدند و برای حفظ محرمانگی، به هر مشارکت‌کننده یک کد اختصاصی (p1 تا p20) داده شد. تحلیل داده‌ها با استفاده از نرم‌افزار مکس کیودی^۱ و در مراحل زیر انجام گرفت:

^۱MAXQDA

نخست، متن مصاحبه‌ها چندین بار خوانده شد (غوطه‌وری در داده)؛ سپس واحدهای معنایی (جملات یا پاراگراف‌های حاوی یک مفهوم مستقل) به صورت کدگذاری باز شناسایی و برچسب‌گذاری شدند که در این مرحله حدود ۲۵۰ کد اولیه به دست آمد.

در گام بعد، کدها بازبینی شده و کدهای هم‌معنا با یکدیگر ادغام شدند و کدهای نامرتب با هدف پژوهش حذف شدند. پس‌از آن، کدهای باقیمانده براساس تشابه معنایی در زیرمقوله‌ها دسته‌بندی شدند و زیرمقوله‌های مرتبط در قالب مقوله‌های اصلی (شامل ریسک‌های فناوری، ریسک‌های حقوقی و امنیتی، ریسک‌های مدیریتی و حکمرانی و ریسک‌های فرهنگی - اجتماعی) سازماندهی شدند. در پایان این فرایند، ۱۶۷ کد نهایی استخراج شدند.

برای اعتباربخشی به یافته‌ها از معیارهای گوبا و لینکلن^۱ (اعتبار^۲ انتقال‌پذیری^۳، اطمینان‌پذیری^۴ و تأییدپذیری^۵) استفاده شد. برای حصول اعتبار، فرایند تحلیل شامل بازبینی‌های مکرر توسط محقق، کدگذاری مجدد رونوشت‌ها و بررسی مداوم کدهای استخراج شده بود. برای افزایش قابلیت انتقال یافته‌ها، توصیف غنی از زمینه پژوهش و مشخصات دقیق مشارکت‌کنندگان (جدول ۱) ارائه و تنوع خبرگان از سازمان‌ها، تخصص‌ها و سوابق گوناگون رعایت شد. برای اطمینان‌پذیری، تمامی مراحل پژوهش (از طراحی پرسش‌ها تا کدگذاری نهایی) مستندسازی شد. برای تأییدپذیری، یک متخصص بیرونی دیگر فرایند کدگذاری و نتیجه‌گیری را بررسی و اطمینان حاصل کرد که یافته‌ها ریشه در داده‌ها دارند. توافق بین کدگذاران با روش بحث و اجماع و تأیید ارزیاب بیرونی حاصل شد. تمامی مراحل پژوهش با رعایت اصول اخلاقی (اخذ رضایت آگاهانه، تضمین محرمانگی و حق انصراف مشارکت‌کنندگان) انجام پذیرفت.

یافته‌ها

مبتنی بر پرسش اصلی پژوهش، جامعه خبرگان پیرامون چالش‌ها و مخاطرات ناشی از به‌کارگیری هوش مصنوعی در حوزه حفاظت از میراث مستند مورد پرسش قرار گرفتند. یافته‌های حاصل از مصاحبه، که نمایانگر طیف وسیعی از نگرانی‌های تخصصی است، در جدول (۲) به‌طور مشروح ارائه شده است. شایان‌ذکر است که فراوانی در جدول (۲) نشان‌دهنده تعداد دفعات تکرار کدهای مرتبط با هر ریسک در مجموع مصاحبه‌ها است، نه صرفاً تعداد مشارکت‌کنندگانی که به آن ریسک اشاره کرده‌اند. به عبارت دیگر، یک خبره ممکن است در چند نوبت مختلف از مصاحبه به ریسک خاصی اشاره کرده باشد و شدت یا گستردگی توجه خبرگان به هر ریسک را بهتر نشان می‌دهد.

جدول ۲. چالش‌ها و ریسک‌های مرتبط با هوش مصنوعی در حفاظت از میراث مستند

مقوله اصلی	زیرمقوله	تعریف/توضیح کوتاه	فراوانی
ریسک‌های فناوری	محدودیت، کیفیت و سازگاری داده‌ها	خطا در یادگیری و بازیابی ناشی از داده‌های نامرتب و ناقص ضرورت دقت و تناسب داده‌ها با تسک مورد نظر	۱۵

^۱Guba & Lincoln

^۲credibility

^۳transferability

^۴dependability

^۵confirmability

	نیاز به پیش‌پردازش، پاک‌سازی و برجسب‌گذاری دقیق داده‌ها پالایش، تطبیق و یکپارچه‌سازی داده‌ها برای تضمین صحت و یکدستی کفایت، جامعیت و کیفیت داده‌ها در فرآیند آموزش مدل		
۴	محدودیت سازگاری مدل‌های زبانی بزرگ و ابزارهای پردازش متن با زبان فارسی و سایر زبان‌های غیرانگلیسی	محدودیت‌های زبانی	
۸	یادگیری محدود مدل از جنبه‌ای خاص از پدیده تولید نتایج سوگیرانه یا ناقص به دلیل سوگیری داده‌ها تأثیر نامتوازن بودن داده‌های آموزشی بر دقت مدل ناتوانی در شناسایی موارد نادر یا استثنایی ارزیابی عملکرد مدل بر پایه معیارهای نامناسب ارائه تصویر غیرواقعی از کیفیت عملکرد مدل	سوگیری داده‌ها و ارزیابی نادرست مدل‌ها	
۵	وهم و تعمیم نادرست مدل‌های هوش مصنوعی بر اثر داده‌های نابرابر	وهم‌گویی ^۳	
۱۴	کاهش دقت عملکرد و بروز خطاهای الگوریتمی افزایش وابستگی سازمانی به فناوری‌های خودکار اختلال یا وقفه در کارکرد سامانه‌های هوشمند	کاهش قابلیت اعتماد فنی مدل‌ها و وابستگی به سیستم	
۱	ناکامی در انتقال دقیق دانش تخصصی به سامانه‌های هوش مصنوعی	ناکارآمدی عملیاتی انتقال دانش تخصصی به بستر هوش مصنوعی	
۲۵	ریسک نقض حق نشر و افشای داده‌های محرمانه در استفاده از داده‌های دارای محدودیت دسترسی ریسک حریم خصوصی	مسائل حقوقی، محرمانگی و دسترسی	ریسک‌های حقوقی و امنیتی
۱۲	حملات سایبری به سرورها دسترسی‌های غیرمجاز قطع یا اختلال در اتصال شبکه اجرای دستورات مخرب	ریسک‌های امنیتی و محدودیت‌های دسترسی	
۲۸	تأمین و نگهداری سخت‌افزارهای پردازشی قدرتمند (GPU پیشرفته، برق پایدار، سامانه خنک‌کننده) زیرساخت نرم‌افزاری و مدیریت داده‌ها (نصب و آموزش مدل‌ها روی سرور محلی، محدودیت اشتراک‌گذاری داده‌ها به دلیل امنیت) وابستگی به فناوری خارجی (ریسک‌های عملیاتی و تحریمی، دشواری در خرید سخت‌افزار و مدل‌های به‌روز) هزینه‌ها و ریسک‌های عملیاتی در پیاده‌سازی هوش مصنوعی نگهداری، به‌روزرسانی و نظارت مستمر بر مدل‌های هوش مصنوعی	چالش‌های تأمین زیرساخت سخت‌افزاری و نرم‌افزاری	ریسک‌های مدیریتی و حکمرانی
۲	دسترسی و قابلیت استفاده از ابزارهای هوش مصنوعی	دسترسی به ابزارهای هوش مصنوعی	

^۱data bias^۲imbalanced^۳hallucination

۷	کمبود مهارت فنی و تخصصی کمبود نیروی انسانی ماهر دشواری جذب و حفظ نیروی متخصص	ریسک دانش و مهارت منابع انسانی و پایداری نیروی متخصص	
۱۹	نبود اولویت‌بندی روشن برای به‌کارگیری هوش مصنوعی فقدان نیازسنجی نظام‌مند ضعف پایداری و مدیریت پروژه‌ها کمبود فرهنگ سازمانی مناسب ارزش‌گذاری ناکافی به فناوری‌های نوظهور	ضعف در برنامه‌ریزی، سیاست‌گذاری، اولویت‌بندی و مدیریت پروژه‌های هوش مصنوعی	
۴	انتظارات اغراق‌آمیز و غیرواقعی	انتظارات غیرواقع‌بینانه	
۲۰	مقاومت سازمانی و کندگی پذیرش فناوری‌های نوین کمبود مهارت و ترس از دشواری یادگیری فناوری ضعف همکاری بین بخشی	چالش‌های فرهنگی و سازمانی	ریسک‌های فرهنگی
۳	دقت و صحت اطلاعات ابهام در سهم هوش مصنوعی و انسان در تولید محتوا	چالش اصالت و اعتبار متون تولیدشده توسط هوش مصنوعی	اجتماعی
۱۶۷	جمع		

در ادامه نقل‌قول‌هایی از چالش‌ها و ریسک‌های بیان‌شده در مصاحبه‌ها ارائه می‌شود:

ریسک‌های فناوری

محدودیت، کیفیت و سازگاری داده‌ها

هسته مرکزی و بنیادین هر سامانه هوش مصنوعی را داده‌هایی تشکیل می‌دهند که مدل بر مبنای آن‌ها آموزش می‌بیند. در حوزه حفاظت از میراث مستند، این اصل از اهمیتی مضاعف برخوردار است، زیرا کیفیت خروجی‌های سیستمی مستقیماً بر اصالت، تفسیر و انتقال صحیح میراث تاریخی تأثیرگذار است. مقوله «کیفیت و سازگاری داده‌ها» به‌عنوان یکی از چالش‌های اساسی در مرحله آماده‌سازی داده‌های آموزشی مطرح می‌شود. هرگونه نقص، خطا، ناسازگاری یا کیفیت پایین در داده‌های ورودی، نه تنها منجر به آموزش ناکارآمد مدل و تولید نتایج غیرقابل اعتماد خواهد شد، بلکه پتانسیل ایجاد تحریف معنایی، تداوم سوگیری‌های تاریخی و در نهایت، به خطر انداختن تمامیت و اعتبار میراث مستند را به همراه دارد (Mannheimer et al., 2024a). محور اصلی گفت‌وگوهای مشارکت‌کنندگان نیز حول این موضوع می‌چرخید.

مشارکت‌کننده (۱) یکی از چالش‌های مهم در به‌کارگیری هوش مصنوعی برای بازیابی منابع کتابخانه‌ای را کیفیت و یکپارچگی داده‌ها عنوان کرد. ایشان بیان داشت: «به نظر من، پیش از استفاده از داده‌ها در سامانه‌های هوش مصنوعی، لازم است پاک‌سازی و یکسان‌سازی متادیتا انجام شود. در حال حاضر، داده‌ها یکپارچه نیستند، به‌عنوان مثال تفاوت‌های نوشتاری بین حروف فارسی و عربی، اشتباهات فاصله یا نیم‌فاصله و خطاهای ناشی از ورود سریع اطلاعات وجود دارد». مشارکت‌کننده (۲) نیز هم‌راستا با دیدگاه ایشان چنین اظهار داشت: «بسیاری از داده‌هایی که با آن‌ها کار می‌کنیم، ناقص یا نامنظم هستند؛ در متون فارسی - انگلیسی، وجود ترکیب زبان‌ها، فرمول‌ها و علائم خارجی، داده‌ها را کثیف می‌کند و نیازمند فاز پیش‌پردازش^۱ برای پاک‌سازی و یکدست‌سازی داده‌هاست. بااین‌حال، تجربه من نشان می‌دهد که با ظهور مدل‌های زبانی بزرگ (LLM)، اثر کیفیت داده‌ها

^۱pre-processing

نسبت به گذشته کمتر، و مدیریت آن ساده‌تر شده است؛ زیرا این مدل‌ها از قبل روی مجموعه داده‌های بزرگ آموزش دیده‌اند و توانایی تحلیل جمله‌ها و ویژگی‌ها را با مقاومت بیشتری نسبت به خطاهای جزئی دارند. با این وجود، هنگامی که بخواهیم این مدل‌ها را برای یک تسک خاص فاین‌تیون^۱ کنیم و از داده‌های محدود استفاده کنیم، کیفیت داده‌ها دوباره اهمیت پیدا می‌کند و باید داده‌ها دقیق و پاکسازی شده باشند تا مدل بتواند به‌طور مؤثر عملکرد موردنظر را ارائه دهد^۲.

مشارکت‌کننده (۴ و ۶) با توجه به تجربه عملیاتی خود بیان کرد: «الگوریتم‌ها و پردازش‌های هوش مصنوعی به‌خودی‌خود پیچیده نیستند، اما ایجاد یک دیتای تمیز و قابل استفاده، بخش عمده انرژی ما را مصرف کرد. بنابراین، پیش‌پردازش و پاکسازی زبان فارسی به‌سبب ویژگی‌های خاص خود (شامل نگارش‌های متفاوت، غلط‌های املائی، نیم‌فاصله‌ها و اختلافات یونیکد کاراکترها)، یکی از ریسک‌های مهم در کاربرد هوش مصنوعی محسوب می‌شود».

مشارکت‌کننده (۸) نیز به تفاوت کیفیت و سازگاری داده‌های موجود در سامانه‌هایی مانند نرم‌افزار رسا و کتابخانه دیجیتال در کتابخانه ملی اشاره کرد و گفت: «داده‌های موجود در سامانه‌هایی مانند نرم‌افزار رسا و کتابخانه دیجیتال ممکن است کیفیت یکسان نداشته باشند و نیازمند پالایش و هماهنگی با مدل‌های زبانی و سایر الگوریتم‌های هوش مصنوعی باشند»؛ مشارکت‌کننده (۴) نیز تصریح کرد: «حجم و تنوع داده‌ها برای آموزش حیاتی است؛ داده‌های کم، اغلب فایده خاصی ندارند و برای دستیابی به دقت بالاتر، افزایش داده‌ها معمولاً اولین راهکار است. با این حال، در بسیاری از موارد دسترسی به داده‌های گسترده ممکن نیست».

مشارکت‌کننده (۱۶) داده‌های کتابخانه ملی ایران را مرکز توجه قرار داد و بیان داشت: «در خصوص کتاب‌های کتابخانه ملی ایران، اگر منظور پردازش محتوای کتاب‌ها باشد، شاید داده‌های کتابخانه‌ها کافی به نظر برسد. اما باید توجه داشت که داده‌های زبانی یک ملت تنها به محتوای رسمی مانند کتاب‌ها محدود نمی‌شود. بلکه داده‌های غیررسمی مانند محتوای تولیدشده در شبکه‌های اجتماعی، پلتفرم‌های چت و فضای مجازی نیز بخش حیاتی و بزرگی از زبان زنده را تشکیل می‌دهند. البته نکته بنیادی‌تری نیز وجود دارد: حتی اگر ما داده‌های فارسی کافی را جمع‌آوری و پردازش کنیم، این برای ساختن یک مدل هوش مصنوعی پیشرفته کافی نیست. در پارادایم فعلی هوش مصنوعی، لازم است به‌صورت چندزبانه^۲ فکر کنیم. مدل‌های موفق، دانش را در یک فضای بُرداری مشترک یاد می‌گیرند که در آن معانی، مستقل از زبان خاص، کنار هم قرار می‌گیرند. بنابراین، مشکل اصلی ما تنها داشتن داده برای آموزش صحبت کردن به زبان فارسی نیست، بلکه داشتن داده‌های کافی و متنوع - از منابع چندزبانه - برای آموزش «درک مفهوم» توسط مدل است. تمام صحبت‌های من ناظر بر چالش کلان برای توسعه مدل‌های هوش مصنوعی بومی و کارآمد است».

با توجه به نقل قول‌های خبرگان، هسته مرکزی چالش‌های هوش مصنوعی در میراث مستند، کیفیت و یکپارچگی داده‌ها است. داده‌های ناقص، ناسازگاری‌های یونیکد، ترکیب زبان‌ها و خطاهای املائی در متون فارسی، فرایند پیش‌پردازش را زمان‌بر و هزینه‌بر ساخته است. همچنین، نبود داده‌های کافی برای زبان فارسی (به‌ویژه داده‌های غیررسمی و چندزبانه) به‌عنوان ریسک ساختاری برای توسعه مدل‌های بومی مطرح شده است.

محدودیت‌های زبانی

^۱ فاین‌تیونینگ (Fine-Tuning) در حوزه هوش مصنوعی و یادگیری ماشین به فرایندی گفته می‌شود که در آن مدلی از پیش آموزش دیده (Pretrained Model) براساس داده‌ها و نیازهای خاص یک پروژه، بهینه‌سازی و تنظیم دقیق می‌شود. به عبارت ساده‌تر، مدل پیش‌تر روی یک مجموعه داده بزرگ و عمومی آموزش دیده، و فاین‌تیونینگ به آن کمک می‌کند تا ویژگی‌ها و پاسخ‌های خود را برای کاربرد مشخصی شخصی‌سازی کند.

^۲ multilingual

یکی از چالش‌های محوری در به‌کارگیری فناوری‌های هوش مصنوعی برای حفاظت و پردازش میراث مستند ایران، محدودیت‌های زبانی ذاتی این فناوری‌هاست. هسته اصلی بسیاری از سامانه‌های هوش مصنوعی، به‌ویژه مدل‌های زبانی بزرگ، بر پایه داده‌های عظیم انگلیسی بنا شده و در نتیجه، این مدل‌ها در درک، پردازش و تولید محتوای زبان فارسی و دیگر زبان‌های رایج در گنجینه اسناد تاریخی ایران (مانند عربی و ترکی) با کاستی‌های جدی مواجه‌اند. مضامین مشترکی در بررسی مصاحبه‌ها قابل‌شناسایی بود. مشارکت‌کنندگان (۲، ۹، ۱۶) تصریح کردند: «مدل‌های موجود در بازار جهانی، عمدتاً بر پایه داده‌های انگلیسی آموزش دیده‌اند و در نتیجه، هنگام مواجهه با زبان فارسی، از کمیت و تنوع داده‌های آموزشی کمتری برخوردارند که این امر به ایجاد سوگیری در خروجی‌های آن‌ها می‌انجامد. هرچند مدل‌های بزرگی مانند چت‌جی‌پی‌تی^۱ ذاتاً چندزبانه^۲ توسعه یافته‌اند و به لطف حجم و توان پردازشی فوق‌العاده‌شان، عملکرد نسبتاً قابل قبولی در تولید متن^۳ به زبان فارسی دارند، اما زمانی که به وظایف تخصصی‌تر پردازش زبان طبیعی می‌رسیم، ضعف آن‌ها آشکار می‌شود. منظور من وظایفی فراتر از تولید متن است؛ مانند تولید Embedding‌های باکیفیت (تبدیل متن به فضای بُرداری معنادار)، استخراج رابطه‌های معنایی^۴ و پیوند موجودیت‌ها^۵. عملکرد مدل‌های قابل‌دسترس در این حوزه‌ها برای زبان فارسی به مراتب ضعیف‌تر است. ریشه این مشکل دوگانه است: اول، کمبود داده‌های متنی بزرگ و متنوع به زبان فارسی و دوم، عدم وجود زیرساخت‌های پردازشی قدرتمند و موتورهای جستجوی پیشرفته در سطح ملی برای گردآوری و پردازش این داده‌ها».

مشارکت‌کننده (۳) با توجه به تجربه خود بیان داشت: «بیشتر داده‌ها در قالب پی‌دی‌اف^۶ فارسی هستند، که برای استخراج متن مناسب نیستند. تبدیل پی‌دی‌اف به متن با اُسی‌آر^۷ فارسی مشکل است، زیرا نرم‌افزارهای موجود کیفیت پایین دارند و به نوع چاپ، فونت، رنگ و حتی جدول‌ها حساس هستند».

سوگیری داده‌ها^۸ و ارزیابی نادرست مدل‌ها

اگر داده‌های آموزشی، نمای کامل و عادلانه‌ای از واقعیت تاریخی و فرهنگی ارائه نکنند، خروجی مدل نه تنها مفید نخواهد بود، بلکه می‌تواند به تحریف تاریخ و تداوم بخشیدن به حذف‌شدگی گروه‌های به حاشیه رانده‌شده بینجامد (NIST, 2023). بررسی داده‌های مصاحبه نشان از همسویی دیدگاه برخی مشارکت‌کنندگان در این زمینه داشت. مشارکت‌کننده (۱) بیان داشت: «به نظر من وقتی از هوش مصنوعی استفاده می‌کنیم، سه بخش اصلی وجود دارد: داده، مدل و ارزیابی. اگر مجموعه داده‌ها کیفیت و تعادل لازم را نداشته باشند، مدل نمی‌تواند به درستی یاد بگیرد. به عنوان مثال اگر تنها اطلاعاتی که به مدل داده می‌شود جنبه‌ای خاص از یک موضوع را پوشش دهد، نتیجه سوگیری خواهد بود و مدل تصویری نادرست یا یک‌جانبه ارائه می‌دهد انتخاب معیارهای ارزیابی مناسب، مانند دقت^۹ و بازیابی^{۱۰} برای داده‌های نامتوازن و همچنین تعادل بخشی به داده‌ها پیش از آموزش مدل نیز، برای حصول اطمینان از عملکرد درست مدل ضروری است». مشارکت‌کننده (۶) نیز به نمونه مشابهی اشاره کرد و گفت: «مثالی که

^۱ ChatGPT

^۲ multilingual

^۳ text generation

^۴ relation extraction

^۵ entity linking

^۶ PDF

^۷ OCR

^۸ data bias

^۹ precision

^{۱۰} recall

اخیراً رخ داد، مربوط به دستورالعملی بود که باعث شد یک هوش مصنوعی (گروک^۱)، متعلق به شرکت xAI، زیرمجموعه ایکس متعلق به ایلان ماسک، درباره هولوکاست (نسل‌کشی سفیدپوستان در آفریقای جنوبی) پاسخ‌هایی نامربوط ارائه دهد. شرکت xAI این پاسخ را ناشی از «خطای برنامه‌نویسی» دانست که توسط یک کارمند خاطی ایجاد شده و عذرخواهی کرد. در جنگ‌های رسانه‌ای مانند مناقشه ایران و اسرائیل، به‌صورت عمدی سوگیری‌ها را در داده‌ها یا پاسخ‌های هوش مصنوعی مشاهده کردیم؛^۲ مشارکت‌کننده (۶) نیز تصریح کرد: «درباره چهره‌های سیاسی ملی، هوش مصنوعی ایرانی ممکن است سخنانی با بار تعریفی و حماسی ارائه دهد که در دیگر سامانه‌ها متفاوت است. نکته مهم دیگر، محدودیت مدل‌ها در ترکیب اطلاعات است؛ اگر مدل بخواهد چند دیدگاه متفاوت را تلفیق کند، ممکن است نتیجه‌ای نادرست یا گمراه‌کننده تولید شود. به همین دلیل، مسائل اخلاقی و کنترل سوگیری، بخشی حیاتی از توسعه سامانه‌های هوش مصنوعی در این حوزه هستند» (م. ۶). مرور نظر خبرگان در این زیرمقوله به دو نوع سوگیری اشاره دارد: سوگیری تاریخی (ناشی از ناقص بودن یا یک‌جانبه بودن داده‌های آموزشی) و سوگیری الگوریتمی (ناشی از جهت‌گیری‌های سیاسی یا فرهنگی در طراحی مدل). انتخاب معیارهای ارزیابی نامناسب برای داده‌های نامتوازن نیز به‌عنوان ریسک فرعی مطرح شد.

وهم‌گویی مدل‌های هوش مصنوعی

ریشه اصلی وهم‌گویی مدل‌های هوش مصنوعی اغلب در «نابرابری و کمیت داده‌های آموزشی» نهفته است. هنگامی که یک مدل، با حجم ناکافی یا مجموعه‌ای نامتوازن از داده‌ها (مانند اسناد متعلق به دوره تاریخی خاص یا یک زبان محدود) آموزش می‌بیند، برای پر کردن خلأهای دانش خود، به «حدس‌زنی» و «تولید» اطلاعات می‌پردازد (Yin, 2024). اشتراک نظر قابل توجهی در صحبت‌های مشارکت‌کنندگان در این خصوص دیده می‌شود. مشارکت‌کننده (۱) خاطرنشان ساخت: «زمانی که داده‌های آموزش مدل، غیرساختاریافته باشند، ممکن است مدل برداشت‌های نادرستی داشته باشد، مانند اشتباه گرفتن نام‌های مشابه (به‌عنوان مثال حسن و حسین) یا تفسیر نادرست متن. اما اگر داده‌ها ساختاریافته باشند و روابط بین موجودیت‌ها به‌درستی تعریف شده باشد، مانند گراف دانش، دقت بازیابی اطلاعات بسیار افزایش می‌یابد و احتمال خطا کاهش می‌یابد». مشارکت‌کننده (۳) نیز با توجه به تجربه خود اظهار داشت: «سازمان‌ها معمولاً توان خرید سخت‌افزارهای پیشرفته را ندارند و بنابراین به سراغ مدل‌های کوچک متن‌باز می‌روند. این مدل‌ها یک محدودیت اساسی دارند و داده‌های بومی ما را ندیده‌اند. به همین دلیل، وقتی چنین مدل‌هایی با پرسشی روبه‌رو شوند که به این داده‌ها نیاز دارد، عملاً پاسخ دقیقی نمی‌دهند و به‌جای آن دچار «توهم» می‌شوند».

مشارکت‌کنندگان (۷ و ۱۰) به یکی از ریسک‌های مهم، یعنی بحث اعتمادپذیری سامانه‌ها اشاره کرده و گفتند: «در حوزه چت‌بات‌ها، پرسش اصلی این است که خروجی تا چه اندازه برای پژوهشگران قابل اتکا است. برای کاهش این ریسک، ما تلاش کرده‌ایم با بهره‌گیری از فناوری مانند RAG^۲ (تولید تقویت‌شده بازیابی)، دامنه بازیابی اطلاعات را به منابع معتبر خودمان محدود کنیم تا دقت و اعتبار خروجی‌ها افزایش یابد».

طبق نظر خبرگان، ریشه اصلی وهم‌گویی در «نابرابری و کمبود داده‌های آموزشی» است. مدل برای پر کردن خلأ دانش خود، به حدس‌زنی و تولید اطلاعات جعلی می‌پردازد. کاهش این ریسک با محدود کردن دامنه بازیابی به منابع معتبر (RAG) و ساختاردهی داده‌ها در قالب گراف دانش پیشنهاد شده است.

کاهش قابلیت اعتماد فنی مدل‌ها و وابستگی به سیستم

^۱Grok

^۲Retrieval-Augmented Generation

با خودکارسازی فرایندها، دانش و مهارت‌های سنتی مدیریت اسناد ممکن است به تدریج کمرنگ شود و سازمان برای انجام امور اصلی خود به طور فزاینده‌ای به فناوری وابسته شود و هرگونه قطعی، باگ نرم‌افزاری یا حمله سایبری و خطای اجتناب‌ناپذیر الگوریتمی که منجر به ازکارافتادن سامانه هوش مصنوعی شود، می‌تواند به دلیل همین وابستگی بالا، کل چرخه خدمات و مدیریت محتوای آرشیو را متوقف کند. مشارکت‌کننده (۱۴) یکی از چالش‌های جدی در استفاده از فناوری‌های هوش مصنوعی را مسئله دقت سامانه‌ها دانست و اشاره کرد: «این ابزارها همواره با درصدی از خطا همراه هستند و به نظر من، بروز خطا در چنین سامانه‌هایی امری اجتناب‌ناپذیر است». مشارکت‌کننده (۱۲) نیز یکی از چالش‌های مهم را خطاهای الگوریتمی دانست. ایشان تصریح کرد: «سامانه‌ها گاهی در شناسایی متون به زبان‌های قدیمی دچار اشتباه می‌شوند. این مسئله می‌تواند دقت تحلیل‌ها را کاهش دهد و بر کیفیت بازیابی اطلاعات تأثیر منفی بگذارد». مشارکت‌کننده (۸) نیز هشدار داد: «وقتی بحث استفاده از هوش مصنوعی در آرشیوهای دارای منابع حساس، به ویژه اسناد تاریخی، مطرح می‌شود، باید دقت فراوان به کار گرفته شود. این اسناد ممکن است از دوره‌های مختلف باشند و در دسترس قرار گرفته باشند، اما پردازش آن‌ها با هوش مصنوعی می‌تواند با ریسک‌های قابل توجهی همراه باشد. به عنوان مثال، ممکن است اطلاعات مهمی از دست برود یا الگوریتم، تفسیر اشتباه ارائه دهد». به باور مشارکت‌کننده (۲) نیز «برخی ابزارهای هوش مصنوعی به طور تبلیغی ادعا می‌کردند که توانایی پردازش زبان فارسی را دارند، اما در عمل نتایج مطلوب و قابل اعتماد ارائه نمی‌کردند».

مشارکت‌کنندگان (۵، ۹ و ۱۹) معتقدند: «اعتماد و وابستگی بیش‌ازحد به هوش مصنوعی می‌تواند مشکل‌ساز باشد. ماشین‌ها ممکن است اشتباه کنند یا از کار بیفتند و اگر همه چیز به آن‌ها وابسته باشد، اختلال گسترده رخ می‌دهد». مشارکت‌کننده (۳) نیز یکی از پرسش‌های مهم درباره این مدل‌ها را میزان قابلیت اتکای به آن‌ها عنوان کرد. ایشان تصریح کرد: «در فضای عمومی اغلب درباره توانایی‌های چشمگیرشان صحبت می‌شود، اما وقتی وارد سطح عملیاتی می‌شویم، ضعف‌های جدی آشکار می‌شود». مشارکت‌کننده (۸) نیز به پیامد وابستگی کاربران به این ابزارها اشاره کرد و گفت: «یکی از چالش‌های مهم استفاده از هوش مصنوعی، تبلی فکری کاربران است. کاربران به جای جستجو و تحلیل منابع خود، انتظار دارند که چت‌بات‌ها تمام پرسش‌های آن‌ها را پاسخ دهند. این مسئله باعث کاهش مراجعه مستقیم به منابع و کاهش توانایی تفکر مستقل می‌شود». به باور مشارکت‌کننده (۱۳) «ابزارهای هوش مصنوعی عامل سطحی شدن دانش و کم‌عمق شدن پژوهش‌ها به واسطه اتکای بیش‌ازحد به ابزارهایی مانند چت‌بات‌ها شده است».

در این زیرمقاله دو نگاه متفاوت در میان خبرگان دیده شد: گروهی بر خطاهای اجتناب‌ناپذیر الگوریتمی و اختلال در سامانه تأکید داشتند و گروهی نگران وابستگی بیش‌ازحد سازمانی و کمرنگ شدن مهارت‌های سنتی بودند. «تبلی فکری کاربران» و «سطحی شدن پژوهش‌ها» از پیامدهای وابستگی افراطی به چت‌بات‌ها ذکر شد.

ناکارآمدی عملیاتی انتقال دانش تخصصی به بستر هوش مصنوعی

دانش متخصصان حوزه کتابخانه و آرشیو، شامل شناخت زمینه تاریخی، شیوه‌های تفسیر محتوا و ظرافت‌های ارزیابی منابع و اسناد غالباً از نوع دانش ضمنی است که تجربی است و نظام‌مندسازی آن دشوار می‌باشد. مشارکت‌کننده (۱۴) در این زمینه چنین بیان کرد: «گاهی برای بازتعریف دانش تخصصی موجود با کمک هوش مصنوعی، نتیجه به گونه‌ای می‌شود که کاربرد و کارکرد عملی آن از دست می‌رود. این مسئله را دست‌کم در برخی پروژه‌هایی که خودمان انجام داده‌ایم، مشاهده کرده‌ایم. همچنان یکی از دغدغه‌های اصلی، عملیاتی شدن دقیق فرایندهایی است که اکنون به صورت دستی انجام می‌شوند».

ریسک‌های حقوقی

مسائل حقوقی، محرمانگی و دسترسی

اسناد و منابعی که در گنجینه کتابخانه‌ها و آرشیوها نگهداری می‌شوند، اغلب مشمول قوانین مالکیت فکری و حق نشر، مقررات حفاظت از داده‌های شخصی و طبقه‌بندی محرمانگی هستند. استفاده از این داده‌های حساس در چرخه عمر هوش مصنوعی (از جمع‌آوری و آموزش مدل تا استقرار و بهره‌برداری)، بدون تدوین سیاست‌های شفاف و راهبردهای محکم، سازمان را در معرض ریسک نقض حقوق مالکیت و افشای غیرعمدی اطلاعات محرمانه قرار می‌دهد. اشتراک نظر قابل‌توجهی در تحلیل مصاحبه‌ها در این زمینه مشاهده شد. مشارکت‌کنندگان (۱، ۸، ۱۴، ۱۶) یکی از مهم‌ترین ریسک‌های حوزه دیجیتال‌سازی میراث مستند را موضوع مالکیت معنوی و کپی‌رایت برشمردند. ایشان تصریح کردند: «بسیاری از کتاب‌ها و منابعی که امروز دیجیتال می‌شوند، ممکن است بدون اخذ اجازه از صاحبان اثر یا ناشران وارد پایگاه‌های داده و حتی مدل‌های زبانی بزرگ شوند. کاربر احتمالاً می‌تواند با استفاده از پرامپت‌های خاص در مدل‌های زبانی بزرگ، متن کامل اثر را استخراج کند و این امر عملاً به نقض گسترده حقوق ناشران و نویسندگان می‌انجامد». ایشان در بخش دیگری از صحبت‌های خود به خلاهای ملی اشاره کرد و گفت: «برخلاف بسیاری از کشورها که سازمان‌های مدیریت جمعی برای حق مؤلف دارند، در ایران چنین نهادی وجود ندارد یا در مراحل اولیه است. این سازمان‌ها در کشورهای دیگر واسطه‌ای میان ناشر، کتابخانه و نهادهای فرهنگی هستند و علاوه بر تضمین حقوق پدیدآورنده، امکان بهره‌برداری قانونی و منظم از آثار را فراهم می‌کنند. اگر سازمان معتبر و عمومی مانند کتابخانه ملی مسئول این فرایند شود، اعتماد و اطمینان بیشتری حاصل خواهد شد».

مشارکت‌کننده (۱۹) نیز یکی از مسائل اصلی در استفاده از هوش مصنوعی برای پردازش محتوا را حقوق مؤلف و ناشر دانست. ایشان با ذکر مثال عینی عنوان کرد: «فرض کنید ما متن کامل کتاب‌هایی که برای فیپا ارائه می‌شوند را دریافت کنیم و بخواهیم در مدل‌های زبانی از آن‌ها استفاده کنیم، در این صورت نقض قانونی رخ می‌دهد، حتی اگر کار ما صرفاً تحقیقاتی باشد و داده‌ها منتشر نشوند».

مشارکت‌کننده (۷) نیز دغدغه‌مندی سازمان خود را نسبت به بحث مالکیت داده‌ها مورد تأکید قرار داد و بیان کرد: «همواره سعی کرده‌ایم با صاحبان آثار و کاربران توافق‌نامه‌های مشخصی داشته باشیم و مجوزهای لازم را به‌طور رسمی اخذ کنیم. این مسئله در دو سطح قابل طرح است: نخست در فرآیند دیجیتال‌سازی و عرضه عمومی محتوا که ملاحظات حقوقی پایه‌ای دارد و دوم در استفاده از این داده‌ها برای آموزش مدل‌های هوش مصنوعی. در این سطح اخیر، پرسش اصلی این است که آیا داده‌ها می‌توانند در فرآیند یادگیری ماشین مورد استفاده قرار گیرند یا خیر. برخی معتقدند که چون داده‌ها به‌صورت ضمنی وارد مدل می‌شوند، خطر خاصی وجود ندارد؛ اما دیدگاهی دیگر هشدار می‌دهد که ممکن است این داده‌ها در خروجی مدل بازسازی شوند و به این ترتیب در معرض سوءاستفاده قرار گیرند. به همین دلیل ما این موضوع را بسیار جدی گرفته‌ایم».

مشارکت‌کننده (۳) به بحث محدودیت دسترسی به داده‌ها در سطح کلان اشاره کرد و گفت: «بسیاری از داده‌ها، مثل اطلاعات پرونده‌های قضایی یا داده‌های وزارتخانه‌ها، به دلیل محرمانه بودن، دسترسی دشواری دارند».

تحلیل مصاحبه‌ها موضوع محرمانگی اطلاعات و حریم خصوصی را هم مهم جلوه داد. مشارکت‌کنندگان (۵ و ۱۵) تأکید کردند: «برخی داده‌ها محرمانه هستند و پیش از استفاده در مدل‌ها، باید پاک‌سازی شوند تا از انتشار ناخواسته اطلاعات جلوگیری شود». مشارکت‌کننده (۱۶) تصریح کرد: «مدل‌های زبانی بزرگ قادرند مانند انسان عمل کنند و از این توانایی می‌توان برای فعالیت‌های مخرب استفاده کرد؛ به عنوان مثال حدس زدن یا پیشنهاد رمزهای عبور با اتکا به دیتاهای موجود. همچنین بسیاری از سازوکارهای ساده وریفیکیشن انسانی - آنچه برای تشخیص انسان از ربات به کار می‌رود - اکنون قابل حل شدن توسط این مدل‌ها است که از منظر حفاظت خطر جدی است». ایشان یکی دیگر از مهمترین چالش‌ها را حریم خصوصی دانست و گفت: «با ورود گسترده چت‌بات‌ها و استفاده عمومی از آن‌ها، بسیاری از اطلاعات شخصی افراد به خارج از کشور منتقل شده‌اند و

در نتیجه، ما حکمرانی مشخصی روی داده‌های داخلی خود نداریم. برای مثال، داده‌هایی که شرکت‌های بزرگ مانند گوگل از افراد جمع‌آوری می‌کنند، بسیار فراتر از آن چیزی است که حتی اعضای خانواده می‌دانند؛ این شامل موقعیت، علاقه‌مندی‌ها و تعاملات افراد می‌شود.

مشارکت‌کننده (۸) به بحث حریم خصوصی در کتابخانه ملی اشاره کرد و گفت: «هنگام استفاده از داده‌های عضویت کاربران، باید مشخص باشد که آیا فرد اجازه داده است اطلاعاتش برای تحلیل‌های هوشمند استفاده شود یا خیر. تنها در چهارچوب قانونی مشخص، می‌توان از داده‌ها بهره برد».

طی این زیرمقوله سه سطح ریسک حقوقی شناسایی شد: نقض حق نشر (استفاده از داده‌های دارای کپی‌رایت در آموزش مدل‌ها)، حریم خصوصی (احتمال بازسازی اطلاعات شخصی در خروجی مدل) و فقدان حکمرانی داده (نبود سازمان مدیریت جمعی برای حق مؤلف در ایران). خبرگان بر ضرورت اخذ مجوز رسمی و شفافیت در زنجیره استفاده از داده‌ها تأکید داشتند.

ریسک‌های امنیتی و محدودیت‌های دسترسی

پیاده‌سازی سامانه‌های هوش مصنوعی، سازمان را در معرض پارادایم امنیتی جدیدی قرار می‌دهد و به‌طور هم‌زمان دامنه تهدیدات سایبری را نیز گسترش می‌دهد. این ریسک‌ها ناظر بر تهدیدات داخلی و خارجی است که محرمانگی، یکپارچگی و دسترس‌پذیری (سه رکن اصلی امنیت اطلاعات) داده‌های حساس کتابخانه‌ای و آرشیوی و خود سامانه‌های هوشمند را هدف قرار می‌دهند (NIST, 2023). برای یک نهاد حافظ میراث ملی مانند سازمان اسناد و کتابخانه ملی، که منابع آن اغلب منحصر به فرد و غیرقابل جایگزین هستند، پیامدهای نقص امنیتی می‌تواند به معنای تحریف تاریخ، نقض حریم خصوصی افراد یا حذف دائمی بخشی از حافظه ملی باشد. به باور مشارکت‌کننده (۹) «یکی از ریسک‌های مهم در استفاده از فناوری‌های هوش مصنوعی، وابستگی به ابزارهای غیربومی است. به‌ویژه در حوزه محتوای فارسی، باید محدودیت‌ها و حدود دسترس‌پذیری را دقیقاً مشخص کنیم تا در صورت بروز آسیب‌های امنیتی، اطلاعات خصوصی و داده‌های حساس، دچار افشا یا نفوذ نشوند». در این خصوص مشارکت‌کننده (۱۶) اظهار داشت: «خیلی از سازمان‌ها تمایل ندارند داده‌های خود را به مدل‌های زبانی بزرگ خارجی بسپارند. به‌عنوان مثال، کتابخانه ملی ممکن است داده‌های حساس و خصوصی‌ای داشته باشد که قابل‌ارسال به مدل زبانی بزرگ خارجی نیست. اگر سازمان زیرساختی داشته باشد که بتواند مدل زبانی را در داخل خود میزبانی کند، می‌توان این مدل را به‌طور مؤثر با فعالیت‌های سازمان یکپارچه کرد و از آن بهره برد؛ اما این امر نیازمند سرمایه‌گذاری و هزینه‌های بالایی است».

مشارکت‌کننده (۵) نیز به تهدیدهای سایبری اشاره کرد و گفت: «سامانه‌های هوش مصنوعی می‌توانند هدف هک یا حملات قرار گیرند و این موضوع ممکن است باعث اختلال در عملکردشان یا سوءاستفاده از اطلاعات ارزشمند شود». مشارکت‌کننده (۲) به محدودیت‌های امنیتی در سازمان‌های دولتی تأکید کرد و بیان کرد: «بسیاری از سازمان‌ها به‌منظور جلوگیری از نفوذ و حفاظت از داده‌ها، دستورالعمل‌ها و رویه‌های سخت‌گیرانه‌ای دارند که دسترسی به سرورها و سرویس‌ها را محدود می‌کند». مشارکت‌کننده (۴) معتقد است: «از نظر امنیتی، نگرانی عمده زمانی مطرح می‌شود که داده‌های حساس یا محرمانه در اختیار داشته باشیم. اما داده‌های عمومی، مانند کتاب‌ها، پایان‌نامه‌ها، مقالات و سایر منابع منتشرشده، از نظر امنیتی مشکل‌ساز نیستند». مشارکت‌کننده (۱۰) نیز به دستورالعمل‌های مخرب اشاره کرد و هشدار داد: «در مدل‌های زبانی وقتی کاربر پرسشی می‌پرسد، مدل

زبانی این درخواست را به یک کوثری تبدیل می‌کند و آن را اجرا می‌کند. در این مرحله، ممکن است این کوثری دستخوش تزریق یا دستور مخرب شود و عملیاتی را انجام دهد که مجاز نیست یا خطرناک است.^۱

برخی دیگر از مشارکت‌کنندگان نظر متفاوتی دربارهٔ ریسک‌های امنیتی داشتند. مشارکت‌کنندگان (۶ و ۱۷) معتقد بودند: «از منظر امنیتی، هوش مصنوعی به‌طور مستقیم مسئول حفظ امنیت داده‌ها نیست، چراکه وظیفه اصلی آن، پردازش داده‌ها است. داده‌ها باید در بستری امن ذخیره شوند و هوش مصنوعی صرفاً با دریافت این داده‌ها یا ورودی‌های کاربر، فرایند پردازش را انجام داده و خروجی ارائه می‌دهد. هر استاندارد امنیتی مربوط به حفاظت اطلاعات باید بخشی از چهارچوب امنیتی و عملکردی هوش مصنوعی نیز باشد». مشارکت‌کننده (۱۰) نیز اذعان داشت: «از نظر من مسئله بستگی دارد به اینکه مدل زبانی بزرگ را دقیقاً کجا و برای چه کاری به کار ببرید. اگر مدل صرفاً وظیفه پاسخ‌گویی داشته باشد - یعنی تنها به پرسش‌ها جواب بدهد و به فایل‌ها دسترسی خواندنی داشته باشد - خطر نسبتاً محدود است. اما وقتی مدل به انجام «اکشن» و تعامل با سامانه‌ها واگذار شود (به‌عنوان مثال دانلود، کپی یا حذف فایل‌ها)، آنگاه ریسک‌های امنیتی جدی پدید می‌آیند. یکی از مشکلات کلیدی، تعیین و اعمال سطوح دسترسی است: در پایگاه داده‌های معمولی می‌توان مجوزها را تفکیک کرد، اما در پیاده‌سازی‌های مبتنی بر مدل زبانی بزرگ مدیریت دقیق این‌که چه پرسش‌هایی مجاز به پاسخ‌گویی هستند و چه پرسش‌هایی نه، بسیار دشوارتر است. به‌علاوه، محدودیت‌ها را می‌توان به‌راحتی با بازنگارش یا تغییر گویش پرسش، دور زد و از مدل، پاسخ‌هایی گرفت که نباید ارائه شوند».

در این زیرمقاله، خبرگان به حملات سایبری، دسترسی‌های غیرمجاز، قطع شبکه، تزریق دستورات مخرب و ریسک وابستگی به ابزارهای غیرومی اشاره کردند. همچنین برخی خبرگان این ریسک‌ها را قابل مدیریت با دستورالعمل‌های امنیتی معمول دانستند.

ریسک‌های مدیریتی و حکمرانی

چالش‌های تأمین زیرساخت سخت‌افزاری و نرم‌افزاری

پیاده‌سازی فناوری‌های پیشرفته هوش مصنوعی در حوزه حفاظت از میراث مستند، باوجود وعده‌های تحول‌آفرین، با چالش‌های جدی در حوزه‌های مالی، زیرساخت فناوریانه و حکمرانی مواجه است. این چالش‌ها به‌عنوان پیش‌نیازهای اساسی مطرح می‌شوند که عدم‌توجه به آن‌ها می‌تواند کل پروژه را در مرحله طراحی یا مرحله آزمایشی با شکست مواجه کند. اشتراک نظر قابل توجهی در تحلیل مصاحبه‌ها حول این محور دیده می‌شد. مشارکت‌کنندگان (۱ و ۱۴) یکی از چالش‌های اصلی در به‌کارگیری هوش مصنوعی برای حفاظت و بازیابی منابع را زیرساخت سخت‌افزاری و نرم‌افزاری و مدیریت هزینه‌ها می‌دانستند. مشارکت‌کننده (۱۷) به چالش‌های مالی در این خصوص اشاره کرد و گفت: «به نظر من، جدی‌ترین ریسک در استفاده از هوش مصنوعی برای حفاظت از میراث مستند به موضوع دسترسی به GPUها و در صورت نبود آن، اتکا به APIهای خارجی مدل‌های زبانی بزرگ بازمی‌گردد. GPUها هم بسیار پرهزینه‌اند و هم نگهداری و مدیریت آن‌ها نیازمند زیرساخت‌هایی مانند GPU Farm^۲ است که فراهم کردن آن برای بسیاری دشوار است. از سوی دیگر، استفاده از APIهای خارجی نیز با موانعی مانند تحریم‌ها، محدودیت‌های دسترسی و دشواری‌های امنیتی همراه است». مشارکت‌کننده (۱) با ذکر مثال عینی عنوان کرد: «سرور فعلی ما دارای GPUهای قدیمی است و نمی‌توانیم مدل‌های سنگین را به‌طور عملیاتی اجرا کنیم. اجرای این مدل‌ها نیازمند GPU

^۱ injection

^۲ منظور از GPU Farm یا «مزرعه GPU» زیرساخت محاسباتی بزرگ است که از تعداد زیادی کارت گرافیک قدرتمند (GPU) به‌صورت موازی و شبکه‌ای استفاده می‌کند.

قدرتمند، منبع تغذیه مطمئن، سامانه خنک‌کننده و محیط مناسب است. علاوه بر این، به دلیل مسائل امنیت داده و تحریم‌ها، نمی‌توانیم از سرویس‌های ابری استفاده کنیم و باید مدل‌ها را محلی نصب و آموزش دهیم. به همین دلیل، اجرای عملیاتی هوش مصنوعی در حال حاضر بیشتر جنبه آزمایشی دارد». هم‌راستا با این دیدگاه، مشارکت‌کننده (۹) بیان کرد: «پیش‌تر، پردازنده‌های معمولی (CPU) کافی بودند، اما اکنون برای پردازش‌های سریع و موازی به GPU و TPU نیاز داریم. متأسفانه دسترسی به این شتاب‌دهنده‌ها محدود است و کمبود آن‌ها، علاوه بر کندی پردازش، ریسک‌هایی مانند نفوذپذیری و مسائل امنیت اطلاعات را افزایش می‌دهد. محدودیت در انرژی و تجهیزات کنترل دما می‌تواند به آسیب سخت‌افزاری و کاهش پایداری سامانه منجر شود. سرمایه‌گذاری ناکافی روی تراشه‌های کم‌مصرف و مبدل‌های انرژی نیز این ریسک‌ها را تشدید می‌کند. از سوی دیگر، محدودیت ذخیره‌سازی و مدیریت ناکافی بکاپ، احتمال از دست رفتن داده‌های ارزشمند میراث مستند را افزایش می‌دهد».

به باور مشارکت‌کننده (۲ و ۳) برای پروژه‌های ملی، منابع سخت‌افزاری قابل‌توجهی لازم است: به‌عنوان مثال مدل‌هایی مانند چت‌جی‌پی‌تی برای آموزش از حدود ۱۰۰ هزار GPU استفاده کرده‌اند، درحالی‌که مدل‌های کوچک‌تر ممکن است با ۱۰ هزار GPU آموزش داده شوند. هزینه هر کارت تخصصی هوش مصنوعی اکنون بیش از ۲ میلیارد تومان است. بنابراین سازمان‌ها باید براساس تعداد کاربران و کاربرد مدل، نیاز سخت‌افزاری و بودجه خود را تعیین کنند». مشارکت‌کننده (۱۴) افزود: «با توجه به ماهیت مجموعه‌های مرتبط با میراث مستند، امکان تعریف پروژه‌های بزرگ و گسترده مبتنی بر هوش مصنوعی برای سازمان‌ها فراهم نیست. معمولاً این مجموعه‌ها ناگزیرند با بودجه‌های محدود، پروژه‌هایی کوچک‌تر و سبک‌تر طراحی کنند؛ پروژه‌هایی که در عمل اغلب اثربخشی چندانی ندارند و به نتایج قابل‌توجهی منجر نمی‌شوند».

مشارکت‌کننده (۹) نیز با توجه به محدودیت‌های مالی سازمان‌ها تصریح کرد: «بسیاری از نرم‌افزارهای کتابخانه‌ای موجود اساساً آمادگی لازم برای پذیرش پردازش‌های پیشرفته، مانند بازیابی معنایی را ندارند و تغییر در این زیرساخت‌ها نیازمند بازطراحی اساسی است».

مشارکت‌کننده (۳) به سطح کلان کشور در این زمینه اشاره کرد و گفت: «من اطلاع دارم که در ایران زیرساخت‌های گسترده برای هوش مصنوعی تقریباً وجود ندارد. سازمان هوش مصنوعی که قرار بود هزار کارت پردازشی بخرد و خدمات متمرکز ارائه دهد، منحل شد. دانشگاه‌ها هم منابع محدودی دارند به نظر من، اصلی‌ترین چالش در استفاده از هوش مصنوعی، کمبود زیرساخت و توان سخت‌افزاری است و تقریباً به‌طور کامل وابسته به وارداتیم».

مشارکت‌کننده (۷) به تجربه سازمانی خود در این زمینه اشاره کرد و تصریح کرد: «در برخی موارد، از سرویس‌های خارجی استفاده می‌کنیم؛ البته تنها زمانی که ملاحظات محرمانگی داده، حجم اطلاعات و صرفه اقتصادی اجازه دهد و کیفیت مورد انتظار نیز تضمین شود. اما هر جا این شرایط برقرار نباشد، ناچاریم زیرساخت‌های داخلی را خودمان فراهم کنیم. به همین دلیل طی دو سال گذشته سروی مجهز به هشت GPU پیشرفته H200 تهیه کرده‌ایم که در سطح کشور کم‌نظیر است. در مجموع، هزینه‌های سنگین سخت‌افزاری و ضرورت به‌روزرسانی مداوم از مهم‌ترین ریسک‌ها و موانع ما در این حوزه محسوب می‌شوند». مشارکت‌کننده (۳ و ۱۰) به بحث چالش‌های زیرساختی و مالی در کتابخانه ملی اشاره کرد و اظهار داشت: «داده‌های کتابخانه ملی بسیار غنی است، اما پرسش اساسی این است که با چه زیرساختی قرار است این داده‌ها را آموزش دهیم؟ برای یک آموزش کامل حداقل به ۱۰ تا ۲۰ GPU نیاز داریم و برای تنظیم دقیق نیز دست‌کم ۵ GPU قدرتمند لازم است. البته این هزینه مقطعی نیست، بلکه نیازمند سرمایه‌گذاری مستمر است»؛ به اعتقاد مشارکت‌کننده (۳) نیز کتابخانه ملی مجموعه داده‌های بسیار گسترده‌ای از منابع و کتاب‌ها را در اختیار دارد. اما به گفته ایشان، «گام اُسی‌آر آن یکی از چالش‌برانگیزترین مراحل کار است. با توجه به

تنوع گسترده داده‌ها و حجم زیاد آن‌ها، کار بسیار پیچیده است و اجرای آن نیازمند منابع GPU گسترده است. مشارکت‌کننده (۵) در این خصوص افزود: «در سال ۱۳۹۰، برای بهبود دقت اُسی آر در منابع نشریات دیجیتال‌شده کتابخانه ملی، از مازول‌های تخمین عبارت‌های فارسی استفاده کردیم. این اقدام باعث شد دقت نتایج از حدود ۸۲ درصد به ۹۲ درصد افزایش یابد. در این مورد دو دسته چالش جدی وجود داشت که یکی از آن‌ها چالش فنی بود که به محدودیت‌های ابزار و ناکافی بودن کیفیت خروجی مربوط می‌شد».

مشارکت‌کننده (۹) به دیگر زیرساخت‌های ضروری در بهره‌گیری از توان هوش مصنوعی اشاره کرد و گفت: «یکی از چالش‌های اصلی در بازیابی دانش، توانایی ما در شناسایی رابطه بین داده‌ها و موجودیت‌ها در زبان فارسی است. دلیل این محدودیت، نبود اصطلاح‌نامه‌های به‌روز در حوزه‌های تخصصی است و هستی‌شناسی دقیق یا آنتولوژی جامع نداریم. در نتیجه، بازیابی معنایی اطلاعات به شکل مؤثر ممکن نیست. در کتابخانه ملی، اقدامات خوبی در زمینه داده‌های پیوندی صورت گرفته است، اما کافی نیست. اگر بتوانیم بازیابی هوشمند و معنایی اطلاعات را پیاده کنیم، می‌توانیم مشکلات موجود در داده‌ها را کاهش دهیم. با افزودن استنتاج فازی، سامانه‌های پیشنهاددهنده مبتنی بر روابط داده‌ها و یکپارچه‌سازی با سامانه‌های موجود، می‌توان بازیابی معنایی را بهبود بخشید. انتقال از استانداردهای سنتی مانند MARC به الگوهای مفهومی نظیر BIBFRAME و LRM، زیرساخت لازم برای رویکردهای هوشمندانه را فراهم می‌کند و پیاده‌سازی آن‌ها را آسان‌تر می‌سازد».

در این زیرمقاله سه دسته چالش اصلی مطرح شد: (۱) تأمین GPU پیشرفته، برق پایدار و سامانه خنک‌کننده (هزینه بالا و محدودیت تحریمی، (۲) وابستگی به فناوری خارجی (عدم امکان استفاده از سرویس‌های ابری خارجی به دلیل امنیت و تحریم، (۳) هزینه‌های جاری نگهداری و به‌روزرسانی.

ریسک دانش و مهارت منابع انسانی و پایداری نیروی متخصص

یکی از بنیادی‌ترین چالش‌ها در مسیر استقرار موفقیت‌آمیز فناوری‌های نوین مانند هوش مصنوعی، متوجه سرمایه انسانی سازمان است. فناوری هوش مصنوعی، صرفاً یک ابزار نیست، بلکه اکوسیستم پیچیده‌ای است که مدیریت و بهره‌برداری از آن مستلزم وجود مهارت‌های تخصصی و میان‌رشته‌ای است. مشارکت‌کننده (۱۴) «یکی از چالش‌های مهم در زمینه استفاده از هوش مصنوعی را کمبود مهارت و توانمندی‌های لازم در سازمان‌ها دانست. اغلب اوقات کارها به ارگان‌های ثالث سپرده می‌شود و این وابستگی به نهادهای بیرونی می‌تواند محدودیت‌هایی در کنترل و مدیریت فرآیندها ایجاد کند». مشارکت‌کننده (۳) نیز استخدام متخصص را یکی دیگر از ریسک‌های اصلی در این حوزه قلمداد کرد و گفت: «متخصصان ماهر کم هستند و ممکن است شما منابع خود را روی یک استارت‌آپ یا فردی سرمایه‌گذاری کنید و در نهایت نتیجه مطلوب را دریافت نکنید. بسیاری از کسانی که در پروژه‌ها شرکت می‌کنند حتی تحصیلات مرتبط با هوش مصنوعی ندارند. علاوه بر تخصص، مشکل دیگر مربوط به همکاری و تعامل تیمی است؛ بعضی از افراد اشتیاق یا توانایی لازم برای درک فرآیندهای سازمان را ندارند و زبان مشترکی برای ارتباط با کارشناسان پروژه وجود ندارد. این ترکیب کمبود تخصص و ضعف در همکاری باعث می‌شود پروژه‌ها با ریسک بالا مواجه شوند و خروجی نهایی گاهی صرفاً یک «سره‌مبندی» باشد». مشارکت‌کننده (۱۳) ریشه کمبود مهارت‌ها و تخصص‌های جدید را به دانشگاه‌ها نسبت داد و معتقد بود: «در حال حاضر، بسیاری از پروژه‌ها و پایان‌نامه‌های دانشگاهی عمق علمی و نوآوری کافی ندارند».

مشارکت‌کننده (۷) نیز صراحتاً به بحث فرار مغزها اشاره کرد و گفت: «به‌طور کلی در کشور در بخش فناوری اطلاعات با مشکل کمبود نیروی ماهر مواجه‌ایم. بسیاری از متخصصان یا مهاجرت می‌کنند یا ترجیح می‌دهند به‌صورت دورکاری برای پروژه‌های خارجی فعالیت کنند».

مشارکت‌کننده (۱۹) اظهار نگرانی کرد و گفت: «به نظر من مهم‌ترین ریسک این است که در فرآیند استقرار این ابزارها از افراد متخصص استفاده نشود. بسیاری از ما، مثل کتابداران یا پژوهشگران، کاربران ابزارهای آماده هستیم و لزوماً متخصص این حوزه به شمار نمی‌آییم. اما پیاده‌سازی و ایجاد زیرساخت‌های هوش مصنوعی در سازمان‌ها نیازمند مهندسان و متخصصانی است که دانش عمیق و به‌روز داشته باشند و با لبه‌های علمی و فنی این حوزه در سطح جهانی آشنا باشند».

کمیود نیروی ماهر، فرار مغزها، ناپایداری نیروی انسانی (ترک پروژه در میانه کار) و ضعف در همکاری میان‌رشته‌ای از جمله چالش‌های مطرح در این زیرمقوله بود. همچنین نبود آموزش‌های به‌روز در دانشگاه‌ها و وابستگی به افراد یا استارت‌آپ‌های بیرونی به‌عنوان ریسک مدیریتی مطرح شد.

ضعف در برنامه‌ریزی، سیاست‌گذاری، اولویت‌بندی و مدیریت پروژه‌های هوش مصنوعی

در حوزه حفاظت از میراث مستند، پیاده‌سازی موفق هوش مصنوعی تنها مسئله‌ای فنی نیست، بلکه در درجه نخست مسئله‌ای راهبردی و حکمرانی است. این مقوله به ریسک‌هایی می‌پردازد که ریشه در لایه‌های مدیریتی، تصمیم‌سازی و فرهنگ سازمانی دارند. فقدان چهارچوب حکمرانی منسجمی برای فناوری اطلاعات می‌تواند تمامی سرمایه‌گذاری‌ها در حوزه هوش مصنوعی را به ورطه شکست بکشاند (NIST, 2023). مشارکت‌کننده (۶) به‌صراحت به این نکته اشاره کرد و گفت: «یکی از بزرگ‌ترین موانع در کشور ما، مقاومت مدیران و نیروی انسانی در برابر پذیرش نوآوری و نبود حکمرانی منسجم در حوزه هوش مصنوعی است. درحالی‌که بسیاری از کشورهای منطقه مانند قطر، عمان و امارات، اسناد ملی و برنامه‌های بلندمدت روشنی در این زمینه تدوین کرده‌اند درحالی‌که ما در کشور شاید به صد دستگاه مدرن هم دسترسی نداشته باشیم. ما هنوز بر سر تأسیس یک نهاد ملی هوش مصنوعی و تعیین متولی اصلی آن دچار اختلاف هستیم. تحریم‌ها و هزینه‌های سنگین نیز این کمیود را تشدید کرده است. باین‌حال، مشکل اصلی تنها در سخت‌افزار نیست، بلکه در نگاه بسته و مقاومت ساختاری مدیران و کارکنان است».

مشارکت‌کننده (۳) درباره سیاست‌گذاری در حوزه هوش مصنوعی در سطح کلان ابراز نگرانی کرد و گفت: «در ایران، وضعیت هوش مصنوعی هنوز جایگاه مشخصی ندارد. سرمایه‌گذاری خاصی صورت نگرفته و دانش این حوزه تازه و نوظهور است. بنابراین، انتظار خروجی بزرگ از منابع محدود غیرواقعی است. چهارچوب‌های حاکمیت داده تعریف نشده و تصمیم‌گیری اغلب غیرعلمی و سلیقه‌ای است».

مشارکت‌کننده (۱۹) نیز به اهمیت مسئله مدیریت در فناوری اطلاعات در سازمان‌ها تأکید کرد و اظهار داشت: «ما نیاز به حکمرانی فناوری اطلاعات داریم و این حکمرانی متأسفانه در حال حاضر در سازمان و حتی در سطح کشور، ناقص یا با ابهام است. بنابراین، نبود مدیریت فعال، انگیزه ناکافی و ضعف در اعمال استانداردها، باعث می‌شود حتی در دسترس بودن ابزار و فناوری‌های پیشرفته نتواند به‌طور مؤثر مورد استفاده قرار گیرد».

مشارکت‌کننده (۱۱) بحث سیاست‌گذاری را مهم دانست. ایشان هشدار داد: «پیش از هر چیز، سازمان باید به بلوغ لازم برای تغییر سیاست‌ها برسد و زمان مناسب اصلاح سیاست‌ها را درک کند. این تصمیم باید هماهنگ با مسیر کلان و اهداف سازمان باشد، در غیر این صورت، بهره‌برداری مؤثر از فناوری‌های نوین با ریسک مواجه خواهد شد».

تأمین بودجه و مدیریت آن نیز در صحبت‌های برخی مشارکت‌کنندگان مشاهده شد. مشارکت‌کننده (۸) اذعان داشت: «اجرای پروژه‌های هوش مصنوعی نیازمند بودجه هنگفت و نیروی متخصص است و اگر مدیران توانمند بر پروژه نظارت نداشته باشند، ممکن است منابع، صرف امور غیرمرتبط شود». مشارکت‌کننده (۱۸ و ۱۹) به همین موضوع در قالب بی‌ثباتی پست‌های مدیریتی اشاره کردند و گفتند: «تغییرات مکرر در پست‌های مدیریتی سازمان‌های دولتی باعث می‌شود بسیاری از پروژه‌ها نیمه‌کاره رها شوند و منابع ارزشمند سازمانی هدر برود. اگر ما همانند شرکت‌های خصوصی برای هر پروژه مدیر مستقلاً تعیین کنیم که فراتر از تغییرات سیاسی، وظیفه تکمیل پروژه را بر عهده داشته باشد، می‌توانیم از اتلاف منابع جلوگیری کنیم». مشارکت‌کننده

(۵) نیز به تجربه خود در این زمینه اشاره کرد و گفت: «تجربه دیگری که داشتم مربوط به استفاده از یک سامانه پیشنهاددهنده^۱ در بستر دیجیتال در کتابخانه ملی بود. من به عنوان مسئول فنی، استقرار این ابزار را بر روی سایت بر عهده داشتم و حدود شش ماه داده گردآوری شد. اما با تغییر مسئولیت و واگذاری کار به فرد دیگر، این روند متوقف شد».

به باور مشارکت‌کننده (۱۹) «ما شاهد پروژه‌های بسیار بزرگ و پرهزینه‌ای هستیم که تحت عنوان «مگا پروژه» تعریف می‌شوند؛ اما یا به نتیجه نمی‌رسند، یا در صورت اجرا، هزینه‌های نگهداری آن قدر بالاست که تداوم و پایداری‌شان دشوار یا حتی غیرممکن می‌شود». مشارکت‌کننده (۱۸) نیز تأکید و تصریح کرد: «همیشه مشکل نبود امکانات نیست. برای نمونه، چند سال پیش یک سامانه اُسی آر ظاهراً با هزینه قابل توجهی در کتابخانه ملی خریداری شده، اما استفاده مؤثری از آن صورت نگرفته است. این نشان می‌دهد که علاوه بر خرید نرم‌افزار یا سخت‌افزار، نحوه بهره‌برداری و به‌کارگیری نیز اهمیت اساسی دارد».

مشارکت‌کننده (۱۹) به صراحت به ناکارآمدی و بی‌تخصصی مدیران اشاره کرد و معتقد بود: «به نظر من بزرگ‌ترین مخاطره هوش مصنوعی در سازمان ما نه فناوری است و نه الگوریتم‌ها، بلکه مدیران ناکارآمد و بی‌تخصص در حوزه فناوری اطلاعات هستند. بنابراین حتی با ابزارهای پیشرفته، فقدان حاکمیت فناوری و فرهنگ سازمانی مناسب مانع بهره‌برداری مؤثر از هوش مصنوعی می‌شود».

مشارکت‌کننده (۱) بر نبود اولویت‌بندی روشن برای به‌کارگیری هوش مصنوعی در کتابخانه ملی اشاره کرد و گفت: «تصمیم‌گیری و اجرای پروژه‌ها اغلب از سمت فرد یا تیم داخلی است، بدون راهنمایی یا سفارش رسمی از مدیریت. سازمان هنوز این موضوع را به عنوان اولویت رسمی در نظر نگرفته است».

مشارکت‌کننده (۳) بر زمان‌بندی نادرست پروژه‌های هوش مصنوعی اشاره کرد و گفت: «به دلیل نبودن کار و عدم آشنایی کامل با سازمان، تخمین دقیق زمان لازم دشوار است. به همین ترتیب، بودجه‌بندی پروژه هم می‌تواند اشتباه باشد، زیرا حجم واقعی داده‌ها و میزان پردازش مورد نیاز اغلب بیشتر از آن چیزی است که ابتدا پیش‌بینی می‌کنیم».

خبرگان در این زیرمقوله بر فقدان حاکمیت هوش مصنوعی، نبود اولویت‌بندی در سازمان، بی‌ثباتی مدیریتی، تصمیم‌گیری سلیقه‌ای و ضعف فرهنگ سازمانی تأکید داشتند.

انتظارات غیرواقع‌بینانه

در فضای کنونی که تبلیغات گسترده حول هوش مصنوعی، آن را به عنوان راه‌حلی جادویی و همه‌فن‌تصویر می‌کند، این خطر وجود دارد که ذی‌نفعان، انتظاراتی اغراق‌آمیز و فاقد پشتوانه فنی از این فناوری پیدا کنند. چنین انتظاراتی می‌تواند منجر به تعریف اهداف غیرقابل دستیابی، تخصیص نادرست منابع و در نهایت، ناامیدی و عقب‌نشینی کامل از پروژه‌ها شود. مشارکت‌کننده (۳) در این زمینه چنین بیان کرد: «واقعیت این است که هوش مصنوعی آن‌طور که تبلیغ می‌شود، هوشمند کامل نیست. ما هنوز در مراحل ابتدایی آن هستیم و بسیاری از بحث‌ها پیرامون آن بیشتر جنبه رسانه‌ای و غیرعلمی دارد. بنابراین، من احساس نمی‌کنم که این ابزارها بتوانند جایگزین کامل سرویس‌ها و تجربه انسانی شوند. در عمل، این فناوری‌ها بیشتر به عنوان ابزار کمکی کنار انسان به کار می‌روند و تجربه و قضاوت انسانی قابل جایگزینی نیست. پیش‌بینی من این است که در نهایت، ما شاهد همزیستی انسان و هوش مصنوعی خواهیم بود».

مشارکت‌کننده (۱۰) به انتظارات غیرواقع‌بینانه در سطح مدیران اشاره کرد و گفت: «بزرگ‌ترین چالشی که می‌بینم این است که مدیران درک درستی از هوش مصنوعی ندارند و توقع دارند که یک مدل زبانی بزرگ یا ربات بتواند همه کارهایی که انسان انجام می‌دهد را به‌طور کامل جایگزین کند. این انتظارات غیرواقعی باعث می‌شود که پروژه‌ها به‌سختی پیش بروند یا حتی اجرا نشوند». مشارکت‌کننده (۳) نیز به نکته مشابهی اشاره کرد و اذعان داشت: «نباید انتظار داشته باشیم هر پروژه‌ای در ایران بتواند

^۱recommender

به سطحی مانند چت جی پی تی برسد. چت جی پی تی با هزینه‌ای در حدود یک تریلیون دلار ساخته شده است. تا زمانی که ایران نتواند حتی کسری از چنین سرمایه‌ای را تأمین کند، انتظار ایجاد محصولی مشابه چندان واقع‌بینانه نیست». مشارکت‌کننده (۴) افزایش دانش عمومی مدیران ارشد نسبت به هوش مصنوعی و کاربردهای آن را جهت اتخاذ تصمیم‌گیری‌های واقع‌بینانه و راهبردی ضروری دانست و بیان کرد: «دانش و نگاه مدیریتی نسبت به هوش مصنوعی محدود است. تجربه برخی خطاهای فناوری در پردازش متون فارسی، باعث می‌شود برخی مدیران نسبت به کاربرد هوش مصنوعی شک کنند و از آن صرف‌نظر کنند. بنابراین، افزایش دانش عمومی مدیران ارشد نسبت به هوش مصنوعی و کاربردهای آن ضروری است بدون آنکه لازم باشد وارد جزئیات فنی شوند».

تبلیغات رسانه‌ای و بزرگنمایی توانمندی‌های هوش مصنوعی باعث می‌شود مدیران انتظار جایگزینی کامل انسان و حل یک‌باره مسائل را داشته باشند. این انتظارات منجر به تعریف اهداف دست‌نیافتنی، تخصیص نادرست منابع و در نهایت ناامیدی می‌شود.

دسترسی به ابزارهای هوش مصنوعی

ناتوانی سازمان در تأمین، پیکربندی و استفاده بهینه از ابزارها و پلتفرم‌های هوش مصنوعی، چه به دلیل محدودیت‌های فنی، مالی یا تحریمی، یکی از معضلات مهم کشورهای در حال توسعه است. مشارکت‌کننده (۲) تصریح کرد: «بسیاری از ابزارها برای ما قابل دسترس نیستند، یا در صورت استفاده، به محض شناسایی IP ایران، دسترسی مسدود می‌شود». مشارکت‌کننده (۹) نیز با تمرکز بر متون تخصصی به عنوان یک چالش چنین گفت: «بیشتر ابزارهای NLP فعلی بر متون عمومی متمرکز هستند و برای پردازش متون تخصصی نیاز به مدل‌ها یا ابزارهای خاص داریم. مدل‌های زبانی بزرگ توانایی پردازش گسترده دارند، اما چون اغلب بر داده‌های عمومی آموزش دیده‌اند، برای متون تخصصی باید تنظیم یا توسعه داده شوند».

ریسک‌های فرهنگی و اجتماعی

چالش‌های فرهنگی و سازمانی

پیشرفت فناوری تنها به معنای حل مسائل فنی نیست، بلکه مستلزم تغییر در بستر فرهنگی و سازمانی است. در مسیر استقرار هوش مصنوعی، گاه بزرگ‌ترین موانع، نه در سخت‌افزارها و الگوریتم‌ها، که در سنگ‌اندازی‌های نامحسوس فرهنگ سازمانی و ساختارهای مدیریتی نهفته است. مشارکت‌کننده (۱۴) تصریح کرد: «تجربه نشان داده که در ادوار مختلف با رویکردهای متفاوت، مقاومت در برابر نوآوری مشاهده شده و به همین دلیل ایجاد تغییرات سریع و پیاده‌سازی فناوری‌های جدید بسیار دشوار است». به باور مشارکت‌کننده (۶) نیز کارکنان و افراد مرتبط با این حوزه اغلب در برابر تغییرات مقاومت می‌کنند. ایشان روایت کرد: «اگر در پذیرش و کنترل این ابزار تعلل کنیم، ممکن است از دست ما خارج شود». مشارکت‌کننده (۱۹) نیز در این خصوص معتقد است: «بعضی از پرسنل مشتاق تغییر و نوآوری هستند اما گروهی دیگر تمایلی به تحول ندارند و در مقابل تغییر مقاومت می‌کنند. انتظار می‌رود افراد فنی با انگیزه بالا و آمادگی برای یادگیری باشند، اما در عمل بسیاری تنها به نگهداری سامانه‌ها بسنده می‌کنند و هیچ پشتیبانی فنی مؤثری وجود ندارد».

مشارکت‌کنندگان (۱۳ و ۱۸) نیز با توجه به تجربه خود گفتند: «گاهی حتی تغییرات ساده در سامانه‌های نرم‌افزاری می‌تواند برای کاربران به یک مانع بزرگ تبدیل شود».

مشارکت‌کننده (۶) بر ذهنیت‌های مقاوم و سستی در آرشیوها و کتابخانه‌ها اشاره کرد و تصریح کرد: «چالش اصلی نه کمبود بودجه یا فناوری، بلکه غلبه بر ذهنیت‌های مقاوم و سستی است که مانع بهره‌برداری از ظرفیت‌های واقعی هوش مصنوعی در آرشیوها و کتابخانه‌ها می‌شود. این قابلیت، نگرانی‌هایی را در میان کارکنان ایجاد کرده است؛ بسیاری از آن‌ها می‌ترسند شغل خود را از دست بدهند».

مشارکت‌کنندگان (۲ و ۱۷) نیز دغدغه امنیت شغلی را مهم برشمردند و معتقدند: «مقاومت کارکنان به دلیل ورود هوش مصنوعی طبیعی است. این ابزارها ممکن است جایگزین برخی فعالیت‌های انسانی شوند و افراد ناخودآگاه احساس تهدید کنند». مشارکت‌کننده (۷) با توجه به تجربه خود در این زمینه چنین گفت: «در مواردی که سامانه‌های هوشمند جایگزین فعالیت‌های دستی شدند - مثل اعراب‌گذاری متون - طبیعی بود که نگرانی‌هایی در میان نیروهای انسانی ایجاد شود مبنی بر اینکه ماشین‌ها در حال تصاحب وظایف آن‌ها هستند. این موضوع نه تنها در سطح کارکنان بلکه در سطوح مدیریتی نیز مشهود بود». مشارکت‌کننده (۳) به اهمیت فرهنگ و ساختار سازمانی در این زمینه اشاره کرد و گفت: «گاهی حتی اگر ریاست سازمان موافق باشد، فرهنگ و ساختار سازمانی ممکن است با روند پروژه همسو نباشد. بنابراین، برای موفقیت، لازم است که کارشناسان، معاونان و دیگر اعضای سازمان همسو و همکار باشند و مسیر پروژه با واقعیت‌های سازمانی هماهنگ شود». مشارکت‌کننده (۴) به محافظه‌کاری در مدیریت داده‌ها و فقدان دیدگاه باز و کاربردی از جمله اشتراک برون‌داد و داده‌های پژوهشی نیز اشاره کرد و گفت: «بسیاری از داده‌های پژوهشی حساسیت بالایی دارند. درحالی‌که در کشورهای دیگر انتشار آزاد داده‌ها و داده‌های باز سال‌هاست نهادینه شده، ما هنوز در مرحله تصمیم‌گیری درباره انتشار مقالات و نسخه‌های داده‌هایمان هستیم».

مقاومت در برابر تغییر (به دلیل ترس از دست دادن شغل یا دشواری یادگیری)، ضعف همکاری بین‌بخشی، محافظه‌کاری در اشتراک داده و عدم سازگاری فرهنگ سازمانی با نوآوری از موانع اصلی در این زیرمقوله عنوان شد.

چالش اصالت و اعتبار متون تولیدشده توسط هوش مصنوعی

یکی از بنیادی‌ترین چالش‌های پیش روی سازمان‌های حافظ میراث مستند، مانند سازمان اسناد و کتابخانه ملی، در مواجهه با فناوری هوش مصنوعی، مسئله اصالت و اعتبار است. مشارکت‌کننده (۲) یکی از چالش‌های اصلی در این زمینه را چنین بیان کرد: «مشخص نیست یک متن تا چه اندازه توسط انسان تولید شده و تا چه حد توسط هوش مصنوعی ایجاد شده است. اما به نظر من، اصالت یک متن زمانی که توسط انسان تولید شده باشد، همچنان بالاتر از متنی است که توسط هوش مصنوعی تولید شده است». اشتراک نظر قابل توجهی در تحلیل صحبت‌ها در این زمینه به چشم می‌خورد. مشارکت‌کنندگان (۸ و ۱۹) بر تجربه خود در استفاده از هوش مصنوعی اشاره کرده و گفتند: «من تجربه زیادی در استفاده از هوش مصنوعی، به‌ویژه چت‌جی‌پی‌تی، در پژوهش و بازیابی اطلاعات دارم، اما گاهی خروجی‌ها کاملاً اشتباه یا ناقص هستند. به عنوان مثال لینک‌ها یا منابعی که ارائه می‌دهد، درست نیستند. حتی گاهی عنوان مقاله‌ای که ذکر می‌کند با عنوان واقعی آن متفاوت است». مشارکت‌کننده (۱۶) بر سطحی و کم‌عمق بودن دانش تولید شده اشاره کرد و گفت: «مدل‌های زبانی بزرگ، اطلاعات را با سرعت بسیار بالایی ارائه می‌دهند، اما عمق دانشی که منتقل می‌کنند کمتر است و مسیر یادگیری آن‌ها تا حدی تصادفی است. تفاوت اصلی بین استفاده از کتاب و تعامل با مدل زبانی بزرگ در همین موضوع است: کتاب‌ها یادگیری عمیق و مسیرمند ارائه می‌دهند، درحالی‌که مدل‌های زبانی بزرگ سریع و پویا هستند اما کم‌عمق».

مهم‌ترین نگرانی در این زیرمقوله، ابهام در مرز انسان و ماشین در تولید محتواست. خروجی‌های هوش مصنوعی ممکن است حاوی خطاهای فاحش، منابع جعلی و دانش سطحی باشد. سازمان‌های حافظ میراث به دلیل رسالت‌شان در حفظ اصالت، باید در استفاده از محتوای تولیدشده توسط هوش مصنوعی احتیاط بیشتری کنند.

بحث و نتیجه‌گیری

پژوهش حاضر با هدف شناسایی و تحلیل نظام‌مند چالش‌ها و ریسک‌های کاربردی هوش مصنوعی در حفاظت از میراث مستند در کتابخانه‌ها و آرشیوها، با تأکید بر سازمان اسناد و کتابخانه ملی ایران به عنوان کتابخانه مادر کشور، انجام شد. این پژوهش از

طریق مصاحبه‌های نیمه‌ساختاریافته با ۲۰ نفر از خبرگان حوزه‌های علم اطلاعات، آرشیو، فناوری اطلاعات و هوش مصنوعی، به استخراج ۱۶۷ کد نهایی در قالب چهار مقوله اصلی ریسک‌های فناوری، حقوقی و امنیتی، مدیریتی و حکمرانی، و فرهنگی - اجتماعی دست یافت. یافته‌های این پژوهش نشان داد که ریسک‌های مدیریتی و حکمرانی (با بیشترین فراوانی) و ریسک‌های فناوری (با فراوانی قابل توجه) نقش تعیین‌کننده‌تری در موفقیت یا شکست پروژه‌های هوش مصنوعی در نهادهای حافظ میراث مستند ایران دارند. در ادامه، یافته‌های پژوهش در پرتو پیشینه‌های نظری و تجربی تحلیل و تفسیر می‌شود.

ریسک‌های فناوری: یافته‌های این پژوهش نشان داد که چالش‌های مرتبط با کیفیت و سازگاری داده‌ها، محدودیت‌های زبانی، سوگیری الگوریتمی، و هم‌گویی مدل‌ها، قابلیت اعتماد فنی مدل‌ها و وابستگی به سامانه و ناکارآمدی در انتقال دانش تخصصی، از مهم‌ترین ریسک‌های فناوری در حفاظت از میراث مستند محسوب می‌شوند. این یافته‌ها تأکید می‌کنند که مدیریت داده‌ها، پالایش، برچسب‌گذاری و پایش مدل‌ها از اجزای کلیدی تضمین عملکرد صحیح سامانه‌های هوش مصنوعی هستند. هم‌راستا با این یافته، امز و لویس (۲۰۲۰) نیز در پژوهش خود با چالش‌هایی در مورد مدیریت و ذخیره‌سازی داده‌ها مواجه بودند و تأکید کردند کیفیت پایین داده می‌تواند منجر به خروجی‌های نادرست شود و یکپارچگی مجموعه‌های میراث فرهنگی را تحت‌تأثیر قرار دهد. به اعتقاد لیو^۱ (۲۰۲۴) نیز پایداری و دقت الگوریتم‌های هوش مصنوعی نیازمند بهبود است، به‌ویژه در مواجهه با اسناد پیچیده و اطلاعات چندزبانه که چالش‌هایی را ایجاد می‌کند. اوکوروما (۲۰۲۴)، تیل (۲۰۲۴)، تیریبلی^۲ و همکاران (۲۰۲۴) و امز و لویس (۲۰۲۰) نیز بیان کردند الگوریتم‌های هوش مصنوعی اگر به‌دقت نظارت و تنظیم نشوند، می‌توانند سوگیری را تداوم بخشند. کتابخانه‌ها باید اطمینان حاصل کنند که الگوریتم‌هایی که استفاده می‌کنند منصفانه و شفاف هستند و دسترسی برابر به منابع و خدمات را برای تمامی کاربران، صرف‌نظر از پیشینه آن‌ها، فراهم می‌کنند. تیل (۲۰۲۴) نیز تأکید قابل‌توجهی بر لزوم بررسی دقیق داده‌های تولیدشده توسط هوش مصنوعی قبل از ارائه دسترسی به مراجعین دارد. این امر برای اطمینان از صحت و قابلیت اطمینان اطلاعات ارائه شده بسیار مهم است، زیرا خطاها در خروجی‌های هوش مصنوعی می‌تواند کاربران را گمراه کند یا تمامیت کتابخانه و آرشیو را به خطر بیندازد.

باین‌حال، آنچه یافته‌های این پژوهش را متمایز می‌سازد، تأکید بر چالش‌های خاص زبان فارسی است که به‌دلیل کمبود داده‌های آموزشی استاندارد، فقدان ابزارهای پردازش زبان طبیعی کارآمد و نبود زیرساخت‌های متن‌باز، به مانع جدی در استقرار هوش مصنوعی در نهادهای ایرانی تبدیل شده است. این مسئله را می‌توان به‌عنوان شکاف بومی در نظر گرفت که در پژوهش‌های بین‌المللی کمتر به آن پرداخته شده است. همچنین پدیده «وهم‌گویی» در مدل‌های زبانی، به‌ویژه در مواجهه با داده‌های کم‌بسامد تاریخی و چندزبانه، یکی از چالش‌های جدی است که می‌تواند به تحریف یا جعل اطلاعات تاریخی بینجامد و اعتبار نهادهای حافظ میراث را با خطر مواجه سازد.

ریسک‌های حقوقی و امنیتی: مهم‌ترین ریسک‌های حقوقی شناسایی شده در این پژوهش شامل نقض حق نشر و مالکیت فکری، محرمانگی و نقض حریم خصوصی و فقدان شفافیت در استفاده از داده‌های دارای محدودیت دسترسی و ریسک‌های امنیتی است. این یافته‌ها با پژوهش‌های اوکوروما (۲۰۲۴) و تیل (۲۰۲۴) که بر پیامدهای اخلاقی و حقوقی پیاده‌سازی هوش مصنوعی تأکید داشته‌اند، همسو است. این یافته‌ها حاکی از آن است که حفاظت از داده‌های شخصی، رعایت قوانین حق نشر و کاهش آسیب‌پذیری سایبری، باید به‌طور هم‌زمان و سازمان‌یافته مورد توجه قرار گیرد. به اعتقاد تیریبلی و همکاران (۲۰۲۴) نیز هوش مصنوعی مولد با بازتعریف مفهوم مالکیت معنوی و اصالت، چالش‌های حقوقی، اخلاقی و اقتصادی جدیدی را برای هنر و میراث مستند به همراه آورده و نیازمند تدوین قوانین و استانداردهای نوین در این حوزه است.

^۱Liu^۲Tiribelli

پژوهش حاضر این مسئله را در بافت خاص سازمان‌های دولتی و نهادهای حافظ میراث ملی بررسی کرده و نشان داده است که نبود سازمان مدیریت جمعی برای حق مؤلف در ایران و همچنین نبود قوانین شفاف درباره استفاده از داده‌های دارای حق نشر در آموزش مدل‌های هوش مصنوعی، این ریسک‌ها را تشدید می‌کند. در سطح امنیتی، وابستگی به فناوری‌های خارجی و کمبود زیرساخت‌های ابری داخلی، سازمان‌های ایرانی را در معرض ریسک‌هایی چون حملات سایبری، دسترسی‌های غیرمجاز و افشای داده‌های حساس قرار داده است. این یافته‌ها ضرورت تدوین قوانین بومی حقوقی و امنیتی برای استفاده از هوش مصنوعی در حوزه میراث مستند را آشکار می‌سازد.

ریسک‌های مدیریتی و حکمرانی: براساس یافته‌های پژوهش، چالش‌های مالی و زیرساختی با بالاترین فراوانی و ضعف در برنامه‌ریزی، سیاست‌گذاری، اولویت‌بندی و مدیریت پروژه‌های هوش مصنوعی، مهم‌ترین ریسک‌های مدیریتی شناسایی شدند. تأمین سخت‌افزارهای پیشرفته (GPU)، هزینه‌های نگهداری و به‌روزرسانی، وابستگی به فناوری خارجی و ... از مهم‌ترین موانع پیاده‌سازی هوش مصنوعی در سازمان‌های ایران هستند. ریسک دانش و مهارت منابع انسانی و پایداری نیروی متخصص نیز حائز است و نشان‌دهنده اهمیت مدیریت، حکمرانی و سیاست‌گذاری راهبردی در موفقیت هوش مصنوعی است. مقربی منظری (۱۳۹۸) و محمدصالحی و همکاران (۱۴۰۱) نیز در پژوهش خود ریسک‌های نیروی انسانی، ریسک‌های فنی، ریسک‌های امنیت اطلاعات و ریسک‌های زیرساختی را در رتبه‌های بالاتر گزارش دادند. علی و همکاران (۲۰۲۴) نیز، استفاده از این فناوری را نیازمند سرمایه‌گذاری در زمینه تأمین مالی، زمان و نیروی انسانی دانست و از کمبود آن ابراز نگرانی نمود. به اعتقاد پاسیکوفسکا-اشناس و لیم^۱ (۲۰۲۳، ص. ۹) نیز به‌کارگیری هوش مصنوعی در حوزه میراث مستند و میراث فرهنگی نیازمند سرمایه‌گذاری در زمینه‌های متعددی است که زیرساخت‌ها، تجهیزات و منابع انسانی ماهر از مهمترین آن‌ها محسوب می‌شوند. ایشان ریسک سرمایه‌گذاری و در نتیجه دشواری در دسترسی به منابع مالی را یکی از چالش‌ها عنوان کرد و بیان کرد این خطر وجود دارد که سازمان‌های فرهنگی نیاز به سرمایه‌گذاری در هوش مصنوعی را کمتر ضروری تلقی کنند. امز و لویس (۲۰۲۰) نیز یکی از خطرات اصلی در این زمینه را کمبود تخصص کارکنان در برنامه‌های هوش مصنوعی دانستند. پژوهش ربیعی و همکاران (۱۴۰۱) نیز نشان داد بودجه محدود می‌تواند توانایی نهادها را در دیجیتالی کردن و نگهداری از اشیاء میراث فرهنگی به‌طور مؤثر محدود کند.

پژوهش حاضر ابعاد جدیدی را نیز آشکار نمود. ازجمله وابستگی به فناوری خارجی و ریسک تحریم که در پژوهش‌های بین‌المللی کمتر مورد توجه قرار گرفته است. این مسئله به چالش ژئوپلیتیکی تبدیل شده که سازمان‌های ایرانی را از به‌روزترین مدل‌ها و سخت‌افزارها محروم می‌کند و آن‌ها را به استفاده از راهکارهای جایگزین و اغلب ناکارآمد مجبور می‌سازد. همچنین ریسک‌های مدیریتی نظیر ضعف در برنامه‌ریزی و مدیریت پروژه‌ها، بی‌ثباتی پست‌های مدیریتی و کمبود فرهنگ سازمانی مناسب، از دیگر چالش‌های جدی است. بااین‌حال، پژوهش حاضر نشان می‌دهد که ضعف حکمرانی فناوری اطلاعات در سطح ملی و سازمانی، یکی از مهم‌ترین موانع ساختاری است که ریشه بسیاری از مشکلات دیگر ازجمله نبود بودجه پایدار، عدم توانایی در جذب و حفظ نیروی متخصص و نبود استانداردهای شفاف را تشکیل می‌دهد.

ریسک‌های فرهنگی - اجتماعی: مقاومت سازمانی در برابر تغییر، ترس از دست دادن شغل، نبود فرهنگ سازمانی مناسب، از مهم‌ترین ریسک‌های فرهنگی - اجتماعی شناسایی شده در این پژوهش هستند. این یافته‌ها با پژوهش اوکوروما (۲۰۲۴) که بر پیامدهای فرهنگی و اجتماعی پیاده‌سازی هوش مصنوعی تأکید داشت، همسو است. پژوهش حاضر نشان می‌دهد که در سازمان‌های دولتی ایران، مقاومت در برابر تغییر ریشه‌ای عمیق‌تر دارد و به عواملی چون ساختارهای بوروکراتیک، نبود انگیزه‌های کافی برای نوآوری و نبود شفافیت در فرایندهای تصمیم‌گیری بازمی‌گردد. تحلیل این داده‌ها نشان می‌دهد که حتی

^۱Pasikowska-Schnass & Lim

اگر زیرساخت‌های فناوری و داده‌ای و سیاست‌های مدیریتی به‌خوبی طراحی شوند، فقدان آمادگی فرهنگی و سازمانی می‌تواند مانع بهره‌برداری مؤثر از سامانه‌های هوش مصنوعی شود. همچنین چالش اصالت و اعتبار متون تولیدشده توسط هوش مصنوعی، به‌ویژه در حوزه میراث مستند، نگرانی‌های جدی را در میان خبرگان ایجاد کرده است. تیریلی و همکاران (۲۰۲۴) نیز معتقدند بازتولید سه‌بعدی تعاملی آثار تاریخی نمی‌تواند اصالت اثر را به‌درستی نمایش دهد و حتی ممکن است به‌نوعی «غیرانسانی‌سازی» میراث فرهنگی منجر شود. بنابراین، طراحی سازوکارهای نظارت انسانی و اعتبارسنجی دقیق برای محتوای تولیدشده توسط هوش مصنوعی، ضرورت انکارناپذیری است که در پژوهش‌های پیشین کمتر به آن توجه شده است.

پژوهش حاضر نشان داد که پیاده‌سازی موفق هوش مصنوعی در سازمان‌های حافظ میراث مستند، ازجمله سازمان اسناد و کتابخانه ملی ایران، فراتر از انتخاب فناوری مناسب، مستلزم ایجاد اکوسیستم چندلایه‌ای است که ابعاد فنی، حقوقی، مدیریتی و فرهنگی را به‌طور هم‌زمان پوشش دهد.

یافته‌های پژوهش حاضر، طیف گسترده‌ای از کنشگران را در برمی‌گیرد. در صدر این گروه‌ها، سازمان اسناد و کتابخانه ملی ایران به‌عنوان متولی اصلی حفاظت از میراث مستند کشور و الگوی سایر کتابخانه‌ها و مراکز آرشیوی قرار دارد که می‌تواند با بهره‌گیری از نتایج این پژوهش، گام‌های مؤثری در مسیر پیاده‌سازی هوش مصنوعی بردارد. علاوه بر این، سیاست‌گذاران و برنامه‌ریزان حوزه فناوری اطلاعات در سطح ملی می‌توانند با تکیه بر این یافته‌ها، چهارچوب‌های کلان و قوانین مورد نیاز برای توسعه امن و پایدار هوش مصنوعی را تدوین کنند. همچنین مدیران و کارشناسان کتابخانه‌ها و مراکز آرشیوی سراسر کشور که به دنبال به‌کارگیری هوش مصنوعی در سازمان خود هستند، پژوهشگران حوزه علم اطلاعات، آرشیو و هوش مصنوعی که به دنبال توسعه دانش در این زمینه هستند و نیز شرکت‌های دانش‌بنیان و توسعه‌دهندگان فناوری در حوزه هوش مصنوعی و مدیریت محتوا، همگی ازجمله ذی‌نفعان اصلی این پژوهش به‌شمار می‌روند که می‌توانند از دستاوردهای آن برای بهبود فرآیندها و ارتقای کیفیت خدمات خود بهره ببرند.

پیشنهاد می‌شود سازمان اسناد و کتابخانه ملی ایران با تدوین سند راهبردی هوش مصنوعی مبتنی بر مدیریت ریسک و سرمایه‌گذاری هدفمند در زیرساخت‌های داده و سخت‌افزار (شامل تأمین GPU، توسعه مدل‌های زبانی بومی و ایجاد مجموعه داده‌های استاندارد فارسی)، گام نخست را در مسیر تحول دیجیتال بردارد. سیاست‌گذاران ملی نیز باید با وضع قوانین شفاف برای حق نشر و حریم خصوصی، حمایت از ایجاد GPU Farm ملی و تخصیص بودجه پایدار، بسترهای قانونی و مالی لازم را فراهم آورند. در سطح سازمانی، توانمندسازی نیروی انسانی از طریق آموزش‌های نظام‌مند، اجرای پروژه‌های آزمایشی با نظارت انسانی و سازوکارهای اعتبارسنجی مستمر و ترویج فرهنگ سازمانی مبتنی بر نوآوری و پذیرش تغییر، از ضرورت‌های انکارناپذیر است. این اقدامات، در کنار همکاری با دانشگاه‌ها و شرکت‌های دانش‌بنیان، می‌تواند سازمان را در مسیر بهره‌گیری پایدار، اخلاق‌محور و مأموریت‌گرا از هوش مصنوعی در حفاظت از میراث مستند ملی یاری رساند.

قدردانی

بدین‌وسیله از تمامی خبرگانی که با مشارکت و ارائه دیدگاه‌های تخصصی خود در انجام این پژوهش همکاری کردند، صمیمانه قدردانی می‌شود.

تضاد منافع

نویسنده اعلام می‌کند که در ارتباط با انتشار این مقاله هیچ‌گونه تعارض منافع وجود ندارد.

منابع

- ربیعی، محبوبه، و رضایی شریف‌آبادی، سعید (۱۴۰۱). دیجیتالی شدن اشیاء فرهنگی در بافت میراث فرهنگی (کتابخانه‌ها، موزه‌ها و آرشیوها). *مطالعات دانش‌پژوهی*، ۱ (۲)، ۴۰-۲۰. Doi:10.22034/jkrs.2022.50670.1010
- صمیعی، میترا، طاهری، مهدی، و فرزادی، سمیه (۱۴۰۱). شناسایی و رتبه‌بندی ریسک‌های زنجیره تأمین در کتابخانه‌های دیجیتالی دانشگاه‌های دولتی شهر تهران براساس استاندارد ایزو ۳۱۰۰۰. *پژوهشنامه پردازش و مدیریت اطلاعات*، ۳۷ (۳)، ۷۸۰-۷۴۹. Doi: 10.35050/JIPM010.2022.687
- کریمی، شهبلا (۱۳۹۱). بررسی معیارهای ارزیابی میراث مستند در برنامه حافظه جهانی یونسکو جهت شناسایی مؤلفه‌های موردنظر برای ثبت آثار مستند ایرانی: ارائه یک دستورالعمل پیشنهادی. پایان‌نامه کارشناسی ارشد. دانشگاه الزهرا (س).
- محمدصالحی، راحله، باب‌الحوائجی، فهیمه، صمیعی، میترا، و حریری، نجلا (۱۴۰۱). ارائه الگوی پیشنهادی مدیریت ریسک تولید محتوا در کتابخانه‌های دیجیتالی دولتی شهر تهران. *علوم و فنون مدیریت اطلاعات*، انتشار آنلاین: ۱۵ اسفند ۱۴۰۱. DOI: 10.22091/stim.2022.8091.1775
- مظفری، عظیمة، حیاتی، زهیر، و مظفری، افسانه (۱۴۰۰). به کارگیری مدیریت ریسک و اولویت‌بندی ریسک‌ها در راستای توسعه عملکرد کتابداران کتابخانه‌های عمومی استان فارس. *مطالعات کتابداری و علم اطلاعات (علوم تربیتی و روانشناسی)*، ۱۳ (۱)، ۳۲-۲۱. <https://doi.org/10.22055/slis.2019.18386.1247>
- مقربی منظری، هدی (۱۳۹۸). ارزیابی ریسک‌های منابع انسانی کتابخانه‌های دیجیتالی دانشگاه‌های دولتی شهر تهران. پایان‌نامه کارشناسی ارشد. دانشگاه علامه طباطبائی.

References

- Ames, S., & Lewis, S. (2020). Disrupting the library: Digital scholarship and Big Data at the National Library of Scotland. *Big Data & Society*, 7(2). <https://doi.org/10.1177/2053951720970576>
- Ali, M. Y., Naeem, S. B., Bhatti, R., & Richardson, J. (2024). Artificial intelligence application in university libraries of Pakistan: SWOT analysis and implications. *Global Knowledge, memory and communication*, 73(1/2), 219-234. DOI: [10.1108/GKMC-12-2021-0203](https://doi.org/10.1108/GKMC-12-2021-0203)
- Badawy, W. (2023). Data-driven framework for evaluating digitization and artificial intelligence risk: a comprehensive analysis. *AI and Ethics*, 5(1), 453-478. <https://doi.org/10.1007/s43681-023-00376-4>
- Graneheim, U. H., & Lundman, B. (2004). Qualitative content analysis in nursing research: concepts, procedures and measures to achieve trustworthiness. *Nurse Education Today*, 24(2), 105-112. doi: 10.1016/j.nedt.2003.10.001.
- ISO 31000. (2018). *Risk management — Guidelines*. Retrieved April 19, 2024 from: <https://www.iso.org/obp/ui/#iso:std:iso:31000:ed-2:v1:en>
- ISO/IEC 23894 (2023). *Information technology — Artificial intelligence — Guidance on risk management*. <https://cdn.standards.iteh.ai/samples/77304/cb803ee4e9624430a5db177459158b24/ISO-IEC-23894-2023.pdf>
- Karimi, Sh. (2012). *Investigating the Selection Criteria Document Heritage in UNESCO World Memory Program for identifying considered components to Register Iranian Documentary: a Recommended Manual*. Master Degree, Alzahra University. [In Persian]
- Korka, E., Emmanouloudis, D., Kravari, K., Kokkinos, N., & Dimitriadi, K. (2022). The protection of natural and cultural heritage monuments, museums and archives from risks: Bridging artificial intelligence, risk assessment and stakeholders. In *International Conference on Transdisciplinary Multispectral Modeling and Cooperation for the Preservation of Cultural Heritage* (pp. 17-28). Cham: Springer International Publishing.
- Liu, Z. (2024). The Prospects and Challenges of Artificial Intelligence Technology in Archival Management. *Academic Journal of Management and Social Sciences*, 6(3), 64-66. <https://doi.org/10.54097/wxdsk683>
- Mannheimer, S., Bond, N., Young, S. W., Kettler, H. S., Marcus, A., Slipher, S. K., Clark, J.A., Shorish, Y., Rossmann, D., & Sheehy, B. (2024a). Responsible AI practice in libraries and archives: a review of the literature. *Information Technology and Libraries*, 43(3). DOI: 10.5860/ital.v43i3.17245
- Mannheimer, S., Rossmann, D., Clark, J., Shorish, Y., Bond, N., Scates Kettler, H., Sheehy, B., & Young, S. W. (2024b). Introduction to the Special Issue: Responsible AI in Libraries and Archives. *Journal of eScience Librarianship*, 13(1), DOI: 10.7191/jeslib.860.
- Mogharabi Manzari, H. (2019). *Assessment of human recourse risks in digital libraries of public universities of Tehran city*. Master Degree, Allameh Tabatabai. [In Persian]

- Mohammadsalehi, R., Babalhavaeji, F., Samiei, M., & Hariri, N. (2023). Presenting a proposed model for content production risk management in government digital libraries in Tehran. *Sciences and Techniques of Information Management*. DOI: 10.22091/stim.2022.8091.1775 [In Persian]
- Mozaffari, A., Hayati, Z., & Mozaffari, A. (2021). Applying risk management and prioritizing risks to enhance the performance of public librarians in Fars Province. *Library and Information Science Studies (Educational Sciences and Psychology)*, 13(1), 21-32. <https://doi.org/10.22055/slis.2019.18386.1247>. [In Persian]
- NIST (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*. https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf?isid=enterprisehub_us&ikw=enterprisehub_us_lead%2Fhow-to-responsibly-use-ai-powered-hr-tools_textlink_https%3A%2F%2Fnlpubs.nist.gov%2Fnistpubs%2Fai%2FNIST.AI.100-1.pdf
- Okoroma, F. N. (2024). Artificial intelligence and libraries: Import, risks and prospects. *Journal of ICT Development, Applications and Research*, 6(2), 30–40. DOI: 10.47524/jictdar.v6i2.31.
- Pansoni, S., Tiribelli, S., Paolanti, M., Di Stefano, F., Frontoni, E., Malinverni, E. S., & Giovanola, B. (2023). Artificial intelligence and cultural heritage: Design and assessment of an ethical framework. *International Archives of the Photogrammetry, Remote Sensing and Spatial Information Sciences*, 48, 1149-1155. DOI: 10.5194/isprs-archives-XLVIII-M-2-2023-1149-2023.
- Pasikowska-Schnass, M., & Lim, Y. S. (2023). Artificial intelligence in the context of cultural heritage and museums. *European Parliament briefing, members' research service PE747, 120*, 1-11. [https://www.europarl.europa.eu/RegData/etudes/BRIE/2023/747120/EPRS_BRI\(2023\)747120_EN.pdf](https://www.europarl.europa.eu/RegData/etudes/BRIE/2023/747120/EPRS_BRI(2023)747120_EN.pdf)
- Rabiei, M. and Rezaei sharifabadi, S. (2023). Digitization of Cultural Objects in the Context of Cultural Heritage (Libraries, Archives and Museums). *Journal of Knowledge-Research Studies*, 1(2), 20-40. doi: 10.22034/jkrs.2022.50670.1010 [In Persian]
- Samiei, M., Taheri, M., & Farzadi, S. (2022). Identification and ranking of supply chain risks in digital libraries of public universities in Tehran based on ISO 31000. *Journal of Information Processing and Management*, 37(3), 749–780. Doi: 10.35050/JIPM010.2022.687 [In Persian]
- Teel, Z. (2024). Artificial intelligence's role in digitally preserving historic archives. *Preservation, Digital Technology & Culture*, 53(1), 29-33. DOI: 10.1515/pdtc-2023-0050.
- Tiribelli, S., Pansoni, S., Frontoni, E., & Giovanola, B. (2024). Ethics of Artificial Intelligence for Cultural Heritage: Opportunities and Challenges. *IEEE Transactions on Technology and Society*. 5(3):293-305. DOI: 10.1109/TTS.2024.3432407.
- UNESCO (2025). *EXPANDED DEFINITION OF DOCUMENTARY HERITAGE*. <https://unesdoc.unesco.org/ark:/48223/pf0000395988#:~:text=Documentary%20heritage%20comprises%20those%20single,would%20be%20a%20harmful%20impoverishment>
- Wang, X., & Liu, B. (2023). Research on Artificial Intelligence Risk Prevention and Control System for Smart Libraries. In: *Emerging Technology-Based Services and Systems in Libraries, Educational Institutions, and Non-Profit Organizations*. <https://www.igi-global.com/chapter/research-on-artificial-intelligence-risk-prevention-and-control-system-for-smart-libraries/328666>
- Yazdi, M., Zarei, E., Adumene, S., & Beheshti, A. (2024). Navigating the Power of Artificial Intelligence in Risk Management: A Comparative Analysis. *Safety*, 10, 42. <https://doi.org/10.3390/safety10020042>
- Yin, Z. (2024). A review of methods for alleviating hallucination issues in large language models. *Applied and Computational Engineering*, 76, Article 20240608. <https://doi.org/10.54254/2755-2721/76/20240608>
- Yu, K., Gong, R., Sun, L., & Jiang, C. (2019). The application of artificial intelligence in smart library. In *2019 international conference on organizational innovation (ICOI 2019)* (pp. 708-713). Atlantis Press.