



<https://ecs.ui.ac.ir/?lang=en>

Journal of Endowment & Charity Studies

E-ISSN: 3115-7475

Vol.4, Issue.1, No.7, 2026, pp 113- 144

Received: 17/03/2026 Accepted: 17/05/2026

Research Paper

The Role of Algorithmic Governance in Charitable Organizations' Decision-Making: Opportunities and Threats of Shadow Management

Yasaman Charyani Zanjani 

Ph.D. Candidate, Department of Management, Cha.C., Islamic Azad University, Chalous, Iran
yasaman.charyani@iau.ac.ir

Davood Kia Kojouri * 

Associate Professor, Department of Management, Cha.C., Islamic Azad University, Chalous, Iran
dr.davoodkia@iau.ac.ir

Morteza Mousakhani 

Professor, Department of Management, Science and Research Branch, Islamic Azad University, Tehran, Iran
mosakhani@iau.ac.ir

Introduction

The increasing integration of Artificial Intelligence (AI) governance into social sectors, particularly non-profit organizations, has made algorithms a core component of management processes for aid allocation, needs assessment, and project prioritization. This evolution, while enhancing efficiency, introduces challenges such as algorithmic bias, opacity in decision-making, accountability issues, and potential privacy violations (Sari, 2025: 128; Barnes, 2025: 2). Often, these processes manifest as “shadow governance,” where algorithms invisibly influence organizational decisions, potentially perpetuating inequality or eroding public trust (Aizenberg et al., 2025: 921). Since AI systems are trained on historical data, there is a risk of reinforcing existing biases and leading to unintentional discrimination in resource distribution (Morse et al., 2020: 2; Akter et al., 2021: 2).

While most research has focused on AI applications in fields like education and policy (Yuanyuan, 2024: 105; Rakowski & Kowaliková, 2024: 2), the examination of algorithmic governance within charitable organizations remains limited (Furendal, 2023: 4). Consequently, the central research problem is to understand how algorithmic governance shapes the opportunities and threats of shadow management in non-profits and to identify the necessary mechanisms to enhance transparency, accountability, and fairness in this process. The study aims to provide a conceptual framework for the ethical management of this phenomenon. Its significance lies in mitigating bias risks and bolstering public trust, though the study is limited by a scarcity of empirical studies in this domain.

*Corresponding Author

Charyani Zanjani, Y. Kia Kojouri, D., & Mousakhani, M. (2026). The role of algorithmic governance in charitable organizations' decision-making: opportunities and threats of shadow management. *Journal of Endowment & Charity Studies*, 4(1), 113-144. <https://doi.org/10.22108/ecs.2026.148576.1180>

3115-7475© The Author(s). Published by University of Isfahan

This is an open access article under the CC BY-NC 4.0 License (<https://creativecommons.org/licenses/by-nc/4.0>).



<https://doi.org/10.22108/ecs.2026.148576.1180>

Methodology

This research employed an exploratory mixed-methods approach across two stages to analyze the role of algorithmic governance in charitable organizations' decision-making and to elucidate the opportunities and threats of shadow management. The first stage involved a systematic meta-synthesis of credible academic sources from English (Web of Science, Scopus, Google Scholar, ScienceDirect) and Persian (SID, Noormags, Civilica, Elmnet) databases. From 142 initial articles, 60 English and 6 Persian sources were selected for final analysis after rigorous screening based on title, abstract, quality, and content relevance. This meta-synthesis identified key dimensions, components, and indicators of algorithmic governance's opportunities and threats, culminating in a comprehensive, multidimensional conceptual framework.

The findings from the meta-synthesis informed the design of the data collection instrument for the second, qualitative stage. A semi-structured questionnaire, based on the identified dimensions and components, was developed for expert interviews. These interviews aimed to gather specialized perspectives on the opportunities and threats of shadow management emerging from the expansion of algorithmic governance. This approach revealed hidden aspects of the phenomenon by integrating theoretical foundations with expert experiences. Accordingly, the opportunities and threats of shadow management in charities were investigated using thematic analysis with the participation of 10 experts (professors of public administration and specialists in ethics, governance, and algorithmic management). Data were collected through in-depth, focused interviews, analyzed via open, axial, and selective coding, and the reliability of the analysis was confirmed.

Findings

In the first stage, this study utilized meta-synthesis and the systematic screening of 66 peer-reviewed sources to develop a comprehensive framework of the dimensions, components, and indicators of algorithmic governance. The findings reveal that algorithmic governance presents a complex duality of capacity and risk. Opportunities include enhanced transparency and accountability, improved data and model performance, elevated employee experience, increased organizational agility, and optimized decision-making. Conversely, associated threats involve data quality issues, model errors, algorithmic bias, ethical and privacy violations, psychological stress, and operational security risks. This framework provided the foundation for designing the data collection instrument and the subsequent thematic analysis.

In the second stage, thematic analysis of expert interviews identified the opportunities and threats of "shadow management" within algorithmic governance in charitable organizations. Results demonstrate that the intelligent application of data and algorithms can transform charities from traditional intermediaries into active social governance agents by empowering independent decision-making, fostering synergy with state entities, producing field-based knowledge, and strengthening social capital. However, structural threats were also identified, such as the surreptitious redesign of decision-making processes, monopolization of data interpretation, mission drift, the reproduction of inequality, erosion of accountability, and weakened organizational learning. Ultimately, while shadow management may enhance agility and innovation, it poses significant challenges to the sustainability and legitimacy of charitable organizations if robust governance mechanisms are absent. The final conceptual model is presented in [Figure 1](#).



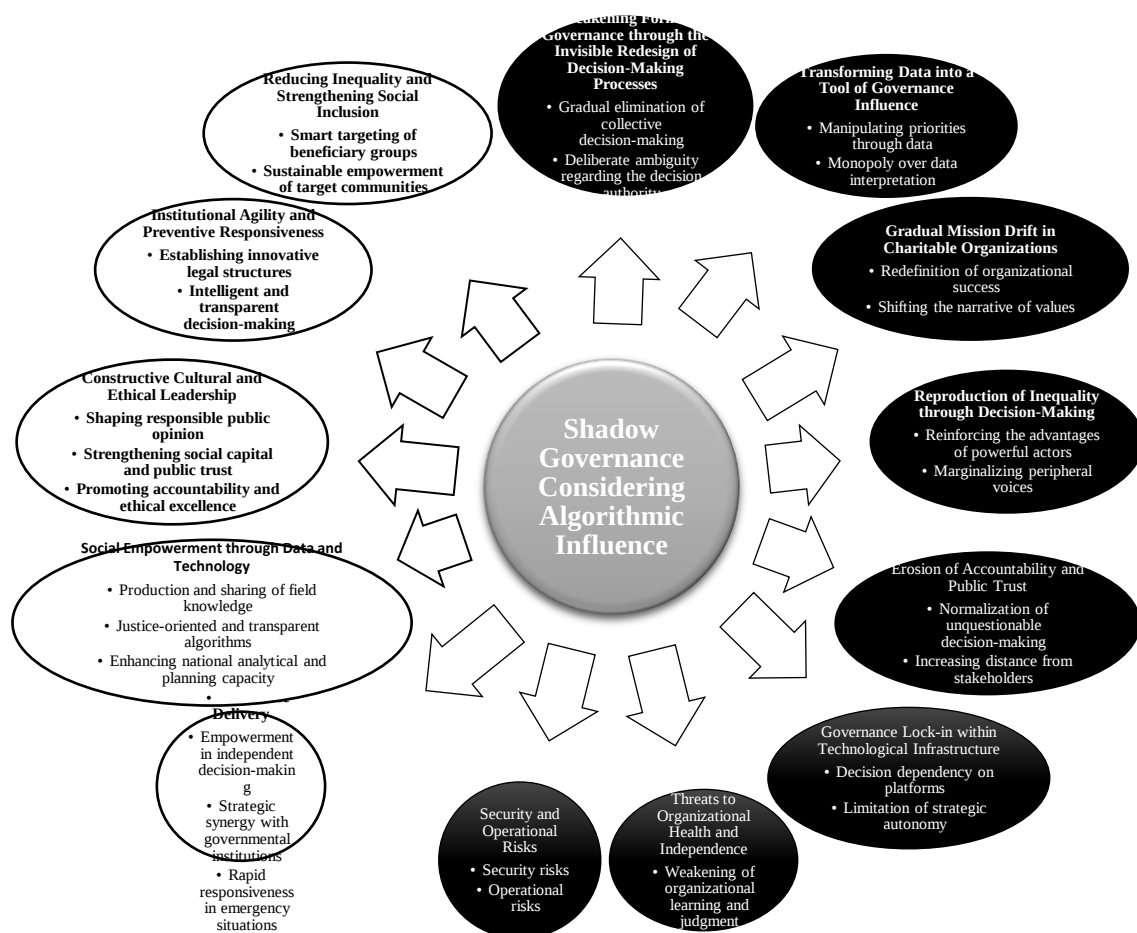


Figure 1. Final Model of the Threats and Opportunities of Shadow Management Considering Algorithmic Governance

Discussion and Conclusion

The findings of this study indicate that algorithmic governance in charitable organizations is a dual-faceted phenomenon that can simultaneously create pathways for enhanced efficiency, equity, and agility, while also enabling the emergence of shadow management and structural threats. Within the proposed framework, we identified opportunities such as strengthened decision-making capacities, the production and use of reliable data, the development of fairness-oriented algorithms, and enhanced institutional collaboration. Conversely, threats such as the weakening of formal decision-making processes, data monopolization, erosion of accountability, reproduction of inequalities, and technological lock-in demonstrate that the unregulated application of algorithms may jeopardize organizational integrity and legitimacy. Accordingly, it is essential for charities to adopt a balanced approach that leverages the transformative potential of data and algorithms while preventing hidden and informal decision-making mechanisms. Based on the results, it is recommended that charities and policymaking bodies develop a framework for responsible algorithmic governance, including standards for transparency, bias auditing, meaningful human oversight, and data protection.

Despite the systematic design of this study, its focus was primarily at the organizational level and based on a limited number of experts, with less attention paid to cultural, legal, and technological contextual factors. Therefore, future research should employ quantitative methods and causal modeling to test these relationships, conduct cross-country comparative studies, and examine the role of regulatory bodies in mitigating threats and reinforcing opportunities. Field-based investigations into employees' interactions with algorithms and studies on small and community-based charities may further enrich the literature in this domain.

Keywords

Algorithmic Governance, Shadow Management, Charitable Organizations, Algorithmic Transparency, Algorithmic Fairness, Responsible AI Governance.



مقاله پژوهشی

نقش حکمرانی الگوریتمی در تصمیم‌گیری خیریه‌ها: فرصت‌ها و تهدیدهای مدیریت سایه

یاسمن چریانی‌زنجانی^۱ ، داود کیاکجوری^۲  و مرتضی موسی‌خانی^۳ 

۱. دانشجوی دکتری گروه مدیریت، واحد چالوس دانشگاه آزاد اسلامی، چالوس، ایران

yasaman.charyani@iau.ac.ir

۲. دانشیار گروه مدیریت، واحد چالوس دانشگاه آزاد اسلامی، چالوس، ایران

dr.davoodkia@iau.ac.ir

۳. استاد گروه مدیریت، دانشگاه آزاد اسلامی، واحد علوم و تحقیقات، تهران، ایران

mosakhani@iau.ac.ir

چکیده

با گسترش حکمرانی الگوریتمی در تصمیم‌گیری سازمان‌های خیریه، پدیده مدیریت سایه به عنوان قدرت نامرئی با علم، اشراف و تأثیرگذاری بر الگوریتم‌ها، فرصت‌ها و تهدیدهایی را با خود به همراه آورده است. پژوهش حاضر، با هدف واکاوی نقش حکمرانی الگوریتمی در تصمیم‌گیری خیریه‌ها و تبیین فرصت‌ها و تهدیدهای مدیریت سایه در این بستر، از رویکرد ترکیبی اکتشافی بهره گرفت. در گام اول، روش فراترکیب به کار گرفته شد. برای این منظور، تحلیل بر روی ۶۶ منبع معتبر (۶۰ انگلیسی و ۶ فارسی) انجام شد. نتایج فراترکیب زمینه‌ساز تدوین پرسشنامه نیمه‌ساختاریافته تحلیل مضمون با مشارکت ۱۰ خبره از حوزه مدیریت دولتی و خیریه شد. در نتیجه تحلیل مضمون، فرصت‌ها و تهدیدهای مدیریت سایه در بستر حکمرانی الگوریتمی در سازمان‌های خیریه استخراج شدند. یافته‌ها نشان‌دهنده ماهیت پارادوکسیکال حکمرانی الگوریتمی در سازمان‌های خیریه بود که از یک سو، فرصت‌هایی همچون حکمرانی و شفافیت الگوریتمی، عدالت، اخلاق و کاهش سوگیری، کیفیت داده و مدل، تأثیر بر کارکنان و فرهنگ سازمانی و کارایی، کارکرد و نتایج کسب‌وکار مثبت را فراهم می‌کند، و از سوی دیگر، سبب تهدیدهایی مانند ریسک‌های کیفیت داده و مدل، ریسک‌های اخلاقی و تبعیض الگوریتمی و پیامدهای منفی مدیریت الگوریتمی بر کارکنان می‌شود. در این بستر بود که مدیریت سایه سبب بروز فرصت‌ها (مانند تعالی مرجعیت خیریه، هدایت‌گری فرهنگی، توانمندسازی اجتماعی، جابجایی نهادی و کاهش نابرابری) و تهدیدها (مانند تضعیف حکمرانی رسمی، چرخش خزنده مأموریت، فرسایش اعتماد عمومی و قفل‌شدگی فناوریانه) شد. ضمن ایجاد ظرفیت‌های تحول‌آفرین، حکمرانی الگوریتمی بدون حکمرانی مسئولانه به تهدیدهای جدی منجر می‌شود. پیشنهادها شامل تدوین منشور اخلاقی، ممیزی مستمر، آموزش سواد الگوریتمی و تقویت نظارت انسانی هستند.

واژه‌های کلیدی

حکمرانی الگوریتمی، مدیریت سایه، سازمان‌های خیریه، شفافیت الگوریتمی، حکمرانی مسئولانه هوش مصنوعی.

*نویسنده مسئول

چریانی‌زنجانی، ی. کیاکجوری، د. و موسی‌خانی، م. (۱۴۰۵). نقش حکمرانی الگوریتمی در تصمیم‌گیری خیریه‌ها: فرصت‌ها و تهدیدهای مدیریت سایه. *مطالعات وقف و امور خیریه*، ۴(۱)، ۱۱۳-۱۴۴. <https://doi.org/10.22108/ecs.2026.148576.1180>



۱- مقدمه و بیان مسأله

با گسترش روزافزون هوش مصنوعی (AI) در حوزه‌های مختلف اجتماعی، از جمله فعالیت‌های خیریه و بشردوستانه، سازمان‌های خیریه با چالش‌های اخلاقی و مدیریتی جدیدی مواجه شده‌اند. استفاده از الگوریتم‌ها در فرایندهایی مانند تخصیص کمک‌ها، ارزیابی نیازمندان، اولویت‌بندی پروژه‌ها و تصمیم‌گیری‌های راهبردی، مسائلی مهم در حوزه شفافیت، سوگیری، حریم خصوصی و مسئولیت‌پذیری ایجاد کرده است (Sari, 2025: 128). حکمرانی الگوریتمی که معمولاً به صورت مدیریت سایه ظاهر می‌شود، به معنای قدرت نامرئی و غیررسمی الگوریتم‌ها در هدایت تصمیم‌هاست که فرصت‌هایی را برای افزایش کارایی و دسترسی به کمک‌ها فراهم می‌کند، اما هم‌زمان، تهدیدهایی مانند بازتولید نابرابری‌ها و کاهش اعتماد عمومی را به همراه دارد (Akter et al., 2021: 2; Aizenberg et al., 2025: 921). در بخش خیریه که معمولاً بر پایه ارزش‌های اخلاقی و اجتماعی بنا شده است، سیستم‌های هوش مصنوعی بر اساس داده‌های تاریخی آموزش می‌بینند و ممکن است سوگیری‌های موجود را تقویت کنند که این امر ممکن است به تبعیض ناخواسته در تخصیص منابع منجر شود (Morse et al., 2021: 2; Akter et al., 2020: 2). علاوه بر این، با وجود پیچیدگی الگوریتم‌ها و ماهیت جعبه‌سیاه آنها، شفافیت کاهش می‌یابد و پاسخ‌گویی دشوار می‌شود (Sari, 2025: 128; Zhong et al., 2025: 3). پایش داده‌های شخصی ذی‌نفعان نیز نگرانی‌های جدی در حوزه حریم خصوصی ایجاد می‌کند (Akter et al., 2021: 3; Shi & Xing, 2025: 2) و اتکای بیش از حد به تصمیم‌های خودکار ممکن است استقلال انسانی و قضاوت اخلاقی را در فعالیت‌های خیریه کاهش دهد (El Arab et al., 2025: 13).

حکمرانی الگوریتمی یعنی تصمیم‌های اجرایی (انتخاب پروژه، اولویت‌بندی دریافت‌کننده، ارزیابی عملکرد و ...) به طور سیستماتیک توسط مدل‌ها و داشبوردها هدایت شود، نه فقط مدیران انسانی (Zhang et al., 2025: 4; Keegan & Meijerink, 2024: 5; Stark & Vanden Broeck, 2025: 396). در ترکیب انسان-الگوریتم، بخشی از پاسخ‌گویی به داشبوردها و معیارهای کمی منتقل می‌شود؛ این امر ممکن است اختیار لایه‌های میانی و امنیت روانی کارکنان را کاهش دهد و نوعی پاسخ‌گویی در سایه بسازد که در آن، قدرت واقعی در منطق شاخص‌هاست، نه ساختار رسمی (Barnes, 2025: 149; Stark & Vanden Broeck, 2024: 2). در سطح کل سازمان، استفاده غیررسمی از ابزارهای هوش مصنوعی خارج از چارچوب‌های مصوب به عنوان هوش مصنوعی سایه^۱ توصیف شده است؛ یعنی کارمندان از ابزارها برای تصمیم‌های حساس (برای مثال، غربال کمک‌گیرندگان) استفاده می‌کنند، بدون اینکه این استفاده در سیاست‌ها، ناظر و ممیزی منعکس شده باشد (Ross et al., 2025: 2).

برای مقابله با این چالش‌ها، سازمان‌های خیریه می‌توانند از اصول اخلاق سازمانی بهره ببرند، مانند ایجاد چارچوب‌های حاکمیت اخلاقی، مدیریت ذی‌نفعان، شفاف‌سازی تصمیم‌های الگوریتمی، حفظ نظارت انسانی و آموزش کارکنان (Schultz, 2025: 921; Seele, 2022: 101; Salehi, 2025: 233; Aizenberg et al., 2025: 921). پژوهش حاضر، با تمرکز بر فرصت‌ها و تهدیدهای ایجادشده به واسطه حکمرانی الگوریتمی در تصمیم‌گیری خیریه‌ها، فرصت‌ها و تهدیدهای مدیریت سایه را بررسی و الگویی برای تعادل نوآوری فناورانه با مسئولیت اخلاقی پیشنهاد می‌کند. این رویکرد می‌تواند به تقویت اعتماد و عدالت در فعالیت‌های خیریه کمک کند؛ جایی که قدرت نامرئی الگوریتم‌ها روابط قدرت و ارزش‌های ادراکی را بازتعریف می‌کند (Shaw, 2025: 118). بدون تعبیه پیش‌دستانه اخلاق در ساختار الگوریتمی، پیامدهایی مانند تثبیت سوگیری‌ها و کاهش اعتماد ذی‌نفعان اجتناب‌ناپذیر خواهند بود (De Cremer & McGuire, 2022: 35; Liu et al., 2025: 345).

¹ Shadow AI

کاربرد الگوریتم تأمین مالی مبتنی بر پیوندهای اجتماعی^۱ در پلتفرم‌های خیریه‌محور به سبک یک پلتفرم تأمین مالی جمعی غیرمتمرکز برای حمایت از پروژه‌های عمومی دیجیتال، با تکیه بر ایده حکمرانی باز توزیع یارانه بر اساس پیوندهای اجتماعی و مطلوبیت‌های نوع‌دوستانه به جای صرف تعداد یا مبلغ کمک‌های انفرادی، توانسته است نسبت به تأمین مالی کلاسیک، به بهبود رفاه اجتماعی و افزایش عدالت در توزیع منابع میان پروژه‌ها منجر شود (Miller et al., 2025: 108). همچنین، به‌کارگیری هم‌زمان بلاک‌چین و یادگیری ماشین در تخصیص منابع خیریه، با استفاده از قراردادهای هوشمند برای شفافیت و قفل‌شدن قواعد و الگوریتم‌ها برای تشخیص تقلب، تطبیق اهداکننده و پیش‌بینی موفقیت کمپین‌ها، موجب افزایش شفافیت و ردیابی تراکشن‌ها، کاهش تقلب و بهبود کارایی تخصیص و تطبیق نیازمندان شده است (Raphael, 2025: 2; Lakshmanan et al., 2024: 849). طراحی ساختارهای تعهد مالی دیجیتال (مانند تعهد سخت یا همراه با ضرب‌الاجل) می‌تواند ریسک نکول در جمع‌آوری کمک‌های مردمی را به طرزی معنادار کاهش دهد و پایداری مالی پروژه‌های خیریه را تقویت کند (عباس‌نیا و عبدی، ۱۴۰۴: ۱۴۱). در کنار این موارد، الگوریتم‌هایی که به قضاوت و جمع‌بندی در خصوص چگونگی جمع‌آوری کمک‌ها و همچنین فرایند راستی‌آزمایی اختصاص دارند، توانسته‌اند درخواست‌های جعلی را کاهش دهند و اعتماد اهداکنندگان را حفظ کنند؛ هرچند همچنان نیازمند نظارت انسانی و پیش‌بینی پروتکل‌های اعتراض و بازبینی هستند (Vasuki, 2025: 4). البته تبلیغات رسانه‌ای و برندسازی دیجیتال مؤسسه‌های خیریه از طریق مراحل آغاز تا بلوغ و توجه به عوامل سازمانی، اجتماعی و فناورانه، می‌تواند اعتماد اهداکنندگان و مددجویان را افزایش دهد و پیامدهایی مثبت برای سازمان و جامعه ایجاد کند (کیاکجوری و چریانی‌زنجانی، ۱۴۰۳: ۹۱؛ صاحب‌الداری و همکاران، ۱۴۰۴: ۱۹۹).

حکمرانی الگوریتمی که خیریه‌ها از آن بهره می‌گیرند، به صورت مدیریت سایه عمل می‌کند که تصمیم‌گیری را از حوزه انسانی به مدل‌های محاسباتی پنهان منتقل می‌کند (Martin, 2019: 837; Jarrahi et al., 2021: 3). این وضعیت چالش‌هایی مانند سازوکارهای عدم پاسخ‌گویی شفاف، واگذاری مسئولیت تصمیم‌گیری به الگوریتم توسط مدیران و ابهام الگوریتمی ایجاد می‌کند (Benlian et al., 2022: 827; Morse et al., 2020: 2). سوگیری الگوریتمی در خیریه‌ها می‌تواند عدالت در تخصیص کمک‌ها را تهدید کند و گروه‌های آسیب‌پذیر را در معرض تبعیض قرار دهد (Akter et al., 2021: 3; Gupta et al., 2022: 3). پژوهش‌های پیشین بر جنبه‌های کلان حکمرانی هوش مصنوعی تمرکز داشته‌اند، مانند کاربردها در سلامت (اقتدار و همکاران، ۱۴۰۲: ۲۷۳)، آموزش (Yuanyuan, 2024: 105) و سیاست (Rakowski & Kowaliková, 2024: 2)، اما کمتر به نقش الگوریتم‌ها در تصمیم‌گیری خیریه‌ها و تهدیدهای مدیریت سایه توجه کرده‌اند (دهقانی و همکاران، ۱۴۰۴: ۲؛ قاسمی، ۱۴۰۰: ۳؛ Furendal, 2023: 4).

فقدان مدلی بومی برای ارتقای اخلاق و بهره‌وری در بهره‌گیری از حکمرانی الگوریتمی در خیریه‌ها شکافی آشکار در ادبیات ایجاد کرده است؛ جایی که سازوکارهای شفافیت، عدالت الگوریتمی و نظارت انسانی نادیده گرفته شده‌اند (Carrillo, 2020: 4; Jensen et al., 2020: 3). پژوهش حاضر این خلأ را پر می‌کند و بر طراحی الگویی تمرکز دارد که رابطه میان شفافیت، پاسخ‌گویی و فرهنگ سازمانی را در مواجهه با مدیریت سایه تبیین کند (Hunkenschroer & Luetge, 2022: 978; Schmitt, 2024: 127; Hernández-Lugo, 2024: 8). این الگو می‌تواند فرصت‌ها و تهدیدهای نوآوری الگوریتمی را در برابر ارزش‌های والای خیریه نمایان کند تا فقط به فرصتی برای تقویت عدالت و اعتماد تبدیل شود (Leicht-Deobald et al., 2022: 3; Liao, 2023: 3). بر این اساس، این پژوهش

¹ Connection-Oriented Quadratic Funding (CO-QF)

به دنبال ارائه مدلی شامل فرصت‌ها و تهدیدهای حکمرانی الگوریتمی و در پی آن، ارائه فرصت‌ها و تهدیدهای مدیریت سایه‌ای است که در پی حکمرانی الگوریتمی در تصمیم‌گیری خیریه‌ها به وجود آمده است.

۲- ادبیات پژوهش

۲-۱ حکمرانی الگوریتمی در تصمیم‌گیری سازمان‌های غیرانتفاعی

حکمرانی الگوریتمی در تصمیم‌گیری سازمان‌های غیرانتفاعی، به ویژه خیریه‌ها، به معنای ادغام سیستم‌های هوش مصنوعی در فرایندهای کلیدی مانند تخصیص منابع، ارزیابی نیازمندان و راهبردهای جمع‌آوری کمک است که می‌تواند از یک سو، شفافیت را کاهش دهد و سوگیری‌های موجود را تقویت کند و از سوی دیگر، کارایی، سرعت و دقت تصمیم‌گیری را افزایش دهد و امکان تخصیص بهینه‌تر منابع را فراهم کند. بر اساس پژوهش‌های موجود، الگوریتم‌ها در کنار ارتقای کارایی بیش از حد، معمولاً بر پایه داده‌های تاریخی آموزش می‌بینند که ممکن است سوگیری‌های اجتماعی را بازتولید کنند که این امر ممکن است به تبعیض ناخواسته در توزیع کمک‌ها منجر شود (Akter et al., 2021: 3; Aizenberg et al., 2025: 922). در بخش خیریه، حکمرانی الگوریتمی می‌تواند تصمیم‌گیری را از قضاوت انسانی به سمت مدل‌های محاسباتی سوق دهد؛ جایی که در کنار قابلیت‌هایی که در بهبود و تسریع تصمیم‌گیری به همراه دارد، پیچیدگی الگوریتم‌ها شفافیت را کاهش می‌دهد و مسئولیت‌پذیری را مبهم می‌کند (Sari, 2025: 129; Zhong et al., 2025: 2).

در سازمان‌های خیریه، حکمرانی الگوریتمی فرصت‌هایی برای بهینه‌سازی فراهم می‌کند، مانند پیش‌بینی نیازهای جامعه یا هدف‌گیری دقیق‌تر کمک‌ها، اما تهدیدهایی مانند کاهش استقلال انسانی و افزایش وابستگی به فناوری را نیز به همراه دارد (El Arab et al., 2025: 2; Shi & Xing, 2025: 3). پژوهش‌ها نشان می‌دهند به‌کارگیری الگوریتم‌ها در تصمیم‌گیری سازمانی می‌تواند از یک سو، دقت تحلیلی، سرعت پردازش داده‌ها و انسجام تصمیم‌ها را افزایش دهد، اما اتکای بیش از حد به الگوریتم‌ها می‌تواند قضاوت حرفه‌ای را فرسوده کند و تصمیم‌گیری را به فرایندی مکانیکی تبدیل کند (De Cremer & McGuire, 2022: 35; Liu et al., 2025: 343). در رابطه با سازمان‌های خیریه، این دوگانگی اهمیت طراحی چارچوب‌های اخلاقی و نظارتی را برجسته می‌کند تا بهره‌گیری از نوآوری‌های فناورانه در عین حفظ ارزش‌های انسانی و مسئولیت اجتماعی صورت گیرد (Schultz & Seele, 2022: 101; Salehi, 2025: 234).

۲-۲ مدیریت سایه و چالش‌های مدیریتی در بخش خیریه

مدیریت سایه در بخش خیریه به معنای ساختارهای قدرت نامرئی و غیررسمی است که تصمیم‌گیری را خارج از چارچوب‌های رسمی هدایت می‌کند و با حکمرانی الگوریتمی می‌تواند به شکلی جدید از کنترل پنهان تبدیل شود (Morse et al., 2020: 3). در خیریه‌ها، این مدیریت می‌تواند در پشت پرده از طریق الگوریتم‌های جعبه‌سیاه که روابط قدرت را بازتعریف می‌کنند، تقویت شود و به تثبیت سوگیری‌های تاریخی و کاهش اعتماد منجر شود (Akter et al., 2021: 3; Gupta et al., 2022). بدون نظارت انسانی، الگوریتم‌ها می‌توانند کنترل‌گری را تشدید و روابط کاری را غیرشخصی کنند (Noponen, 2019: 44; Tursunbayeva, 2024: 118).

چالش‌های مدیریتی در مدیریت سایه خیریه‌ها شامل ابهام در پاسخ‌گویی و کاهش شفافیت هستند؛ جایی که الگوریتم‌ها تصمیم‌ها را بدون توضیح قابل فهم اتخاذ و کارکنان را در برابر نتایج آسیب‌پذیر می‌کنند (Liao, 2023: 3; Martin, 2019).



837). برای مقابله، سازمان‌های خیریه به رویکردهای پیش‌بینانه نیاز دارند که اخلاق را در طراحی الگوریتم ادغام کنند و نقش‌هایی مانند رئیس ارشد هوش مصنوعی را برای پل‌زنی بین فناوری و مسئولیت اجتماعی ایجاد کنند (Hunkenschroer & Luetge, 2022; 2: 979; Schmitt, 2024). کاهش تصدی‌گری دولت و تشکیل شورای ملی تعاون می‌تواند هماهنگی و نظارت بر فعالیت‌های خیریه را بهبود بخشد و از مدیریت سایه جلوگیری کند (کیاکجوری، ۱۴۰۳: ۱۸۵). این امر می‌تواند عدالت سازمانی را تقویت و از تهدیدهای سایه جلوگیری کند (Hernández-Lugo, 2024: 4; Kandasamy et al., 2025: 1077).

۲-۳ فرصت‌ها و تهدیدهای ادغام هوش مصنوعی در تصمیم‌گیری خیریه

فرصت‌های ادغام هوش مصنوعی در تصمیم‌گیری خیریه شامل افزایش کارایی در تخصیص کمک‌ها و پیش‌بینی نیازها هستند؛ جایی که الگوریتم‌ها می‌توانند داده‌های بزرگ را تحلیل و نوآوری را با مسئولیت اخلاقی متوازن کنند (Akter et al., 2025: 2; Kapate et al., 2021: 3). دیجیتالی‌شدن خدمات خیریه همراه با شمول مالی، اثر تعاملی منفی و معناداری بر فقر چندبعدی دارد و می‌تواند به کاهش متوسط ۱۲ درصدی فقر و هم‌افزایی در سیاست‌های مداخله اجتماعی هوشمند کمک کند (کاظمی نجف‌آبادی، ۱۴۰۴: ۱۸۷). پژوهش‌ها نشان می‌دهند چارچوب‌های حاکمیت اخلاقی می‌توانند شفافیت را افزایش دهند و اعتماد را در بخش خیریه تقویت کنند (Aizenberg et al., 2025: 921; Schultz & Seele, 2022: 99). این ادغام می‌تواند سازمان‌های خیریه را به سمت تصمیم‌گیری داده‌محور سوق دهد و عدالت را در توزیع منابع بهبود بخشد (Jarrahi, 2022: 827; Benlian et al., 2021: 2; et al.). تهدیدهای این ادغام شامل بازتولید سوگیری الگوریتمی و کاهش پاسخ‌گویی هستند که می‌تواند در خیریه‌ها به تبعیض در کمک‌رسانی منجر شوند (Akter et al., 2021: 3; Morse et al., 2020: 3). پیچیدگی الگوریتم‌ها می‌تواند شفافیت را کاهش دهد و کارکنان را به تبعیت منفعلانه سوق دهد که این امر تهدیدی برای اخلاق سازمانی است (De Cremer & McGuire, 2022: 35; Leicht-Deobald et al., 2022: 73). برای کاهش تهدیدها، آموزش مستمر و ممیزی الگوریتم‌ها ضروری است تا سازمان‌های خیریه بدون افتادن در دام مدیریت سایه و فقط با جذب فرصت‌های آن بتوانند از فرصت‌ها بهره ببرند (Liao, 2023: 5).

۲-۴ مروری بر پیشینه پژوهش

پژوهش‌های خارجی بر حکمرانی الگوریتمی، مدیریت سایه و چالش‌های اخلاقی هوش مصنوعی در بخش عمومی تأکید دارند و فرصت‌ها و تهدیدهای تصمیم‌گیری الگوریتمی را تحلیل می‌کنند. لیو^۱ و همکاران (2025) با روش ترکیبی، نشان دادند مدیریت الگوریتمی از طریق مالکیت روان‌شناختی، پنهان‌سازی دانش را افزایش می‌دهد که این امر ممکن است به مدیریت سایه منجر شود. آراگانی^۲ و همکاران (2025) با پیمایش و تحلیل کیفی، فرصت‌های هوش مصنوعی در تقویت دموکراسی را بررسی و هم‌زمان، چالش‌های عدالت و اخلاق را برجسته می‌کنند. باتلر^۳ (2025) با تحلیل حقوقی، اصل عدم تفویض اختیار را در مواجهه با تصمیم‌گیری الگوریتمی ناسازگار می‌داند و خطر کاهش پاسخ‌گویی را هشدار می‌دهد.

¹ Liu

² Aragani

³ Butler

همچنین، خاماس^۱ و همکاران (۲۰۲۵) با مدل‌سازی ساختاری، رهبری سایه را به عنوان سازوکاری موازی برای نفوذ توصیف می‌کنند که شفافیت را تهدید می‌کند. این مطالعات بیشتر به بخش دولتی اختصاص دارند و کمتر به سازمان‌های خیریه توجه کرده‌اند، اما تهدیدهای سوگیری الگوریتمی و کاهش شفافیت را برجسته می‌کنند که می‌تواند به مدیریت سایه در تصمیم‌گیری خیریه‌ها تعمیم یابد.

پژوهش‌های داخلی عمدتاً بر نقش هوش مصنوعی در حکمرانی و سیاست عمومی تمرکز دارند و چالش‌ها و فرصت‌های آن را در زمینه‌های دولتی بررسی می‌کنند. برای مثال، ابوذری (۱۴۰۲) هوش مصنوعی را به عنوان ابزاری برای قدرت نرم در سیاست عمومی معرفی می‌کند و پیشنهاد می‌دهد دولت‌ها برای کارآمدی حکمرانی از آن بهره ببرند. همچنین، قاسمی (۱۴۰۰) با تحلیل توصیفی-تحلیلی، بر تطبیق هوش مصنوعی با ارزش‌های جامعه و هشدار نسبت به وابستگی فناورانه تأکید دارد، در حالی که روشن و همکاران (۱۴۰۰) با فراترکیب، کاربرد هوش مصنوعی در بخش دولتی را بررسی می‌کنند و بر سیاست‌گذاری اخلاقی برای ارتقای خدمات عمومی تمرکز دارند. این پژوهش‌ها بیشتر به جنبه‌های کلان حکمرانی اشاره دارند و کمتر به حکمرانی الگوریتمی در سازمان‌های غیرانتفاعی مانند خیریه‌ها توجه کرده‌اند، اما چالش‌های اخلاقی و شفافیت را برجسته می‌کنند که ممکن است به فرصت‌ها و تهدیدهای مدیریت سایه مرتبط باشند.

در مجموع، پیشینه پژوهش‌های داخلی و خارجی بر کاربرد هوش مصنوعی در حکمرانی عمومی، چالش‌های اخلاقی، سوگیری الگوریتمی و سازوکارهای سایه تمرکز دارد، اما بیشتر به بخش دولتی محدود است و جنبه‌های خاص سازمان‌های غیرانتفاعی مانند خیریه‌ها را نادیده می‌گیرد. پژوهش‌های داخلی مانند قاسمی (۱۴۰۰) و روشن و همکاران (۱۴۰۰) بر سیاست‌گذاری اخلاقی تأکید دارند، در حالی که خارجی‌ها مانند لیو و همکاران (۲۰۲۵) و خاماس و همکاران (۲۰۲۵) تهدیدهای موجود در حکمرانی الگوریتمی و مدیریت سایه را بررسی می‌کنند، اما هیچ کدام به طور مستقیم به تهدیدها و فرصت‌های موجود در تصمیم‌گیری در خیریه‌ها در عصر حکمرانی الگوریتمی اشاره نمی‌کنند. این شکاف پژوهشی ضرورت بررسی نقش حکمرانی الگوریتمی در تصمیم‌گیری خیریه‌ها را برجسته می‌کند؛ جایی که فرصت‌های بهینه‌سازی توزیع کمک‌ها با تهدیدهای مدیریت سایه مانند کاهش شفافیت و سوگیری در برابر نیازمندان همراه است و پژوهش حاضر با تمرکز بر این حوزه، خلأ موجود را پر می‌کند.

۳- روش پژوهش

در این پژوهش، با هدف واکاوی نقش حکمرانی الگوریتمی در تصمیم‌گیری خیریه‌ها و تبیین فرصت‌ها و تهدیدهای مدیریت سایه، از یک رویکرد ترکیبی و اکتشافی در دو مرحله بهره گرفته شد. در مرحله نخست پژوهش، به منظور شناسایی ابعاد و مؤلفه‌های فرصت‌ها و تهدیدهای حکمرانی الگوریتمی در سازمان‌ها، از روش فراترکیب به صورت نظام‌مند استفاده شد. در این راستا، جست‌وجوی منابع علمی در پایگاه‌های بین‌المللی شامل Scholargoogle, Scopus, Web of Science و Sciencedirect و نیز پایگاه‌های فارسی مانند SID, Noormags, Civilica, Elmnet انجام شد. در مرحله شناسایی اولیه، در مجموع، ۱۴۲ منبع علمی استخراج شدند (شکل ۱). پس از غربالگری اولیه بر اساس عنوان، ۱۰۵ مقاله باقی ماندند که چکیده آنها بررسی شد. در ادامه، با اعمال معیارهای ورود و خروج پژوهش، از جمله ارتباط مستقیم با موضوع حکمرانی الگوریتمی،

¹ Khamas



دسترسی به متن کامل، کیفیت علمی مناسب و ارتباط دقیق با موضوع، تعداد منابع کاهش یافت. در نهایت، ۶۶ منبع شامل ۶۰ مقاله انگلیسی و ۶ منبع فارسی به عنوان مطالعات واجد شرایط انتخاب و برای تحلیل نهایی در فراترکیب استفاده شدند. تحلیل محتوای این متون با استفاده از رویکرد تحلیل اهمیت-عملکرد و تکنیک کدگذاری موضوعی انجام شد که به استخراج چارچوب فرصت‌ها و تهدیدهای حکمرانی الگوریتمی منجر شد. خروجی این مرحله ترسیم تصویری جامع از ساختار چندبُعدی فرصت‌ها و تهدیدهای حکمرانی الگوریتمی بود. یافته‌های حاصل از مرحله فراترکیب (جدول ۳) مبنای طراحی ابزار گردآوری داده‌ها در مرحله کیفی دوم پژوهش قرار گرفت. برای این منظور، بر اساس ابعاد و مؤلفه‌های استخراج‌شده از تحلیل متون، یک پرسشنامه نیمه‌ساختاریافته تدوین شد تا از طریق مصاحبه با خبرگان، دیدگاه‌های تخصصی آنان با هدف شناسایی و تبیین فرصت‌ها و تهدیدهای شکل‌گیری مدیریت سایه در پرتو گسترش حکمرانی الگوریتمی جمع‌آوری شوند تا از طریق تلفیق مبانی نظری و تجربیات خبرگان، ابعاد پنهان این پدیده در سازمان‌های خیریه آشکار شوند. در همین راستا، فرصت‌ها و تهدیدهای مدیریت سایه در سازمان‌های خیریه شناسایی شدند. این بخش با به‌کارگیری روش تحلیل مضمون و از طریق مشارکت ۱۰ خبره شامل اساتید حوزه مدیریت دولتی و صاحب‌نظران دارای سابقه اجرایی و پژوهشی چشمگیر در زمینه اخلاق، حکمرانی و مدیریت الگوریتمی انجام شد (جدول‌های ۴ و ۵). (جدول ۱) مشخصات کلی این افراد را بر اساس رشته تخصصی، سابقه کاری و حوزه پژوهشی ارائه می‌کند.

جدول ۱. معرفی خبرگان مشارکت‌کننده در بخش تحلیل مضمون ناشی از فرصت‌ها و تهدیدهای مدیریت سایه در تصمیم‌گیری خیریه‌ها

Table 1. Introduction of experts participating in the thematic analysis of opportunities and threats arising from shadow management in charitable decision-making

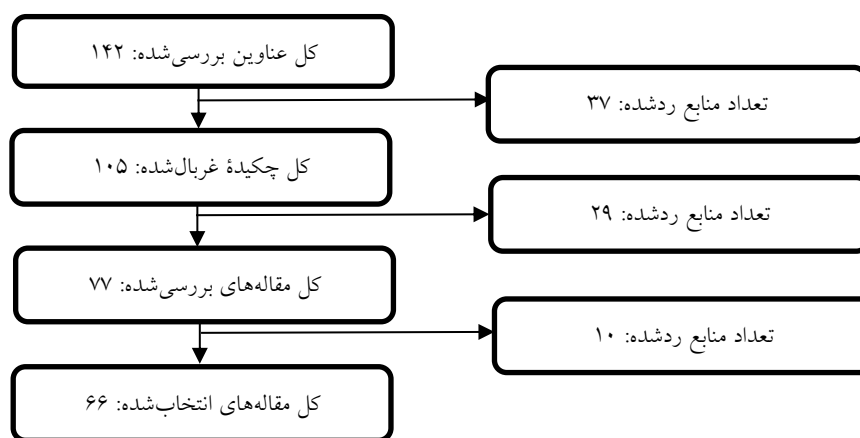
کد خبره	رشته تحصیلی/تخصص علمی	سابقه کار اجرایی/مدیریتی	حوزه‌های پژوهشی و تخصصی
خبره ۱	مدیریت دولتی	۱۵ سال تجربه مدیریتی در بخش دولتی	حکمرانی، سیاست‌گذاری عمومی
خبره ۲	مدیریت فناوری اطلاعات	۱۲ سال تجربه در پروژه‌های فناوری و داده	فناوری‌های دیجیتال، الگوریتم‌ها
خبره ۳	مدیریت دولتی	۱۸ سال تجربه در نهادهای عمومی	حکمرانی و عدالت سازمانی
خبره ۴	سیاست‌گذاری عمومی	۱۰ سال کار اجرایی در حوزه تصمیم‌سازی	سیاست‌گذاری فناورانه
خبره ۵	مدیریت دولتی-خط‌مشی‌گذاری	۱۴ سال تجربه مشاوره سازمانی	اخلاق حرفه‌ای، رفتار سازمانی
خبره ۶	مدیریت دولتی	۲۰ سال مدیریت در سازمان‌های عمومی	حکمرانی و شفافیت
خبره ۷	علوم داده و هوش مصنوعی	۸ سال فعالیت در حوزه الگوریتم‌ها	حکمرانی الگوریتمی
خبره ۸	مدیریت دولتی-خط‌مشی‌گذاری	۱۲ سال فعالیت حرفه‌ای در نهادهای اجتماعی	اخلاق فناوری، تحلیل سیاست
خبره ۹	مدیریت دولتی-خط‌مشی‌گذاری	۲۲ سال سابقه مدیریتی	سیاست‌گذاری اجتماعی
خبره ۱۰	مدیریت فناوری اطلاعات	۹ سال تجربه در حوزه دیجیتال	تحول دیجیتال و الگوریتم‌ها

در واقع، مرحله فراترکیب در این پژوهش صرفاً با هدف مرور ادبیات انجام نشد، بلکه به عنوان مبنای مفهومی و چارچوب اولیه برای مرحله کیفی دوم پژوهش به کار گرفته شد؛ به این معنا که ابعاد، مؤلفه‌ها و شاخص‌های استخراج‌شده از تحلیل نظام‌مند متون در حوزه حکمرانی الگوریتمی به عنوان چارچوب هدایت‌گر برای طراحی پرسشنامه نیمه‌ساختاریافته مصاحبه استفاده شدند. به بیان دیگر، نتایج فراترکیب نقش «لنز تحلیلی» را ایفا کرد که از طریق آن، مصاحبه‌ها با خبرگان هدایت و سامان‌دهی شدند تا مشخص شود این مؤلفه‌های نظری در بستر واقعی سازمان‌های خیریه چگونه می‌توانند زمینه‌ساز شکل‌گیری فرصت‌ها یا تهدیدهای مدیریت سایه شوند. بر این اساس، داده‌های حاصل از مصاحبه‌ها در فرایند تحلیل مضمون

کدگذاری شد و مضامین پایه، سازمان‌دهنده و فراگیر با اتکا به همان چارچوب مفهومی اولیه استخراج شدند. به این ترتیب، مرحله تحلیل مضمون تهدیدها و فرصت‌های مدیریت سایه با در نظر گرفتن حکمرانی الگوریتمی نه به صورت مستقل، بلکه در امتداد و تکمیل نتایج مرحله فراترکیب عمل کرده و پیوند میان ادبیات نظری حکمرانی الگوریتمی و تجربه خبرگان در حوزه تصمیم‌گیری خیریه‌ها را برقرار کرده است. مشارکت‌کنندگان با توجه به مدل حکمرانی الگوریتمی استخراج‌شده، شواهد، تجربیات و دیدگاه‌های خود را درباره پیامدهای مدیریت سایه در بستر خیریه‌ها ارائه کردند. داده‌های حاصل از مصاحبه‌های عمیق و متمرکز، پس از ثبت و گردآوری، مراحل کدگذاری باز، محوری و انتخابی را طی کردند و مضامین اصلی مرتبط با فرصت‌ها و تهدیدها استخراج و دسته‌بندی شدند. برای اطمینان از پایایی و توافق میان کدگذاری‌ها، ضریب کاپای کوهن محاسبه شد که مقدار ۰/۷۹۶ نشان‌دهنده قابلیت اتکای زیاد در فرایند تحلیل است. در نهایت، تلفیق یافته‌های مرحله فراترکیب و تحلیل مضمون چارچوبی تحلیلی از نقش حکمرانی الگوریتمی در تصمیم‌گیری خیریه‌ها ارائه داد که در آن، هم‌زمان فرصت‌های نوظهور و تهدیدهای مدیریت سایه در این عرصه مورد توجه قرار گرفته‌اند.

۴- یافته‌های پژوهش

در مرحله نخست پژوهش که با هدف ارائه چارچوب فرصت‌ها و تهدیدهای حکمرانی الگوریتمی در سازمان‌ها انجام شد، از روش فراترکیب استفاده شد. در این مرحله، فرایند انتخاب و غربالگری منابع به صورت گام‌به‌گام و نظام‌مند مطابق **شکل (۱)** دنبال شد.



شکل ۱. چگونگی غربال مقالات به منظور فراترکیب مقالات مرتبط با حکمرانی الگوریتمی

Figure 1. The screening process of articles for the meta-synthesis of studies related to algorithmic governance

طبق **جدول (۲)**، مولفه‌های اصلی از طریق تحلیل اهمیت- عملکرد محتوای تحقیق از منابع زیر به صورت کدبندی برداشت شده است.

جدول ۲. توصیف منابع فراترکیب ارائه‌چارچوب فرصت‌ها و تهدیدهای حکمرانی الگوریتمی

Table 2. Description of the meta-synthesis sources for developing the framework of opportunities and threats of algorithmic governance

کد	نویسنده	سال	کد	نویسنده	سال	کد	نویسنده	سال
L1	Vaghasia et al.	2025	L21	Radanliev et al.	2024	L41	Ryan & Stahl	2021
L2	Peringa et al.	2025	L22	Singhal et al.	2024	L42	Deldjoo et al.	2021
L3	Wang et al.	2025	L23	El Mestari et al.	2024	L43	Tran & Oh	2021
L4	Piedra-Cascón et al.	2025	L24	Ortega-Bolaños et al.	2024	L44	Stepin et al.	2021
L5	Michelucci & Venturini	2025	L25	Sachan et al.	2024	L45	Fukuda-Parr & Gibbons	2021
L6	Daily et al.	2025	L26	Glavin et al.	2024	L46	Wieringa	2020
L7	Baron et al.	2025	L27	Mensah	2023	L47	Raji et al.	2020
L8	Al-Dulaimi & Mohammed	2025	L28	Pentyala	2023	L48	Bastopcu & Ulukus	2020
L9	Marmolejo-Ramos et al.	2025	L29	Pagano et al.	2023	L49	Xu et al.	2019
L10	Cappelli & Di Marzo Serugendo	2025	L30	Herm et al.	2023	L50	Lin et al.	2019
L11	Pokhidnia	2025	L31	Cooper & Marques-Silva	2023	L51	Sampson et al.	2019
L12	Moon & Ahn	2025	L32	Corso et al.	2023	L52	Martin	2019
L13	Loreti & Visani	2024	L33	Chung & Lee	2023	L53	Jobin et al.	2019
L14	Munch et al.	2024	L34	Ding et al.	2022	L54	Selbst & Barocas	2018
L15	Lim et al.	2024	L35	Izza et al.	2022	L55	Yazdani et al.	2018
L16	El Koshiry et al.	2024	L36	Kim & Routledge	2022	L56	Veale et al.	2018
L17	Shannaq et al.	2024	L37	De Bruijn et al.	2022	L57	Wachter et al.	2017
L18	Akinrinola et al.	2024	L38	Green	2022	L58	Clark et al.	2016
L19	Folorunso et al.	2024	L39	Van Giffen et al.	2022	L59	Arslanturk et al.	2016
L20	Farayola et al.	2024	L40	Langer & Landers	2021	L60	Diakopoulos	2014
P1	دهقانی و همکاران	۱۴۰۴	P3	ناصرمقدسی	۱۴۰۲	P5	ستارند	۱۴۰۲
P2	اقتدار و همکاران	۱۴۰۲	P4	زواری	۱۴۰۲	P6	خرمشاد	۱۳۹۲

شاخص‌های استخراج‌شده از فراترکیب مطالعات پیشین در دو دسته اصلی فرصت‌ها و تهدیدهای حکمرانی الگوریتمی طبقه‌بندی شدند. فرصت‌ها بیانگر پیامدهای مثبت و ظرفیت‌های بهبود عملکرد سازمانی ناشی از به‌کارگیری سیستم‌های الگوریتمی هستند، در حالی که تهدیدها به ریسک‌های اخلاقی، اجتماعی و مدیریتی مرتبط با استفاده از این سیستم‌ها اشاره دارند. بر این اساس، کدهای منابع استخراجی به صورت جدول (۳) نمایش داده شدند.

جدول ۳. چارچوب ابعاد، مؤلفه‌ها و شاخص‌های فرصت‌ها و تهدیدهای حکمرانی الگوریتمی مستخرج از فراترکیب

Table 2. Framework of dimensions, components, and indicators of opportunities and threats of algorithmic governance derived from the meta-synthesis

سازه	بعد	مؤلفه	شاخص	منبع
فرصت‌ها	حکمرانی و شفافیت الگوریتمی	شفافیت عملکرد	میزان ارائه توضیح درباره تصمیم‌های الگوریتم	L34, L44, L57, L35, L37, L30, L14, L46, L50, L7, L54
			وجود مستندات عملکردی مدل برای کارکنان	L34, L46, L51, L47, L54, P2, P4
			درصد فرایندهایی که علت تصمیم‌هایشان توضیح‌پذیر است	L34, L44, L31, L35, L30, L50, L7



سازه	بُعد	مؤلفه	شاخص	منبع
مسئولیت پذیری	پاسخ گویی و مسئولیت پذیری	تعیین مسئول رسمی برای خطاهای الگوریتم	وجود پروتکل پاسخ گویی در خطاها	L60, L52, L36, L46, L8, L47, L13, L54, P6
				L60, L52, L46, L8, L47, L54, P2, P4
				L60, L46, L8, L47, L38
نظارت انسانی	میزان دخالت انسان در تصمیم های کلیدی	میزان دخالت انسان در تصمیم های کلیدی	درصد تصمیم های خودکار که نیازمند تأیید انسانی هستند	L52, L37, L46, L7, L54, P6
				L57, L52, L37, L46, L7, L54
				L34, L46, L51, L47, L54
اخلاق و مسئولیت پذیری الگوریتمی	رعایت اصول اخلاقی	وجود کد اخلاقی برای توسعه/استفاده از الگوریتم	میزان تطابق با استانداردهای اخلاقی بین المللی	L53, L45, L22, L24, L10, L27, L11, L41, L18
				L53, L45, L22, L24, L10, L27, L11, L41
				L16, L15, L49, L48
کیفیت داده و مدل	کیفیت و یکپارچگی داده	میزان بهروزرسانی داده ها	استانداردهای امنیت و محرمانگی داده ها	L16, L19, L20, L20, P2, P4
				L16, L4, L33, L17, P5
				L16, L3, L4, L55
پایش و نگهداری مدل	فاصله زمانی میان آموزش مجدد مدل ها	پایداری عملکرد در زمان های مختلف	زمان پاسخ الگوریتم	L16, L49, L55, L48
				L16, L4, L55, P5
				L25, L9, L26, L40, P4
تأثیر بر کارکنان و فرهنگی سازمانی	تجربه و رفاه کارکنان	رضایت شغلی در محیط الگوریتمی	میزان اعتماد کارکنان به سیستم های الگوریتمی	L26, L40
				L24, L25, L40
				L24, L25, L40
کارایی، کارکرد و نتایج کسب و کار	بهبود عملکرد سازمان	کاهش هزینه های عملیاتی	رضایت کاربر از تصمیم های سیستم	L25, L9, L26, L40
				L25, L26, L40
				L25, L26, L40
کیفیت تصمیم گیری	مقایسه دقت تصمیم های انسانی در برابر تصمیم های الگوریتمی	کاهش زمان انجام کارها در فرایندهای الگوریتمی	افزایش بهره وری کارکنان	L25, L26, L40
				L25, L26, L40
				L25, L26, L40
ریسک های کیفیت داده و خطاهای عملکرد	نقص در کیفیت و یکپارچگی داده	نرخ داده های ناسازگار یا ناقص	سرعت اجرای تغییرات در الگوریتم ها	L25, L26, L40
				L24, L25, L40
				L24, L25, L40, P1
تهدیدها	مدل	خطاهای عملکرد	نرخ خطا در تصمیم های حیاتی	L24, L25, L40, P1
				L16, L58, L17, L59
				L16, L2, L32, L6

سازه	بُعد	مؤلفه	شاخص	منبع
		ضعف در پایش و نگهداری مدل	تعداد هشدارهای ناشی از انحراف داده	L1, L16, L4, L39, L55,
ریسک‌های اخلاقی و تبعیض الگوریتمی	سبب‌گیری و تبعیض الگوریتمی	تفاوت در رفتار الگوریتم برای گروه‌های مختلف کاربران یا کارکنان	تفاوت در رفتار الگوریتم برای گروه‌های مختلف کاربران یا کارکنان	L24, L29, L25, L42, L27, L41, L12
		تعداد موارد مستندشده تبعیض الگوریتمی	امتیازدهی ناعادلانه	L24, L29, L25, L42, L27, L41, L12, P6
		نقض اصول اخلاقی در کاربرد الگوریتم	تعداد موارد تخلف از الزامات اخلاقی	L24, L29, L25, L42, L27, L41, L12
		تهدیدهای حریم خصوصی کارکنان	سطح دسترسی سیستم به داده‌های شخصی	L21, L23, L28, L11, L56, P2
		پیامدهای منفی مدیریت الگوریتمی بر کارکنان	عدم رضایت کارکنان از نحوه جمع‌آوری داده‌ها	L9, L26, L40
			عدم مطابقت سیستم با قوانین حفاظت داده	L21, L23, L28, L11, L56, P2
		فشار روانی ناشی از نظارت الگوریتمی	میزان استرس ناشی از نظارت خودکار	L26, L40
		تنش در روابط کارکنان	نرخ شکایت‌های کارکنان از مدیریت الگوریتمی	L25, L9, L26, L40
		کار در محیط الگوریتمی	میزان احساس بی‌عدالتی در سیستم امتیازدهی و نظارت	L24, L29, L25, L42, L27, L41, L12
			افزایش نرخ ترک خدمت پس از اجرای الگوریتم	L26, L40

لازم است تأکید شود که چارچوب ارائه‌شده در **جدول (۳)** صرفاً یک مرور ادبیات نبود، بلکه به عنوان چارچوب هدایت‌گر مصاحبه‌های نیمه‌ساختاریافته با خبرگان عمل کرد؛ به این معنا که پرسش‌های پرسشنامه بر اساس همین ابعاد و مؤلفه‌ها طراحی شدند تا خبرگان ضمن ارزیابی این فرصت‌ها و تهدیدها در سازمان‌های خیریه، پیامدهای سطح بالاتری را شناسایی کنند که در ادبیات کمتر بررسی شده‌اند. بر اساس یافته‌های فراترکیب، مجموعه‌ای جامع از مؤلفه‌ها و شاخص‌های مرتبط با ابعاد مختلف حکمرانی الگوریتمی در بخش دولتی استخراج شد. این چارچوب مبنای طراحی پرسشنامه نیمه‌ساختاریافته تحلیل مضمون قرار گرفت. بر این اساس، با اتکا به درک دقیق فرصت‌ها و تهدیدهای حکمرانی الگوریتمی، از خلال مصاحبه‌های خبرگان، مضامین مرتبط با فرصت‌ها و تهدیدهای مدیریت سایه در سازمان‌های خیریه، در وضعیتی که این سازمان‌ها تحت تأثیر سازوکارهای حکمرانی الگوریتمی عمل می‌کنند، استخراج می‌شوند.

با اتکا به چارچوب مفهومی و مؤلفه‌های شناسایی‌شده از مرحله فراترکیب، گام بعدی پژوهش به بررسی چگونگی ظهور پدیده مدیریت سایه در بستر خاص سازمان‌های خیریه معطوف شد. در راستای تقویت اعتبار یافته‌های پژوهش، فرایند انتخاب خبرگان به صورت هدفمند غیراحتمالی و بر اساس معیارهای مشخص انجام شد. همه خبرگان دارای مدرک دکتری و عضو هیئت علمی دانشگاه بودند. معیارهای خبرگی شامل تخصص علمی مرتبط با حوزه‌های مدیریت، حکمرانی و فناوری؛ سابقه قابل ملاحظه اجرایی یا مدیریتی در نهادهای عمومی یا اجتماعی؛ فعالیت علمی و پژوهشی در موضوع‌های اخلاقی، سیاست‌گذاری یا

حکمرانی الگوریتمی؛ و توانایی ارائه تحلیل‌های عمیق و نقادانه بودند. بر این اساس، ۱۰ خبره واجد شرایط در پژوهش مشارکت کردند. بنابراین، در گام بعدی و با تکیه بر همین چارچوب مفهومی (جدول ۳)، داده‌های حاصل از مصاحبه‌های نیمه‌ساختاریافته با خبرگان به روش تحلیل مضمون کدگذاری شدند تا مشخص شود هر یک از فرصت‌ها و تهدیدهای حکمرانی الگوریتمی در عمل، چگونه و در قالب چه مضامین مرتبط با مدیریت سایه در سازمان‌های خیریه ظاهر می‌شوند.

در این مرحله، با به‌کارگیری روش تحلیل مضمون و مشارکت خبرگان یادشده، داده‌های کیفی گردآوری و تحلیل شدند تا فرصت‌های نهفته در حکمرانی الگوریتمی برای شکل‌گیری یا تقویت این نوع حکمرانی کشف و دسته‌بندی شوند. هدف ترسیم نقشه‌ای از امکانات و مجرای بود که از طریق آن، خیریه‌ها می‌توانند، گاه به صورت ناخواسته، نقش‌هایی فراتر از مأموریت سستی خود را ایفا کنند. یافته‌های این مرحله که در قالب مضامین فراگیر، سازمان‌دهنده و پایه ساختاردهی شده‌اند، در جدول (۴) ارائه می‌شوند.

جدول ۴: فرصت‌های ناشی از مدیریت سایه در تصمیم‌گیری خیریه‌ها در چارچوب حکمرانی الگوریتمی

Table 4. Opportunities arising from shadow management in charitable decision-making within the framework of algorithmic governance

مضمون فراگیر	مضمون سازمان‌دهنده	مضمون پایه	متخصصان پیشنهاددهنده
تعالی مرجعیت خیریه در	توانمندسازی در تصمیم‌گیری مستقل	ارائه خدمات حمایتی و نیکوکارانه مکمل در کنار نظام رفاه رسمی	خبره ۱، خبره ۴، خبره ۷
خدمت‌رسانی عمومی		شناسایی سریع نیازهای معیشتی و حمایتی گروه‌های آسیب‌پذیر	خبره ۲، خبره ۵
		تخصیص بهینه منابع مالی و کمک‌های مردمی بر اساس تحلیل داده‌محور و شواهد عینی	خبره ۱، خبره ۳، خبره ۶، خبره ۸
هم‌افزایی راهبردی با نهادهای حاکمیتی		ایجاد پل‌های ارتباطی مؤثر و شفاف با مدیران دولتی برای انتقال تجربیات میدانی	خبره ۱، خبره ۲، خبره ۴، خبره ۹، خبره ۱۰
		تبدیل حمایت مالی هدفمند به ابزاری برای پیشبرد برنامه‌های مشترک توسعه‌ای	خبره ۳، خبره ۵، خبره ۷، خبره ۹
		ارائه داده‌ها و تجربیات میدانی خیریه‌ها برای بهبود سیاست‌های حمایتی	خبره ۲، خبره ۶، خبره ۸، خبره ۹، خبره ۱۰
پاسخ‌گویی سریع در شرایط اضطراری		پرکردن خلأهای تصمیم‌گیری و خدمات‌رسانی در شرایط بحرانی	خبره ۳، خبره ۴
توانمندسازی اجتماعی از طریق داده و فناوری	تولید و اشتراک دانش میدانی	ایجاد بانک‌های اطلاعاتی از نیازمندان و گروه‌های آسیب‌پذیر برای هدفمندکردن کمک‌های خیریه	خبره ۱، خبره ۶، خبره ۱۰
		تعریف و اولویت‌بندی مسائل حمایتی بر اساس داده‌های میدانی خیریه‌ها	خبره ۲، خبره ۴، خبره ۷
		اشتراک‌گذاری داده‌های حمایتی و تجربیات امداد رسانی خیریه‌ها با نهادهای سیاست‌گذار برای بهبود برنامه‌ریزی	خبره ۵، خبره ۶، خبره ۸، خبره ۹
الگوریتم‌های عدالت‌محور و شفاف		اولویت‌بندی عادلانه و علمی ذی‌نفعان با مدل‌های امتیازدهی شفاف	خبره ۱، خبره ۲

مضمون فراگیر	مضمون سازمان‌دهنده	مضمون پایه	متخصصان پیشنهاددهنده
		بهینه‌سازی تخصیص منابع با تصمیم‌گیری هوشمند و کاهش خطای انسانی	خبره ۳، خبره ۶، خبره ۷، خبره ۱۰
		افزایش اعتماد عمومی از طریق ارائه تصمیم‌های مبتنی بر شواهد و قابل راستی‌آزمایی	خبره ۱، خبره ۴، خبره ۹، خبره ۸
	ارتقای ظرفیت تحلیل و برنامه‌ریزی ملی	کمک به بهبود تحلیل مسائل فقر و آسیب اجتماعی از طریق داده‌ها و تجربیات خیریه‌ها	خبره ۳، خبره ۲، خبره ۵، خبره ۷
هدایت‌گری فرهنگی و اخلاقی سازنده	شکل‌دهی به افکار عمومی مسئولانه	آگاهی‌بخشی عمومی درباره نیازهای واقعی نیازمندان و اولویت‌های حمایتی	خبره ۵، خبره ۸
		افزایش حساسیت اجتماعی نسبت به فقر و مسائل معیشتی گروه‌های آسیب‌پذیر	خبره ۶، خبره ۹، خبره ۱۰
		ترویج فرهنگ نیکوکاری و مسئولیت اجتماعی در حل مسائل اجتماعی	خبره ۲، خبره ۴، خبره ۱
	تقویت سرمایه اجتماعی و اعتماد عمومی	کسب مشروعیت و اعتماد از سوی جامعه به واسطه عملکرد خیرخواهانه و مؤثر	خبره ۷، خبره ۳، خبره ۸، خبره ۱۰، خبره ۱
		دعوت به همکاری و مشارکت عمومی از طریق تبیین ارزش‌های مشترک	خبره ۸، خبره ۴
	ترویج مسئولیت‌پذیری و تعالی رفتاری	توانمندسازی دریافت‌کنندگان کمک از طریق برنامه‌های توسعه‌ای و مهارت‌آموزی	خبره ۶، خبره ۹، خبره ۵، خبره ۷
		ایجاد بسترهای حمایتی که به تدریج به خوداتکایی و رشد فردی منجر شود	خبره ۳، خبره ۵، خبره ۱، خبره ۹، خبره ۱۰، خبره ۲
چابکی نهادی و پاسخ‌گویی پیشگیرانه	ایجاد ساختارهای حقوقی نوآورانه	تعریف ساختارهای سازمانی منعطف برای فعالیت مؤثر خیریه‌ها در حوزه حمایت اجتماعی	خبره ۴، خبره ۱۰، خبره ۶
		شفاف‌سازی و تفکیک مسئولیت‌ها در پروژه‌های مشارکتی برای افزایش اعتماد	خبره ۸، خبره ۳، خبره ۸، خبره ۵
	تصمیم‌گیری هوشمند و شفاف	طراحی نظام تصمیم‌گیری چندلایه با تفکیک وظایف و افزایش دقت	خبره ۶، خبره ۹، خبره ۴
		افزایش قابلیت ردیابی و بازرینی تصمیم‌ها با بهره‌گیری از سامانه‌های هوشمند	خبره ۵، خبره ۱، خبره ۷، خبره ۱۰
		ایجاد فضای آزمون (پایلوت) برای سیاست‌های نوآورانه با کمترین هزینه اجتماعی	خبره ۱، خبره ۳، خبره ۲، خبره ۵، خبره ۱۰
کاهش نابرابری و تقویت فراگیری اجتماعی	تمرکز هوشمند بر گروه‌های هدف	توجه ویژه و مؤثر به گروه‌های دارای اولویت اجتماعی با برنامه‌ریزی هدفمند	خبره ۴، خبره ۸
		طراحی برنامه‌های حمایتی خیریه‌محور برای مسائل اجتماعی مغفول‌مانده	خبره ۶، خبره ۱۰، خبره ۹
	توانمندسازی پایدار جوامع هدف	گذار از کمک‌رسانی مقطعی به برنامه‌های جامع توانمندسازی	خبره ۷، خبره ۱، خبره ۶، خبره ۱۰
		افزایش آگاهی و ظرفیت مطالبه‌گری ساختاری ذی‌نفعان از طریق مشارکت‌دهی هدفمند	خبره ۲، خبره ۵، خبره ۹، خبره ۴، خبره ۸

همان‌طور که در **جدول (۴)** مشاهده می‌شود، تحلیل عمیق مضامین کیفی این پژوهش نشان داد حکمرانی الگوریتمی در سازمان‌های خیریه صرفاً یک ابزار فناورانه برای پردازش داده‌ها نیست، بلکه می‌تواند به عنوان بستری برای ارتقای کیفیت تصمیم‌گیری، افزایش شفافیت و بهبود کارایی تخصیص منابع در فعالیت‌های حمایتی عمل کند. خبرگان مشارکت‌کننده باور داشتند بهره‌گیری هوشمندانه و اخلاقی از داده و الگوریتم می‌تواند خیریه‌ها را از جایگاه سستی یک واسطه خیرخواه به بازیگرانی راهبردی و تأثیرگذار در عرصه حکمرانی اجتماعی تبدیل کند. این فرصت‌ها نه فقط در خدمت ارتقای کارایی و اثربخشی داخلی این سازمان‌ها قرار دارند، بلکه می‌توانند به تقویت جامعه مدنی، توانمندسازی دولت در تصمیم‌گیری، کاهش نابرابری‌ها و هدایت مسئولانه افکار عمومی منجر شوند. **جدول (۴)** این ظرفیت‌های تحول‌آفرین را در قالب مضامین فراگیر و سازمان‌دهنده ارائه می‌دهد. در مقابل فرصت‌های شناسایی‌شده، تحلیل مضامین کیفی به‌وضوح نشان داد همین سازوکارهای حکمرانی الگوریتمی، در صورت فقدان نظارت و چارچوب‌های اخلاقی مناسب، می‌تواند به بستری برای شکل‌گیری مدیریت سایه در فرایندهای تصمیم‌گیری تبدیل شوند. این تهدیدها ناظر بر خطراتی هستند که گسترش مدیریت سایه در بستر حکمرانی الگوریتمی می‌تواند برای سلامت نهادی سازمان‌های خیریه، شفافیت تصمیم‌گیری و اعتماد عمومی به این نهادها ایجاد کند؛ حتی اگر این سازوکارها در کوتاه‌مدت کارآمد به نظر برسند. **جدول (۵)** این تهدیدها را در قالب مضامین فراگیر، سازمان‌دهنده و پایه طبقه‌بندی و ارائه می‌کند.

جدول ۵. تهدیدهای ناشی از مدیریت سایه در تصمیم‌گیری خیریه‌ها در چارچوب حکمرانی الگوریتمی

Table 5. Threats arising from shadow management in charitable decision-making within the framework of algorithmic governance

مضمون فراگیر	مضمون	مضمون پایه	مصاحبه‌شوندگان
	سازمان‌دهنده	پیشنهاددهنده	
تضعیف حکمرانی رسمی از طریق	حذف تدریجی	حذف تدریجی تصمیم‌گیری جمعی	خبره ۲، خبره ۵، خبره ۷، خبره ۹
بازطراحی نامرئی فرایند	تصمیم‌گیری جمعی		
تصمیم‌گیری		حاشیه‌نشینی نهادهای ناظر راهبردی	خبره ۱، خبره ۴، خبره ۸
		تصاحب تدریجی جهت‌گیری‌های برنامه‌ای توسط سیستم‌های ارزیابی و گزارش‌دهی مالی و عملکردی	خبره ۳، خبره ۱۰
ایجاد ابهام عامدانه		چندلایه‌سازی فرایند تصمیم به گونه‌ای که مسئول نهایی قابل شناسایی نباشد	خبره ۵، خبره ۷
در مرجع تصمیم		ارجاع تصمیم‌های چالش‌برانگیز به گزارش‌های تحلیلی یا ارزیابی‌های کارشناسی بیرونی برای کاهش مسئولیت مدیریتی	خبره ۱، خبره ۶
تبدیل داده به ابزار نفوذ حکمرانی	دستکاری اولویت‌ها	انتخاب‌گزینی اطلاعات و آمارهای مربوط به حوزه‌های حمایتی جذاب برای اهداکنندگان و ایجاد سوگیری	خبره ۵، خبره ۷
از طریق داده		حذف یا کم‌رنگ‌سازی داده‌های مربوط به گروه‌های کم‌صدا	خبره ۴، خبره ۱۰
انحصار تفسیر داده		محدودکردن دسترسی به داشبوردها و گزارش‌های تحلیلی به حلقه‌ای خاص	خبره ۲، خبره ۴، خبره ۶
		استفاده از زبان فنی برای بستن مسیر پرسشگری سایر ذی‌نفعان	خبره ۲، خبره ۳، خبره ۶، خبره ۸، خبره ۹
چرخش خزنده مأموریت خیریه	بازتعریف موفقیت سازمانی	هم‌ارزسازی اثربخشی خیریه با شاخص‌های کمی	خبره ۱، خبره ۴، خبره ۷
		کوتاه‌مدت	



مضمون فراگیر	مضمون سازمان‌دهنده	مضمون پایه	مصاحبه‌شوندگان پیشنهاددهنده
		کنارگذاشتن پروژه‌های پرریسک اما تحول‌آفرین به دلیل امتیاز الگوریتمی پایین	خبره ۵، خبره ۱۰
	تغییر روایت ارزش‌ها	بازنمایی تصمیم‌های الگوریتمی به عنوان عقلانیت اخلاقی	خبره ۲، خبره ۳، خبره ۶، خبره ۱۰
		حاشیه‌نشینی ارزش‌هایی مانند کرامت مددجو در گزارش‌های رسمی	خبره ۲، خبره ۴، خبره ۶، خبره ۸
بازتولید نابرابری از مسیر تصمیم‌گیری	تقویت مزیت بازیگران قدرتمند	اولویت‌دهی سیستماتیک به مناطق یا گروه‌های دارای داده و شبکه قوی‌تر	خبره ۳، خبره ۹
		تسهیل دسترسی اهداکنندگان بزرگ به مسیرهای اثرگذاری تصمیم‌ها	خبره ۱، خبره ۷
	نادیده‌سازی صداها	حذف روایت‌های کیفی مددجویان به دلیل ناسازگاری با قالب داده‌ای	خبره ۱، خبره ۳، خبره ۹
فرسایش پاسخ‌گویی و اعتماد عمومی	عادی‌سازی تصمیم‌گیری غیرقابل پرسش	جاافتادن این تصور که سیستم اشتباه نمی‌کند در فرهنگ سازمانی کاهش تمایل مدیران به توضیح یا دفاع از تصمیم‌ها	خبره ۱، خبره ۵، خبره ۸، خبره ۶، خبره ۹
	فاصله‌گیری ذی‌نفعان	کاهش مشارکت داوطلبان در فرایند تصمیم‌سازی	خبره ۳، خبره ۴، خبره ۱۰
		تردید اهداکنندگان خرد نسبت به عدالت تخصیص منابع	خبره ۳، خبره ۶، خبره ۷
قفل‌شدگی حکمرانی در زیرساخت فناوریانه	وابستگی تصمیمی به پلتفرم	از دست رفتن امکان تصمیم‌گیری مستقل در نبود سیستم	خبره ۲، خبره ۵، خبره ۷، خبره ۹
		ناتوانی در توقف یا اصلاح تصمیم‌ها به دلیل پیچیدگی فنی	خبره ۱، خبره ۸
	محدودشدن اختیار راهبردی	تحمیل مسیرهای از پیش تعریف‌شده تصمیم‌گیری توسط نرم‌افزار	خبره ۳، خبره ۶
تهدید سلامت و استقلال سازمانی	تضعیف یادگیری و قضاوت سازمانی	تحلیل ظرفیت تحلیلی و قضاوت تخصصی کارکنان به دلیل اتکای بیش از حد به خروجی سیستم	خبره ۴، خبره ۶، خبره ۸، خبره ۹، خبره ۱۰
		نادیده‌گرفتن شهود و دانش زمینه‌ای کارکنان قدیمی در تصمیم‌گیری‌ها	خبره ۱، خبره ۲، خبره ۵
		ایجاد فرهنگ پیروی کورکورانه و اجتناب از مسئولیت‌پذیری فردی	خبره ۳، خبره ۷
مخاطرات امنیتی و عملیاتی	مخاطرات امنیتی	آسیب‌پذیری در برابر حملات سایبری که ممکن است کل عملیات خیریه را مختل کند	خبره ۲، خبره ۸
	مخاطرات عملیاتی	اتکای تصمیم‌های حیاتی به سیستم‌های ناپایدار یا با قابلیت اطمینان نامشخص	خبره ۵، خبره ۷
		نقض احتمالی حریم خصوصی داده‌های حساس اهداکنندگان و مددجویان	خبره ۱، خبره ۳، خبره ۴، خبره ۶، خبره ۹، خبره ۱۰

بر این اساس، مدل نهایی تهدیدها و فرصت‌ها به صورت **شکل (۲)** نمایش داده شد.

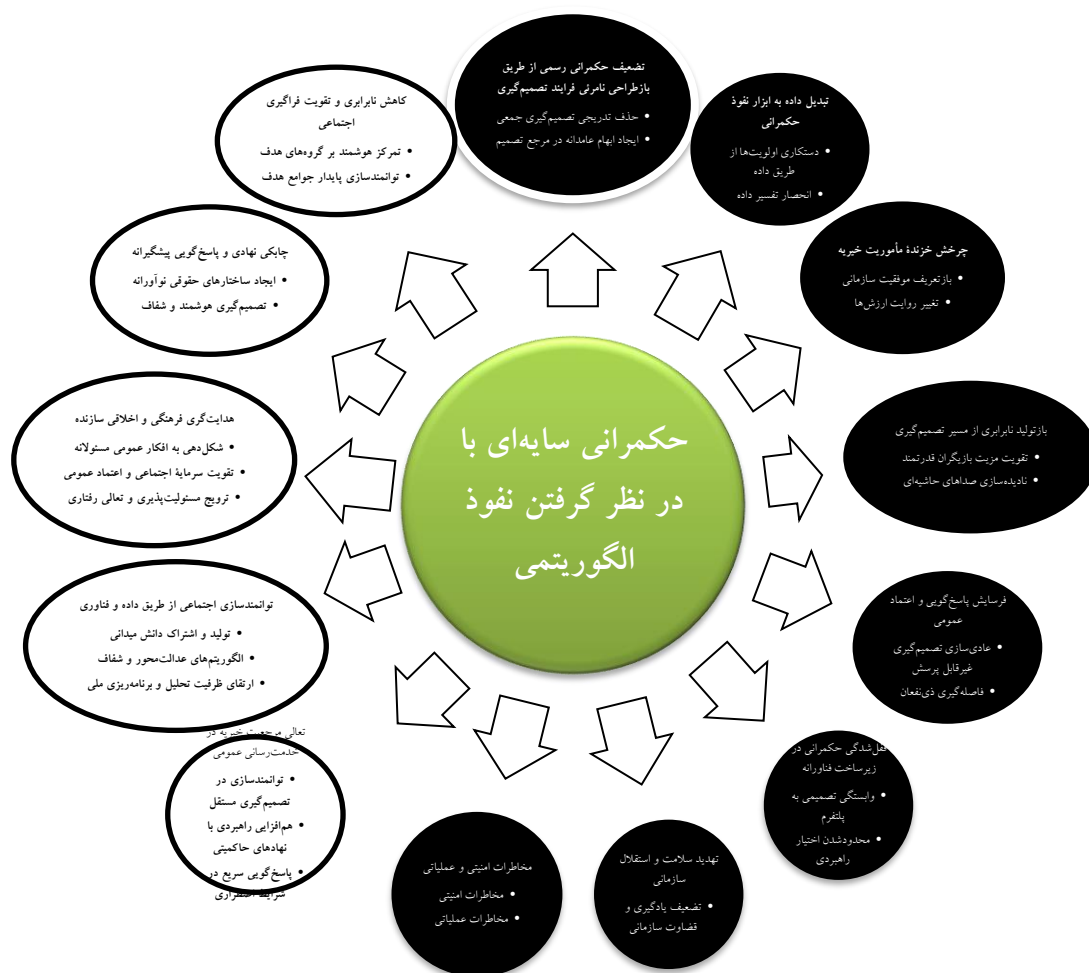
جدول ۶. وضعیت شاخص کاپا و نتایج آماره ضریب توافق کاپای کوهن چالش‌ها و فرصت‌های ناشی از مدیریت سایه با در نظر گرفتن حکمرانی الگوریتمی



Table 6. Kappa index status and the results of Cohen's Kappa agreement coefficient for the challenges and opportunities Arising from shadow management considering algorithmic governance

وضعیت توافق	مقدار عددی شاخص کاپا	نتایج آماره (ضریب توافق کاپای کوهن)
ارزش	کمتر از ۰	۰/۷۹۶
تعداد نمونه‌ها	۰/۲ - ۰	۵۹ (کد باز)
معنیاداری	۰/۲۱ - ۰/۴	۰/۰۰۰۱
	۰/۴۱ - ۰/۶	
	۰/۶۱ - ۰/۸	
	۰/۸۱ - ۱	

با توجه به نتایج جدول (۶)، مقدار ضریب توافق کاپای کوهن برابر ۰/۷۹۶ به دست آمد که بر اساس معیارهای تفسیر، نشان‌دهنده سطح توافق بالا و پایداری قابل قبول کدگذاری‌هاست. همان‌طور که در جدول (۵) مشهود است.

**شکل ۲.** مدل نهایی تهدیدها و فرصت‌های ناشی از مدیریت سایه با در نظر گرفتن حکمرانی الگوریتمی**Figure 2.** Final model of threats and opportunities arising from shadow management considering algorithmic governance

تهدیدهای ناشی از مدیریت سایه از سطح فرایندهای تصمیم‌گیری شروع می‌شوند و تا عمیق‌ترین لایه‌های مأموریت،

ارزش‌ها، اعتماد عمومی و استقلال سازمانی نفوذ می‌کنند. این یافته‌ها تصویری هشداردهنده ترسیم می‌کنند که در آن، کاربست الگوریتم‌ها بدون ملاحظات حکمرانی مناسب نه فقط ممکن است به تضعیف نهادهای رسمی و پاسخ‌گو بینجامد، بلکه در درون خود خیریه نیز به فرسایش سرمایه‌های انسانی، اخلاقی و اجتماعی آن منجر می‌شود. این مجموعه تهدیدها مؤید آن است که مدیریت سایه، با وجود دارا بودن پتانسیل‌های عمل‌گری سریع، در بلندمدت می‌تواند پایداری و مشروعیت ذاتی بخش خیریه را به مخاطره بیندازد.

۵- بحث و نتیجه‌گیری

پژوهش حاضر با هدف واکاوی نقش حکمرانی الگوریتمی در تصمیم‌گیری سازمان‌های خیریه و تبیین فرصت‌ها و تهدیدهای پدیده مدیریت سایه، با اتخاذ رویکردی ترکیبی و اکتشافی در دو مرحله انجام شد. در مرحله نخست، با استفاده از روش فراترکیب و غربال نظام‌مند ۶۶ منبع علمی معتبر، ابعاد و مؤلفه‌های حکمرانی الگوریتمی در سازمان‌های خیریه در قالب فرصت‌ها (شامل حکمرانی و شفافیت الگوریتمی؛ عدالت، اخلاق و کاهش سوگیری؛ کیفیت داده و مدل؛ تأثیر بر کارکنان و فرهنگ سازمانی؛ کارایی، کارکرد و نتایج کسب‌وکار) و تهدیدها (شامل ریسک‌های کیفیت داده و مدل؛ ریسک‌های اخلاقی و تبعیض الگوریتمی؛ پیامدهای منفی مدیریت الگوریتمی بر کارکنان) دسته‌بندی شدند. این چارچوب نظری جامع زیرساخت لازم برای ورود به مرحله دوم پژوهش را فراهم آورد و با علم به این چارچوب، پرسشنامه نیمه‌ساختاریافته تحلیل مضمون برای استخراج فرصت‌ها و تهدیدهای مدیریت سایه، برآمده از نحوه کاربست حکمرانی الگوریتمی تدوین شد و با خبرگان مرتبط مصاحبه انجام شد. در مرحله بعد، با تدوین پرسشنامه‌ای برگرفته از نتایج مرحله نخست، با خبرگان مصاحبه باز انجام شد. در این مرحله، از طریق تحلیل مضمون و با مشارکت ۱۰ خبره از حوزه‌های مدیریت دولتی، فناوری اطلاعات و سیاست‌گذاری، داده‌های کیفی عمیقی گردآوری و تحلیل شدند. پایایی فرایند تحلیل با ضریب کاپای کوهن ۰/۷۹۶ تأیید شد. خروجی مرحله اول یافته‌ها استخراج دو دسته از مضامین مرتبط با فرصت‌هایی بود که حکمرانی الگوریتمی برای خیریه‌ها ایجاد می‌کند و تهدیدهایی که در صورت عدم نظارت، به پیدایش مدیریت سایه منجر می‌شوند. یافته‌های پژوهش نشان می‌دهد حکمرانی الگوریتمی در سازمان‌های خیریه پدیده‌ای ذاتاً دوگانه و پارادوکسیکال است. از یک سو، این پدیده می‌تواند با ایجاد ظرفیت‌هایی همچون تعالی مرجعیت خیریه در خدمت‌رسانی عمومی، توانمندسازی اجتماعی از طریق داده و فناوری، هدایت‌گری فرهنگی و اخلاقی سازنده، چابکی نهادی و کاهش نابرابری، این سازمان‌ها را از جایگاه سستی یک واسطه خیرخواه به بازیگرانی راهبردی و تأثیرگذار در عرصه حکمرانی اجتماعی تبدیل کند. از سوی دیگر، همین سازوکارها در صورت فقدان چارچوب‌های نظارتی، اخلاقی و حقوقی متناسب، به تهدیدهایی جدی برای حاکمیت، مأموریت و مشروعیت این سازمان‌ها تبدیل می‌شوند. تهدیدهایی مانند تضعیف حکمرانی رسمی، تبدیل داده به ابزار نفوذ، چرخش خزننده مأموریت، بازتولید نابرابری، فرسایش پاسخ‌گویی و اعتماد عمومی، قفل‌شدگی حکمرانی در زیرساخت فناوریانه و مخاطرات امنیتی و عملیاتی، تصویری هشداردهنده از آسیب‌پذیری‌های این حوزه ترسیم می‌کنند. آنچه مدیریت سایه را شکل می‌دهد، نه خود الگوریتم‌ها، بلکه نحوه کاربست آنها در فضایی است که ابهامات حقوقی، عدم تقارن دانشی، نابرابری در دسترسی به داده و فقدان پاسخ‌گویی شفاف، زمینه را برای گذار از کارآمدی عملیاتی به سمت شکل‌گیری سازوکارهای پنهان و غیررسمی تصمیم‌گیری فراهم می‌آورند. تبیین نهایی این پژوهش حاکی از آن است که ما با یک گذار پارادایمی در حکمرانی بخش سوم مواجه هستیم؛ گذاری که در آن، مشروعیت و قدرت عمل به‌تدریج از نهادهای رسمی و پاسخ‌گو به

سمت بازیگرانی منتقل می‌شود که از مزیت راهبردی داده و الگوریتم برخوردار هستند، حتی اگر مسئولیت‌پذیری مستقیم و دموکراتیک آنان به طور کامل تعریف نشده باشد. در این چشم‌انداز نوظهور، خیریه‌ها به واسطه دسترسی بی‌واسطه به داده‌های میدانی، توانایی تحلیل پیشرفته و انعطاف‌پذیری نهادی، می‌توانند به مراجعی اثرگذار در تعریف مسائل اجتماعی، اولویت‌بندی تخصیص منابع و حتی هدایت افکار عمومی تبدیل شوند. این تغییر ماهیت، اگرچه فرصت‌هایی بی‌سابقه برای کارآمدی و اثربخشی فراهم می‌آورد، در عین حال، ضمن گشودن افق‌هایی جدید برای ارتقای شفافیت، پاسخ‌گویی و عدالت توزیعی، می‌تواند زمینه‌ساز شکل‌گیری نگرانی‌هایی جدید درباره محدود شدن نظارت عمومی، تشدید نابرابری‌های داده‌ای و تضعیف سلامت دموکراتیک در عرصه عمومی نیز باشد. از این رو، ضرورت طراحی و استقرار یک حکمرانی الگوریتمی مسئولانه برای خیریه‌ها بیش از پیش آشکار می‌شود؛ نوعی حکمرانی که ضمن بهره‌گیری از ظرفیت‌های تحول‌آفرین داده و الگوریتم، بتواند از طریق مکانیسم‌های شفافیت‌بخشی، نظارت انسانی معنادار، تضمین عدالت و کاهش سوگیری و مشارکت‌دهی ذی‌نفعان، از بروز تهدیدهای مدیریت سایه جلوگیری کند و این سازمان‌ها را در مسیر ایفای نقش‌های نوین خود، به عنوان نهادهایی پاسخ‌گو، اخلاقی و توسعه‌بخش، هدایت کند.

جمع‌بندی یافته‌های پژوهش نشان می‌دهد حکمرانی الگوریتمی در سازمان‌های خیریه پدیده‌ای دوجبه‌ای است که می‌تواند مسیرهایی متفاوت را برای آینده این نهادها رقم بزند. از یک سو، ظرفیت‌های داده‌محور و قابلیت‌های تحلیلی الگوریتم‌ها امکان ارتقای کارایی، تقویت توان تصمیم‌سازی، افزایش عدالت در تخصیص منابع و ایفای نقش مؤثرتر در حل مسائل اجتماعی را فراهم می‌آورند و خیریه‌ها را از یک کنش‌گر حمایتی سنتی به بازیگری راهبردی در حکمرانی اجتماعی ارتقا می‌دهند. از سوی دیگر، همین سازوکارها اگر بدون ملاحظات حکمرانی و نظارت شفاف به کار گرفته شوند، می‌توانند به شکل‌گیری مدیریت سایه، تضعیف فرایندهای رسمی تصمیم‌گیری، بازتولید نابرابری و فرسایش پاسخ‌گویی و سرمایه اجتماعی بینجامند. بر این اساس، بهره‌گیری از الگوریتم‌ها در بخش خیریه نیازمند رویکردی متوازن است که ضمن استفاده از فرصت‌های تحول‌آفرین آن، مانع از گسترش پیامدهای پرخطر و پنهان در لایه‌های مأموریت، ارزش‌ها و اعتماد عمومی شود. با عنایت به یافته‌های پژوهش که نشان‌دهنده ظرفیت‌های تحول‌آفرین و در عین حال، آسیب‌پذیری‌های جدی حکمرانی الگوریتمی در خیریه‌هاست، نخستین گام برای اعتلای این بخش تدوین و استقرار نظام حکمرانی الگوریتمی مسئولانه در سطح ملی و سازمانی است. پیشنهاد می‌شود نهادهای سیاست‌گذار مانند سازمان اوقاف و امور خیریه، با همکاری مراکز علمی و فناوریانه، نسبت به تدوین منشور اخلاقی و استانداردهای شفاف برای طراحی و به‌کارگیری سامانه‌های الگوریتمی در خیریه‌ها اقدام کنند. این چارچوب باید شامل الزاماتی مانند قابلیت توضیح‌دهی و شفافیت عملکرد الگوریتم، ممیزی مستمر برای شناسایی و کاهش سوگیری‌ها، تعیین مسئول مشخص برای خطاهای سیستمی و تضمین امنیت و محرمانگی داده‌های اهداکنندگان و مددجویان باشد. ایجاد یک نهاد ناظر تخصصی با مشارکت نمایندگان خیریه‌ها، دولت، دانشگاه و جامعه مدنی می‌تواند ضمن پایش مستمر، زمینه را برای هم‌افزایی راهبردی و جلوگیری از حکمرانی موازی و غیرشفاف فراهم آورد.

در سطح سازمانی، به مدیران و متولیان خیریه‌ها توصیه می‌شود تا به موازات بهره‌گیری از فرصت‌های الگوریتمی برای افزایش کارایی و اثربخشی، نسبت به توانمندسازی سرمایه انسانی خود به طور جدی اقدام کنند. نتایج پژوهش نشان می‌دهد اتکای بیش از حد به سیستم‌های الگوریتمی می‌تواند به تحلیل ظرفیت قضاوت تخصصی، تضعیف یادگیری سازمانی و حاشیه‌نشینی دانش زمینه‌ای کارکنان کهنسال منجر شود. بنابراین، ضروری است برنامه‌های آموزشی مستمر برای ارتقای سواد الگوریتمی کارکنان و مدیران طراحی و اجرا شوند تا آنان بتوانند به عنوان ناظران آگاه و تصمیم‌گیران نهایی، خروجی

سیستم‌ها را با هوشمندی ارزیابی و در موارد ضروری اصلاح کنند. همچنین، بازتعریف فرایندهای تصمیم‌گیری به گونه‌ای که تعامل پویا و متوازن میان قضاوت انسانی و توصیه‌های سیستمی برقرار باشد، از جمله راهکارهای مؤثر برای حفظ و تقویت ظرفیت‌های تحلیلی و اخلاقی درون‌سازمانی است.

سومین محور پیشنهادی بر تقویت اعتماد عمومی و مشارکت فعال ذی‌نفعان از طریق توسعه سامانه‌های الگوریتمی شفاف، فراگیر و پاسخ‌گو تأکید دارد. یافته‌ها نشان می‌دهد اگرچه الگوریتم‌ها می‌توانند دقت، چابکی و عدالت تصمیم‌گیری خیریه‌ها را افزایش دهند، فقدان شفافیت و مشارکت می‌تواند زمینه‌ساز مدیریت سایه، انحصار داده و تضعیف اعتماد عمومی شود. از این رو، توصیه می‌شود خیریه‌ها با انتشار گزارش‌های دوره‌ای درباره نحوه تخصیص منابع و معیارهای اولویت‌بندی، فرایندهای تصمیم‌سازی خود را برای ذی‌نفعان قابل پیگیری کنند. همچنین، ایجاد سازوکارهای مشارکت‌جویانه مانند پنل‌های مشورتی و سامانه‌های بازخورد ضروری است تا روایت‌های کیفی و صدای گروه‌های کم‌برخوردار در تصمیم‌گیری وارد شوند. توجه به کاهش نابرابری داده‌ای و تضمین بازنمایی منصفانه گروه‌های حاشیه‌ای نیز از عناصر کلیدی این رویکرد است. چنین اقداماتی می‌توانند خطرات مدیریت سایه را کاهش دهند، مشروعیت و سرمایه اجتماعی خیریه‌ها را تقویت و گذار از کم‌کسانی مقطعی به توانمندسازی پایدار و عدالت‌محور را امکان‌پذیر کنند.

یافته‌های پژوهش حاضر در خصوص دوگانه فرصت-تهدید حکمرانی الگوریتمی و شکل‌گیری مدیریت سایه با بخشی قابل ملاحظه از ادبیات داخلی و خارجی هم‌راستاست. به ویژه، شناسایی خطر کاهش پاسخ‌گویی و بازطراحی نامرئی فرایند تصمیم‌گیری با تحلیل حقوقی باتلر (2025) درباره تعارض تفویض تصمیم به الگوریتم‌ها و تضعیف مسئولیت‌پذیری نهادی همسو است. همچنین، تأکید این پژوهش بر نفوذ تدریجی و موازی سازوکارهای تصمیم‌یار با یافته‌های خاماس و همکاران (2025) درباره رهبری سایه و سازوکارهای نفوذ غیررسمی انطباق دارد. در بُعد رفتاری و سازمانی نیز نتایج مربوط به انحصار داده و پنهان‌سازی دانش با مطالعه لیو و همکاران (2025) هم‌جهت است که نشان دادند مدیریت الگوریتمی می‌تواند از طریق سازوکارهای روان‌شناختی به تقویت قدرت‌های غیرشفاف منجر شود. از نظر ریسک‌های اخلاقی و تبعیض، یافته‌های حاضر با تأکید مورس^۱ و همکاران (2020) و نیز سلبست و بارکاس^۲ (2018) بر سوگیری و پیامدهای ناعادلانه تصمیم‌های الگوریتمی هم‌راستاست. در ادبیات داخلی نیز تأکید بر ضرورت سیاست‌گذاری اخلاقی و پرهیز از وابستگی فناورانه در آثار قاسمی (۱۴۰۰) و روشن و همکاران (۱۴۰۰) با بخش تهدیدهای شناسایی‌شده در این پژوهش، به ویژه در زمینه قفل‌شدگی فناورانه و فرسایش قضاوت انسانی، هم‌خوانی دارد. با این حال، نوآوری پژوهش حاضر در انتقال این مباحث از بستر صرفاً دولتی به زمینه خاص سازمان‌های خیریه و تبیین پیوند آن با پدیده مدیریت سایه نمایان می‌شود.

پژوهش حاضر با وجود اتخاذ رویکردی نظام‌مند و ترکیبی، با محدودیت‌هایی مواجه بوده است که می‌توانند مسیر پژوهش‌های آتی را روشن کنند. در این پژوهش، سعی شده است تا عمدتاً بر سطح تحلیل سازمانی تمرکز شود. بر این اساس، به نقش متغیرهای زمینه‌ای مانند بسترهای فرهنگی، حقوقی و فناورانه کشورها توجه محدودی شده است. پیشنهاد می‌شود پژوهش‌های آتی با رویکردهای کمی و آزمون مدل‌های معادلات ساختاری، روابط علی میان مؤلفه‌های حکمرانی الگوریتمی و پیامدهای مدیریت سایه را ارزیابی کنند. همچنین، انجام مطالعات تطبیقی میان‌کشوری و بررسی نقش نهادهای تنظیم‌گر در مهار تهدیدها و تقویت فرصت‌ها می‌تواند افق‌هایی جدید را بگشاید. علاوه بر این، یافته‌های پژوهش حاضر ضرورت بررسی دقیق‌تر در

¹ Morse

² Selbst & Barocas

سطح خرد و نیز تحلیل تعاملات روزمره کارکنان با الگوریتم‌ها را آشکار می‌کند؛ موضوعی که اگرچه در این مطالعه مطرح شد، واکاوی عمیق‌تر آن مستلزم انجام پژوهش‌های مردم‌نگارانه آتی است. در نهایت، پژوهش در بسترهای خاص مانند خیریه‌های کوچک و محلی یا خیریه‌های فعال در بحران‌ها می‌تواند به غنای ادبیات نظری و عملی این حوزه بیفزاید.

۶- منابع فارسی

- ابوذری، م. (۱۴۰۳). هوش مصنوعی ابزار قدرت نرم در حوزه سیاست عمومی. فصلنامه علمی مطالعات قدرت نرم، ۱۳ (۴)، ۸۱-۱۰۰. https://www.spba.ir/article_194152.html
- اقتدار، سامره، مسگرزاده، مریم و اپرناک، فاطمه سارا. (۱۴۰۲). هوش مصنوعی در مراقبت‌های پرستاری و مامایی: راه‌حل جدید یا چالش‌های اخلاقی جدید؟ (نامه به سردبیر). مجله دانشکده پرستاری و مامایی ارومیه، ۲۱ (۴). <http://unmf.umsu.ac.ir/article-1-4897-fa.html>
- خرمشاد، م. ب. (۱۳۹۲). نگرشی دیگر بر نسبت اخلاق و سیاست. مجله علوم سیاسی، ۶۲، ۱-۱۸. https://psq.bou.ac.ir/article_12442.html
- دهقانی، م.، دژکامه، ف.، دهمرده، م.، و خاشی ناصری، ر. (۱۴۰۴). هوش مصنوعی و نقش آن در تحول فناوری‌های نوین. بیست‌وچهارمین کنفرانس ملی پژوهش‌های کاربردی در علوم برق، کامپیوتر و مهندسی پزشکی، شیروان. <https://civilica.com/doc/2491173>
- روشن، س. ع.، یعقوبی، ن. م.، و مؤمنی، ا. ر. (۱۴۰۰). کاربرد هوش مصنوعی در بخش دولتی (مطالعه‌ای فراترکیب). فصلنامه انجمن مدیریت ایران، ۱۶، ۱۱۷-۱۴۵. https://journal.iams.ir/article_349.html?lang=eng
- زواری، س. ع. ح. (۱۴۰۲). هوش مصنوعی و گذار به عصر جدید سیاست خارجی. انتشارات اندیشکده روابط بین‌الملل. https://psq.bou.ac.ir/article_12442.html
- ستارند، ز. (۱۴۰۲). حوزه روان‌شناسی هوش مصنوعی: استفاده اخلاقی از هوش مصنوعی بر انگیزه و سلامت روان دانش‌آموزان. اولین همایش بین‌المللی پیشروان تعلیم و تربیت، خرم‌آباد. <https://civilica.com/doc/2345085>
- صاحب‌الداری، م.، شیرخدایی، م.، یحیی‌زاده فر، م.، عابدین، ب.، و موقر، م. (۱۴۰۴). طراحی الگوی برندسازی دیجیتال مؤسسات خیریه. مطالعات وقف و امور خیریه، ۳ (۲)، ۱۹۹-۲۳۰. <https://doi.org/10.22108/ecs.2025.146304.1142>
- عباس‌نیا، س. م.، و عبدی، م. (۱۴۰۴). تحلیل اثر ساختار تعهدات مالی دیجیتال بر ریسک نکول در جمع‌آوری کمک‌های خیریه مردمی برای توسعه زیرساخت گردشگری مذهبی: شواهدی از یک تجربه میدانی. مطالعات وقف و امور خیریه، ۳ (۲)، ۱۴۱-۱۷۰. <https://doi.org/10.22108/ecs.2025.146744.1149>
- قاسمی، م. ر. (۱۴۰۰). هوش مصنوعی و حکمرانی آینده. حوزه، ۳۸ (۱۲-۱۳)، ۱۶۶-۱۷۹. <https://sid.ir/paper/1049584/fa>
- کاظمی نجف‌آبادی، مصطفی. (۱۴۰۴). تحلیل اثر تعاملی دیجیتالی شدن خدمات خیریه و شمول مالی بر فقر چندبعدی در استان‌های ایران: روش گشتاورهای تعمیم‌یافته سیستمی (Panel System GMM). مطالعات وقف و امور خیریه، ۳ (۱)، ۲۰۶-۱۸۷. <https://doi.org/10.22108/ecs.2025.146178.1136>
- کیاکجوری، د. (۱۴۰۳). شناسایی ابعاد و مؤلفه‌های ساختار تعاون در ایران با رویکرد وقف و امور خیریه مبتنی بر پژوهش آمیخته. مطالعات وقف و امور خیریه، ۲ (۱)، ۱۸۵-۲۰۲. <https://doi.org/10.22108/ecs.2024.140454.1092>



کیاکجوری، د.، و چریانی زنجانی، ی. (۱۴۰۳). تأثیر تبلیغات رسانه و دین‌داری اسلامی بر مشارکت در خیرات و وقف با در نظر گرفتن نقش میانجی نگرش خیران در میان خیران مدرسه‌ساز. *مطالعات وقف و امور خیریه*، ۲(۲)، ۹۱-۱۱۰.

<https://doi.org/10.22108/ecs.2024.141064.1102>

ناصرمقدسی، ع. ر. (۱۴۰۲). هوش مصنوعی و سلامت سیاسی. *روزنامه شرق*، ۴۶۲۲، ۱۱-۱۲.

<https://www.magiran.com/article/4434651>

References

- Abbasnia, S. M., & Abdi, M. (2025). Analysis of the effect of digital financial commitment structures on default risk in collecting public charitable donations for the development of religious tourism infrastructure: Evidence from a field experiment. *Endowment and Charitable Affairs Studies*, 3(2), 141-170. <https://doi.org/10.22108/ecs.2025.146744.1149> [In Persian]
- Abuzari, M. (2024) Artificial intelligence as a soft power tool in the field of public policy. *Quarterly Journal of Soft Power Studies*, 13(4), 81-100. https://www.spba.ir/article_194152.html [In Persian]
- Aizenberg, E., Dennis, M. J., & van den Hoven, J. (2025). Examining the assumptions of AI hiring assessments and their impact on job seekers' autonomy over self-representation. *AI & Society*, 40(2), 919-927. <https://doi.org/10.1007/s00146-023-01783-1>
- Akinrinola, O., Okoye, C. C., Ofodile, O. C., & Ugochukwu, C. E. (2024). Navigating and reviewing ethical dilemmas in AI development: Strategies for transparency, fairness, and accountability. *GSC Advanced Research and Reviews*, 18(3), 050-058. <https://doi.org/10.30574/gscarr.2024.18.3.0088>
- Akter, S., Dwivedi, Y. K., Biswas, K., Michael, K., Bandara, R. J., & Sajib, S. (2021). Addressing algorithmic bias in AI-driven customer management. *Journal of Global Information Management (JGIM)*, 29(6), 1-27. <https://doi.org/10.4018/JGIM.20211101.0a3>
- Al-Dulaimi, A. O. M., & Mohammed, M. A. A. W. (2025). Legal responsibility for errors caused by artificial intelligence (AI) in the public sector. *International Journal of Law and Management*. <https://doi.org/10.1108/IJLMA-08-2024-0295>
- Aragani, V. M., Anumolu, V. R., & Selvakumar, P. (2025). Democratization in the age of algorithms: Navigating opportunities and challenges. *Democracy and Democratization in the Age of AI*, 39-56. <https://www.igi-global.com/chapter/democratization-in-the-age-of-algorithms/370854>.
- Arslanturk, S., Siadat, M. R., Ogunyemi, T., Killinger, K., & Diokno, A. (2016). Analysis of incomplete and inconsistent clinical survey data. *Knowledge and Information Systems*, 46(3), 731-750. <https://doi.org/10.1007/s10115-015-0850-7>
- Barnes, E. (2025). Shadow accountability under hybrid algorithmic-human management: Consequences for psychological safety, autonomy, and middle management. *Advances in Social Sciences and Management*, 3(05), 148-162. <https://doi.org/10.63002/assm.305.1115>
- Baron, S., Latham, A. J., & Varga, S. (2025). Explainable AI and stakes in medicine: A user study. *Artificial Intelligence*, 340, 104282. <https://doi.org/10.1016/j.artint.2025.104282>
- Bastopcu, M., & Ulukus, S. (2020). Information freshness in cache updating systems. *IEEE Transactions on Wireless Communications*, 20(3), 1861-1874. <https://doi.org/10.1109/TWC.2020.3037144>
- Benlian, A., Wiener, M., Cram, W., Krasnova, H., Maedche, A., Möhlmann, M., ..., & Remus, U. (2022). Algorithmic management: Bright and dark sides, practical implications, and research opportunities. *Business & Information Systems Engineering*, 64(6), 825-839. <https://doi.org/10.1007/s12599-022-00764-w>
- Butler, O. (2025). Algorithmic decision-making, delegation and the modern machinery of government. *Oxford Journal of Legal Studies*, 45, 727-752. <https://doi.org/10.1093/ojls/gqaf018>
- Cappelli, M. A., & Di Marzo Serugendo, G. (2025). A semi-automated software model to support AI



- ethics compliance assessment of an AI system guided by ethical principles of AI. *AI and Ethics*, 5(2), 1357-1380. <https://doi.org/10.1007/s43681-024-00480-z>
- Carrillo, M. R. (2020). Artificial intelligence: From ethics to law. *Telecommunications Policy*, 44(6), 101937. <https://doi.org/10.1016/j.telpol.2020.101937>
- Chung, J., & Lee, K. (2023). Credit card fraud detection: An improved strategy for high recall using KNN, LDA, and linear regression. *Sensors*, 23, 7788. <https://doi.org/10.3390/s23187788>
- Clark, P. G., Gao, C., & Grzymala-Busse, J. W. (2016, September). A comparison of mining incomplete and inconsistent data. In *International Conference on Information and Software Technologies* (pp. 414-425). Cham: Springer International Publishing. <https://doi.org/10.5755/j01.itc.46.2.17330>
- Cooper, M. C., & Marques-Silva, J. (2023). Tractability of explaining classifier decisions. *Artificial Intelligence*, 316, 103841. <https://doi.org/10.1016/j.artint.2022.103841>
- Corso, A., Karamadian, D., Valentin, R., Cooper, M., & Kochenderfer, M. J. (2023). A holistic assessment of the reliability of machine learning systems. *arXiv preprint:2307.10586*. <https://doi.org/10.48550/arXiv.2307.10586>
- Daily, J. A., Dalby, S., & Greiten, L. (2025). Cognitive biases in high-stakes decision-making: implications for joint pediatric cardiology and cardiothoracic surgery conference. *Pediatric Cardiology*, 46(3), 536-543. <https://doi.org/10.1007/s00246-024-03462-4>
- De Bruijn, H., Warnier, M., & Janssen, M. (2022). The perils and pitfalls of explainable AI: Strategies for explaining algorithmic decision-making. *Government Information Quarterly*, 39(2), 101666. <https://doi.org/10.1016/j.giq.2021.101666>
- De Cremer, D., & McGuire, J. (2022). Human–algorithm collaboration works best if humans lead (because it is fair!). *Social Justice Research*, 35(1), 33-55. <https://doi.org/10.1007/s11211-021-00382-z>
- Dehghani, M., Dejkameh, F., Dehmardeh, M., & Khashi Naseri, R. (2025). *Artificial intelligence and its role in the evolution of new technologies*. Twenty-Fourth National Conference on Applied Research in Electrical, Computer and Medical Engineering Sciences, Shirvan. <https://civilica.com/doc/2491173> [In Persian]
- Deldjoo, Y., Anelli, V. W., Zamani, H., Bellogin, A., & Di Noia, T. (2021). A flexible framework for evaluating user and item fairness in recommender systems. *User Modeling & User-Adapted Interaction*, 31(3). <https://doi.org/10.1007/s11257-020-09285-1>
- Diakopoulos, N. (2014). *Algorithmic accountability reporting: On the investigation of black boxes*. Columbia University Libraries. <https://doi.org/10.7916/D8ZK5TW2>
- Ding, W., Abdel-Basset, M., Hawash, H., & Ali, A. M. (2022). Explainability of artificial intelligence methods, applications and challenges: A comprehensive survey. *Information Sciences*, 615, 238-292. <https://doi.org/10.1016/j.ins.2022.10.013>
- Eghtedar, S., Mesgarzadeh, M., & Aparnak, F. (2023) Artificial intelligence in nursing and midwifery care: a new solution or new ethical challenges?. *Nursing and Midwifery Journal*, 21(4), 272-276. <http://dx.doi.org/10.61186/unmf.21.4.272> [In Persian]
- El Arab, R. A., Abdulaziz, O., & Sagbakken, M. (2025). Economic, ethical, and regulatory dimensions of artificial intelligence in healthcare: An integrative review. *Frontiers in Public Health*, 13, 1617138. <http://doi.org/10.3389/fpubh.2025.1617138>
- El Koshiry, A. M., Eliwa, E. H. I., Abd El-Hafeez, T., & Khairy, M. (2024). Detecting cyberbullying using deep learning techniques: a pre-trained glove and focal loss technique. *PeerJ Computer Science*, 10, e1961. <https://doi.org/10.7717/peerj-cs.1961>
- El Mestari, S. Z., Lenzini, G., & Demirci, H. (2024). Preserving data privacy in machine learning systems. *Computers & Security*, 137, 103605. <https://doi.org/10.1016/j.cose.2023.103605>
- Farayola, O. A., Olorunfemi, O. L., & Shoetan, P. O. (2024). Data privacy and security in it: A review



- of techniques and challenges. *Computer Science & IT Research Journal*, 5(3), 606-615. <https://doi.org/10.51594/csitrj.v5i3.909>
- Folorunso, A., Mohammed, V., Wada, I., & Samuel, B. (2024). The impact of ISO security standards on enhancing cybersecurity posture in organizations. *World Journal of Advanced Research and Reviews*, 24(1), 2582-2595. <https://doi.org/10.1016/j.cose.2023.103605>
- Fukuda-Parr, S., & Gibbons, E. (2021). Emerging consensus on 'ethical AI': Human rights critique of stakeholder guidelines. *Global Policy*, 12, 32-44. <https://doi.org/10.1111/1758-5899.12965>
- Furendal, M. (Ed.) (2024). Political theory of the digital age: Where artificial intelligence might take us by Mathias Risse. *Political Science Quarterly*, 139(1), 154-155. <https://doi.org/10.1093/psquar/qqad140>
- Ghasemi, M. R. (2021). Artificial intelligence and future governance. *Hawza*, 38(12-13), 166-179. <https://sid.ir/paper/1049584/fa> [In Persian]
- Glavin, P., Bierman, A., & Schieman, S. (2024). Private eyes, they see your every move: workplace surveillance and worker well-being. *Social Currents*, 11(4), 327-345. <https://doi.org/10.1177/23294965241228874>
- Green, B. (2022). The flaws of policies requiring human oversight of government algorithms. *Computer Law & Security Review*, 45, 105681. <https://doi.org/10.1016/j.clsr.2022.105681>
- Gupta, A., Kozłowska, I., & Than, N. (2022). The golden circle: Creating socio-technical alignment in content moderation. *arXiv preprint arXiv:2202.13500*. <https://doi.org/10.48550/arXiv.2202.13500>
- Herm, L. V., Heinrich, K., Wanner, J., & Janiesch, C. (2023). Stop ordering machine learning algorithms by their explainability! A user-centered investigation of performance and explainability. *International Journal of Information Management*, 69, 102538. <https://doi.org/10.1016/j.ijinfomgt.2022.102538>
- Hernández-Lugo, M. D. L. C. (2024). Artificial Intelligence as a tool for analysis in social sciences: Methods and applications. *LatIA*, 2, 1-11. <https://doi.org/10.62486/latia202411>
- Hunkenschroer, A. L., & Luetge, C. (2022). Ethics of AI-enabled recruiting and selection: A review and research agenda. *Journal of Business Ethics*, 178(4), 977-1007. <https://doi.org/10.1007/s10551-022-05049-6>
- Izza, Y., Ignatiev, A., & Marques-Silva, J. (2022). On tackling explanation redundancy in decision trees. *Journal of Artificial Intelligence Research*, 75, 261-321. <https://doi.org/10.1613/jair.1.13575>
- Jarrahi, M., Newlands, G., Lee, M., Wolf, C., Kinder, E., & Sutherland, W. (2021). Algorithmic management in a work context. *Big Data & Society*, 8(2). <https://doi.org/10.1177/20539517211020332>
- Jensen, B. M., Whyte, C., & Cuomo, S. (2020). Algorithms at war: The promise, peril, and limits of artificial intelligence. *International Studies Review*, 22(3), 526-550. <https://doi.org/10.1093/isr/viz025>
- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1(9), 389-399. <https://doi.org/10.1038/s42256-019-0088-2>
- Kandasamy, C. A., Haridharshini, E., Prakash, N., Yoganathan, A., Kalyanasundaram, P., & Anitha, P. (2025, March). AI-powered medication reminder and tracker using reinforcement learning algorithm. In *2025 International Conference on Intelligent Computing and Control Systems (ICICCS)* (pp. 1075-1080). IEEE. <https://books.google.com/books?hl=en&lr=&id=myI6EQAAQBAJ&oi=fnd&pg=PR7&dq=Tursunbayeva+AI&ots=0wklZBbrLs&sig=CYx6wYAkqhgYltX8FT6BP9PMEUE#v=onepage&q=Tursunbayeva%20AI&f=false>
- Kapate, N., Dunne, M., Gottlieb, A. P., Mukherji, M., Suja, V. C., Prakash, S., ..., & Mitragotri, S. (2025). Polymer backpack-loaded tissue infiltrating monocytes for treating cancer. *Advanced Healthcare Materials*, 14(5), 2304144. <https://doi.org/10.1002/adhm.202304144>



- Kazemi Najafabadi, M. (2025). Analysis of the interactive effect of digitalization of charitable services and financial inclusion on multidimensional poverty in Iran's provinces: System generalized method of moments (panel system GMM). *Endowment and Charitable Affairs Studies*, 3(1), 187-206. <https://doi.org/10.22108/ecs.2025.146178.1136> [In Persian]
- Keegan, A., & Meijerink, J. (2025). Algorithmic management in organizations? From edge case to center stage. *Annual Review of Organizational Psychology and Organizational Behavior*, 12(1), 395-422. <https://doi.org/10.1146/annurev-orgpsych-110622-070928>
- Khamas, H., Dalvi, M., Abbas, Z., & Sadeghi, M. (2025). Developing and presenting a shadow leadership model in the ministry of water resources of Iraq. *International Journal of Innovation Management and Organizational Behavior*, 2(2), 12-21. <https://doi.org/10.61838/kman.ijimob.3770>
- Khorramshad, M. B. (2013). Another perspective on the relationship between ethics and politics. *Political Science Journal*, 62, 1-18. https://psq.bou.ac.ir/article_12442.html [In Persian]
- Kiakajuri, D. (2024). Identifying dimensions and components of cooperative structure in iran with a focus on endowment and charitable affairs based on mixed research. *Endowment and Charitable Affairs Studies*, 2(1), 185-202. <https://doi.org/10.22108/ecs.2024.140454.1092> [In Persian]
- Kiakajuri, D., & Zanjani, Y. C. (2024). The impact of media advertising and islamic religiosity on participation in charities and endowment with the mediating role of donors' attitudes among school-building donors. *Endowment and Charitable Affairs Studies*, 2(2), 91-110. <https://doi.org/10.22108/ecs.2024.141064.1102> [In Persian]
- Kim, T. W., & Routledge, B. R. (2022). Why a right to an explanation of algorithmic decision-making should exist: A trust-based approach. *Business Ethics Quarterly*, 32(1), 75-102. <https://doi.org/10.1017/beq.2021.3>
- Lakshmanan, M., Mala, G. A., Poorni, R., Ilamurugan, G., Sriramkumar, R., & Gnanavel, R. (2024, December). Blockchain for secure and efficient crowdfunding: An optimized particle swarm approach. In *2024 9th International Conference on Communication and Electronics Systems (ICCES)* (pp. 848-854). IEEE. <https://doi.org/10.1109/icces63552.2024.10860104>
- Langer, M., & Landers, R. N. (2021). The future of artificial intelligence at work: A review on effects of decision automation and augmentation on workers targeted by algorithms and third-party observers. *Computers in Human Behavior*, 123, 106878. <https://doi.org/10.1016/j.chb.2021.106878>
- Leicht-Deobald, U., Busch, T., Schank, C., Weibel, A., Schafheitle, S., Wildhaber, I., & Kasper, G. (2022). The challenges of algorithm-based HR decision-making for personal integrity. *Business and the Ethical Implications of Technology* (pp. 71-86). Cham: Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-18794-0_5
- Liao, T. H. (2023). The importance of human interactivity in artificial intelligence use in advertising: Development of a new scale. *International Journal of Advertising*, 1-36. <https://doi.org/10.1080/02650487.2025.2541505>
- Lim, W. S., Chen, Y. C., Chu, Y. H., Tu, C. H., & Chang, Y. H. (2024). iSAFE: Enabling evenness of data freshness in multi-priority networked intermittent systems. *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems*, 44(6), 2093-2104. <https://doi.org/10.1109/TCAD.2024.3522211>
- Lin, Z. Q., Shafiee, M. J., Bochkarev, S., Jules, M. S., Wang, X. Y., & Wong, A. (2019). Do explanations reflect decisions? A machine-centric strategy to quantify the performance of explainability algorithms. *arXiv preprint arXiv:1910.07387*. <https://doi.org/10.48550/arXiv.1910.07387>
- Liu, P., Yuan, L., & Jiang, Z. (2025). The dark side of algorithmic management: investigating how and when algorithmic management relates to employee knowledge hiding?. *Journal of Knowledge Management*, 29(2), 342-371. <https://doi.org/10.1108/JKM-04-2024-0507>
- Loreti, D., & Visani, G. (2024). Parallel approaches for a decision tree-based explainability algorithm.



- Future Generation Computer Systems*, 158, 308-322. <https://doi.org/10.1016/j.future.2024.04.044>
- Marmolejo-Ramos, F., Marrone, R., Korolkiewicz, M., Gabriel, F., Siemens, G., Joksimovic, S., ..., & Tejada, J. (2025). Factors influencing trust in algorithmic decision-making: An indirect scenario-based experiment. *Frontiers in Artificial Intelligence*, 7, 1465605. <https://doi.org/10.3389/frai.2024.1465605>
- Martin, K. (2019). Ethical implications and accountability of algorithms. *Journal of Business Ethics*, 160(4), 835-850. <https://doi.org/10.1007/s10551-018-3921-3>
- Mensah, G. B. (2023). Artificial intelligence and ethics: A comprehensive review of bias mitigation, transparency, and accountability in AI Systems. *Preprint*, November, 10(1), 1. <https://doi.org/10.13140/RG.2.2.23381.19685/1>
- Michelucci, U., & Venturini, F. (2025). New statistical framework for extreme error probability in high-stakes domains for reliable machine learning. *arXiv preprint arXiv:2503.24262*. <https://doi.org/10.48550/arXiv.2503.24262>
- Miller, J., Weyl, E., & Kanich, C. (2025). Fair decisions through plurality: results from a crowdfunding platform. In *Proceedings of the 5th ACM Conference on Equity and Access in Algorithms, Mechanisms, and Optimization* (pp. 107-118). <https://doi.org/10.1145/3757887.3763019>
- Moon, D., & Ahn, S. (2025). Metrics and algorithms for identifying and mitigating bias in AI design: A counterfactual fairness approach. *IEEE Access*. <https://doi.org/10.1109/ACCESS.2025.3345860>
- Morse, L., Teodorescu, M. H., Awwad, Y., & Kane, G. (2020). A framework for fairer machine learning in organizations. *arXiv preprint arXiv:2009.04661*. <https://doi.org/10.48550/arXiv.2009.04661>
- Munch, L. A., Bjerring, J. C., & Mainz, J. T. (2024). Algorithmic decision-making: The right to explanation and the significance of stakes. *Big Data & Society*, 11(1). <https://doi.org/10.1177/20539517231222872>
- Naser Moghaddasi, A. (2023). Artificial intelligence and political health. *Shargh Newspaper*, 4622, 11-12. <https://www.magiran.com/article/4434651> [In Persian]
- Noponen, N. (2019). Impact of artificial intelligence on management. *EJBO-Electronic Journal of Business Ethics and Organization Studies*, 24(2), 43-50. https://jyx.jyu.fi/jyx/Record/jyx_123456789_66617
- Ortega-Bolaños, R., Bernal-Salcedo, J., Germán Ortiz, M., Galeano Sarmiento, J., Ruz, G. A., & Tabares-Soto, R. (2024). Applying the ethics of AI: A systematic review of tools for developing and assessing AI-based systems. *Artificial Intelligence Review*, 57(5), 110. <https://doi.org/10.1007/s10462-024-10740-3>
- Pagano, T. P., Loureiro, R. B., Lisboa, F. V., Peixoto, R. M., Guimarães, G. A., Cruz, G. O., ..., & Nascimento, E. G. (2023). Bias and unfairness in machine learning models: A systematic review on datasets, tools, fairness metrics, and identification and mitigation methods. *Big Data and Cognitive Computing*, 7(1), 15. <https://doi.org/10.3390/bdcc7010015>
- Pentyala, D. K. (2023). Cloud-based solutions for AI-enhanced data governance and assurance. *International Journal of Social Trends*, 1(1), 154-178. <https://redcrevistas.com/index.php/Revista/article/view/84>
- Peringa, I. P., Cox, E. G. M., Wiersema, R., van der Horst, I. C., Meijer, R. R., & Koeze, J. (2025). Human judgment error in the intensive care unit: A perspective on bias and noise. *Critical Care*, 29(1), 86. <https://doi.org/10.1186/s13054-025-05315-9>
- Piedra-Cascón, W., Burgos-Artizzu, X. P., González-Martin, Ó., Oteo-Morilla, C., Pose-Rodríguez, J. M., & Gallas-Torreira, M. (2025). Evaluation of the accuracy (trueness, precision) and processing time of different 3-dimensional CAD software programs and algorithms for virtual cast alignment. *Journal of Dentistry*, 155, 105619. <https://doi.org/10.1016/j.jdent.2025.105619>



- Pokhidnia, B. (2025). Ethics in information management: Personal data protection. *Economic Scope*, (197), 212-216. <https://doi.org/10.30838/EP.197.212-216>
- Radanliev, P., Santos, O., Brandon-Jones, A., & Joinson, A. (2024). Ethics and responsible AI deployment. *Frontiers in Artificial Intelligence*, 7, 1377011. <https://doi.org/10.3389/frai.2024.1377011>
- Raji, I. D., Smart, A., White, R. N., Mitchell, M., Gebru, T., Hutchinson, B., ..., & Barnes, P. (2020, January). Closing the AI accountability gap: Defining an end-to-end framework for internal algorithmic auditing. In *Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency* (pp. 33-44). <https://doi.org/10.1145/3351095.3372873>
- Rakowski, R., & Kowaliková, P. (2024). The political and social contradictions of the human and online environment in the context of artificial intelligence applications. *Humanities and Social Sciences Communications*, 11(1), 1-8. <https://www.nature.com/articles/s41599-024-02725-y>
- Raphael, R. (2025). Medical crowdfunding with blockchain and machine learning: A secure and intelligent framework. *International Journal for Research in Applied Science and Engineering Technology*. <https://doi.org/10.22214/ijraset.2025.73061>
- Ross, J., Hibbert, L., & Moss, E. (2025). Shadow AI: Governance, risk, and organisational resilience. *2025 International Conference on Artificial Intelligence, Computer, Data Sciences and Applications (ACDSA)* (pp. 1-9). IEEE. <https://doi.org/10.1109/acdsa65407.2025.11166415>
- Rowshan, S. A., Yaqoubi, N., & Momeni, A. (2021). Application of artificial intelligence in the public sector (meta-combination study). *Iranian Journal of Management Sciences*, 16, 117-145. https://journal.iams.ir/article_349.html?lang=eng [In Persian]
- Ryan, M., & Stahl, B. C. (2021). Artificial intelligence ethics guidelines for developers and users: Clarifying their content and normative implications. *Journal of Information, Communication and Ethics in Society*, 19(1), 61-86. <https://doi.org/10.1108/JICES-12-2019-0138>
- Sachan, V. S., Katiyar, A., Somashekher, C., Chauhan, A. S., & Bhima, C. K. (2024). The role of artificial intelligence in HRM: opportunities, challenges, and ethical considerations. *Educational Administration: Theory and Practice*, 30(4), 7427-7435. <https://doi.org/10.53555/kuey.v30i4.2588>
- Saheb Aldari, M., Shirkhodaei, M., Yahyazadehfard, M., Abedin, B., & Mowgar, M. (2025). Designing a digital branding model for charitable institutions. *Endowment and Charitable Affairs Studies*, 3(2), 199-230. <https://doi.org/10.22108/ecs.2025.146304.1142> [In Persian]
- Salehi, F. (2025). Boardroom AI: the governance of AI-assisted corporate decision-making. *Global Journal of Economic and Finance Research*, 4, 232-235. <https://doi.org/10.55677/GJEFR/08-2025-Vol02E4>
- Sampson, C. J., Arnold, R., Bryan, S., Clarke, P., Ekins, S., Hatswell, A., ..., & Wrightson, T. (2019). Transparency in decision modelling: What, why, who and how?. *Pharmacoeconomics*, 37(11), 1355-1369. <https://doi.org/10.1007/s40273-019-00819-z>
- Sari, D. F. K. (2025). The influence of tiktok social media on adolescent social behavior in the digital era: A case study in Kendari City. *Sinergi International Journal of Communication Sciences*, 3(2), 127-140. <https://doi.org/10.61194/ijcs.v3i2.757>
- Satarand, Z. (2023). *The field of artificial intelligence psychology: Ethical use of artificial intelligence on students' motivation and mental health*. First International Conference on Pioneers of Education, Khorramabad. <https://civilica.com/doc/2345085> [In Persian]
- Schmitt, C. (2024). Romanticismo político. *Ediciones Olejnik*. https://books.google.de/books?hl=en&lr=&id=E74CEQAAQBAJ&oi=fnd&pg=PA7&dq=Schmitt+2024&ots=Egds4HQHXn&sig=FzWPcIKo4kpmnl9GE54tkf_-FJg&redir_esc=y#v=onepage&q&f=false
- Schultz, M. D., & Seele, P. (2023). Towards AI ethics' institutionalization: knowledge bridges from business ethics to advance organizational AI ethics. *AI and Ethics*, 3(1), 99-111. <https://doi.org/10.1007/s43681-022-00150-y>
- Selbst, A. D., & Barocas, S. (2018). The intuitive appeal of explainable machines. *Fordham Law*



- Review, 87(3), 1085–1139. <https://doi.org/10.2139/ssrn.3126971>
- Shannaq, B., Ali, O., Maqbali, S. A., & Al-Zeidi, A. (2024). Advancing user classification models: A comparative analysis of machine learning approaches to enhance faculty password policies at the University of Buraimi. *Journal of Infrastructure Policy and Development*, 8, 9311. <https://doi.org/10.24294/jipd9311>
- Shaw, A. N. (2025). Do the ends justify the means? Not according to donors: An analysis of manipulation in charitable marketing. In *Rethinking advertising: Ethics and effectiveness* (pp. 115-133). Cham: Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-86536-7_7
- Shi, L., Li, S., & Xing, J. (2025). Exploring chinese secondary EFL students' self-regulated learning and task engagement in AI-assisted classrooms: A latent growth curve modelling study. *European Journal of Education*, 60(4), e70241. <https://doi.org/10.1111/ejed.70241>
- Singhal, A., Nevéditsin, N., Tanveer, H., & Mago, V. (2024). Toward fairness, accountability, transparency, and ethics in AI for social media and health care: scoping review. *JMIR Medical Informatics*, 12(1), e50048. <https://doi.org/10.2196/50048>
- Stark, D., & Vanden Broeck, P. (2024). Principles of algorithmic management. *Organization Theory*, 5(2). <https://doi.org/10.1177/26317877241257213>
- Stepin, I., Alonso, J. M., Catala, A., & Pereira-Fariña, M. (2021). A survey of contrastive and counterfactual explanation generation methods for explainable artificial intelligence. *IEEE Access*, 9, 11974-12001. <https://doi.org/10.1109/ACCESS.2021.3051315>
- Tran, X. T., & Oh, H. (2021). A modified generic second order algorithm with fixed-time stability. *ISA Transactions*, 109, 72-80. <https://doi.org/10.1016/j.isatra.2020.10.021>
- Tursunbayeva, A. (2024). Augmenting human resource management with artificial intelligence: Towards an inclusive, sustainable, and responsible future. *Springer Nature*. <https://books.google.com/books?hl=en&lr=&id=myI6EQAAQBAJ&oi=fnd&pg=PR7&dq=Tursunbayeva+AI&ots=0wklZBbrLs&sig=CYx6wYAkqhgYltX8FT6BP9PMEUE>
- Vaghasia, P., Patel, R., Patel, D., Goswami, A., Patel, R., & Vaghasia, R. (2025, June). Improving data security and privacy in cloud-based data analysis: A results-driven approach. In *2025 International Conference on Computing Technologies (ICOCT)* (pp. 1-8). IEEE. <https://doi.org/10.1109/ICOCT64433.2025.11118763>
- Van Giffen, B., Herhausen, D., & Fahse, T. (2022). Overcoming the pitfalls and perils of algorithms: A classification of machine learning biases and mitigation methods. *Journal of Business Research*, 144, 93-106. <https://doi.org/10.1016/j.jbusres.2022.01.076>
- Vasuki, V. (2025). AI-driven medical fundraising verification system to detect and prevent fraudulent treatment requests. *International Scientific Journal of Engineering and Management*, 4(6), 1-9. <https://doi.org/10.55041/isjem04277>
- Veale, M., Binns, R., & Edwards, L. (2018). Algorithms that remember: model inversion attacks and data protection law. *Philosophical Transactions of the Royal Society: A Mathematical, Physical and Engineering Sciences*, 376, 20180083. <https://doi.org/10.1098/rsta.2018.0083>
- Wachter, S., Mittelstadt, B., & Russell, C. (2017). Counterfactual explanations without opening the black box: Automated decisions and the GDPR. *Harvard Journal of Law & Technology*, 31(2), 841. <https://doi.org/10.2139/ssrn.3063289>
- Wang, M., Zhu, S., Liu, X., Wen, S., & Mu, C. (2025). Finite-time input-to-state stability of neural networks with disturbances and prescribed performance. *IEEE Transactions on Systems, Man, and Cybernetics: Systems*, 55(5), 3742–3751. <https://doi.org/10.1109/TSMC.2025.3547499>
- Wieringa, M. (2020, January). What to account for when accounting for algorithms: a systematic literature review on algorithmic accountability. In *Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency* (pp. 1-18). Association for Computing Machinery. <https://doi.org/10.1145/3351095.3372833>



- Xu, C., Yang, H. H., Wang, X., & Quek, T. Q. (2019). Optimizing information freshness in computing-enabled IoT networks. *IEEE Internet of Things Journal*, 7(2), 971-985. <https://doi.org/10.1109/JIOT.2019.2947419>
- Yazdani, D., Nguyen, T. T., & Branke, J. (2018). Robust optimization over time by learning problem space characteristics. *IEEE Transactions on Evolutionary Computation*, 23(1), 143-155. <https://doi.org/10.1109/TEVC.2018.2845361>
- Yuanyuan, D. O. N. G. (2024). Ethical risks and countermeasures of artificial intelligence empowering ideological and political education. *Journal of Beijing University of Aeronautics and Astronautics (Social Sciences Edition)*, 37(1), 160–165. <https://dx.doi.org/10.13766/j.bhsk.1008-2204.2022.0955>
- Zavari, S. A. H. (2023). *Artificial intelligence and the transition to a new era of foreign policy*. International Relations Think Tank Publications. <https://doi.org/10.1111/ntwe.12343> [In Persian]
- Zhang, M., Cooke, F., Ahlstrom, D., & McNeil, N. (2025). The rise of algorithmic management and implications for work and organisations. *New Technology*, 40(3), 659–671. <https://doi.org/10.1111/ntwe.12343>
- Zhong, C., Cai, H., Fang, S., Xue, R., & Shan, Y. (2025). Does artificial intelligence reduce energy intensity in manufacturing? Evidence from country-level data. *Energy Economics*, 108784. <https://doi.org/10.1016/j.eneco.2025.108784>

