

Analysis of the Knowledge Structure, Thematic Evolution, and Emerging and Future Research Directions in User Interface Design for Explainable Artificial Intelligence (XAI) in Decision-Making Dashboards: A Scientometric Study

Sepehr
Noroozi Chakoli¹

Ehsan
Mousavi Khaneghah^{2*}

 1. B.Sc. Student in Computer Engineering, Faculty of Engineering, University of Guilan, Rasht, Iran; And Visiting B.Sc. Student, Faculty of Engineering, Shahed University, Tehran, Iran.
Email: Noroozi.sepehr@gmail.com

 2. Associate Professor, Department of Computer Engineering, Faculty of Engineering, Shahed University, Tehran, Iran, (Corresponding author).

Email: emousavi@shahed.ac.ir

Abstract

Purpose: The rapid diffusion of artificial intelligence (AI) in organizational decision-making environments has intensified concerns about transparency, interpretability, and user trust. As complex machine learning models increasingly support managerial and policy decisions, the need for Explainable Artificial Intelligence (XAI) has become essential to ensure that decision-makers understand and appropriately utilize algorithmic outputs. Decision dashboards have emerged as key interfaces through which AI-driven insights are delivered to managers and analysts. Despite their growing importance, the intellectual structure and thematic evolution of XAI within decision dashboard design have not yet been systematically examined. Therefore, this study aims to analyze the knowledge structure, conceptual foundations, and thematic development of research on Explainable Artificial Intelligence in decision dashboards. Employing scientometric techniques, the study identifies core research themes, emerging trends, and structural relationships shaping this rapidly evolving field. It also highlights how technical advancements in AI intersect with human-centered design, decision support systems, and visualization approaches in explainable decision dashboards.

Methodology: This study employs advanced bibliometric and network analysis techniques within a scientometric framework to map the intellectual landscape of explainable artificial intelligence (XAI) research in decision dashboards. Data were collected from the Scopus and Web of Science databases to ensure comprehensive coverage of peer-reviewed literature. Following the PRISMA protocol for systematic screening and refinement, an initial dataset of 942 records was identified. After applying inclusion criteria and limiting the corpus to original research articles, 251 records from Scopus and 112 from Web of Science were retained. Duplicate and irrelevant records were removed, resulting in a final dataset of 269 articles. Data standardization and descriptive bibliometric analyses were conducted using the Bibliometrix package and the Biblioshiny interface. These tools facilitated the examination of publication trends, leading journals, and author productivity patterns, including the application of Bradford's and Lotka's laws. To explore the conceptual structure of the field, keyword co-occurrence network analysis was performed using VOSviewer. Network, density, and overlay visualizations were generated to identify thematic clusters, intellectual linkages, and temporal changes in research topics. Addi-

Received:
13/02/2026

Final revision:
12/06/2026

Accepted:
20/06/2026

Early online access:
21/06/2026

Published:
01/10/2026



Sepehr
Noroozi Chakoli¹

Ehsan
Mousavi Khaneghah^{2*}

Received:
13/02/2026

Final revision:
12/06/2026

Accepted:
20/06/2026

Early online access:
21/06/2026

Published:
01/10/2026



tionally, thematic evolution mapping and logistic growth modeling were applied to examine the developmental trajectory of the field.

Findings: The findings indicate that research on Explainable Artificial Intelligence (XAI) in decision dashboards is evolving from a purely technical, algorithm-centered domain toward a multidimensional, interdisciplinary knowledge structure. Co-occurrence analysis reveals that the intellectual core of the field is organized around three central concepts: explainable AI, decision making, and deep learning. These themes form the primary hub of the conceptual network, suggesting that explainability is increasingly recognized as a fundamental component of intelligent decision support systems rather than merely a supplementary feature of machine learning models. The results also identify several interconnected thematic clusters, including human ethical considerations, technical methodological development, applied operational contexts, cognitive autonomous systems, and emerging specialized topics. This structure reflects the convergence of two complementary research paradigms: algorithmic performance and engineering efficiency on one hand, and human understanding, trust, and accountability on the other. Density visualization confirms that the highest concentration of research lies at the intersection of explainable AI, decision-making, and deep learning, underscoring the central role of deep learning technologies in the field. However, these technologies require integration with interpretability mechanisms and user interface design to be effective in decision dashboards. Overlay visualization further indicates a thematic transition over time. Earlier research primarily focused on classical machine learning techniques such as neural networks and support vector machines. In contrast, recent studies increasingly emphasize visual deep learning, visualization techniques, trust, and intelligent decision support systems. Another key finding is that concepts like trust and visualization have evolved from peripheral topics to integral components of explainable AI (XAI) systems, highlighting the growing importance of user comprehension and interaction in AI-supported decision-making processes. Logistic growth modeling also indicates that the field is still in a rapid expansion phase and has not yet reached scientific saturation.

Conclusion: The results demonstrate that the research landscape of Explainable Artificial Intelligence (XAI) in decision dashboards is shifting from a technology-centered paradigm toward a human- and decision-centered paradigm. Explainability is increasingly understood not only as a technical capability that clarifies algorithmic outputs but also as a cognitive bridge between complex AI systems and human decision-makers. The findings highlight that the effectiveness of AI-enabled decision dashboards depends not only on predictive accuracy but also on the systems' ability to communicate reasoning processes in an interpretable and trustworthy manner. Consistent with previous studies, the practical value of explainable AI emerges when explanations enhance transparency, foster user trust, and support informed decision-making in complex environments. By mapping the intellectual structure and thematic evolution of this interdisciplinary domain, the study provides insights into its conceptual foundations and future research directions. In particular, it emphasizes the importance of integrating deep learning technologies with explainability mechanisms, visualization techniques, and human-centered interface design to develop effective decision dashboards.

Keywords: Explainable Artificial Intelligence (XAI), Decision dashboards, Decision support systems, Scientometric analysis, Deep learning, Human-centered AI, Visualization, Decision making, Research fronts, Future research directions, Emerging trends

تحلیل ساختار دانشی، تکامل موضوعی و جهت‌گیری‌های آینده و نوظهور در پژوهش‌های طراحی رابط کاربری برای هوش مصنوعی توضیح‌پذیر (XAI) در داشبوردهای تصمیم‌گیری: مطالعه‌ای علم‌سنجی

سپهر نوروزی چاکلی^۱

۱. دانشجوی کارشناسی مهندسی کامپیوتر، دانشکده فنی و مهندسی، دانشگاه گیلان، رشت، ایران؛ و
دانشجوی کارشناسی میهمان، دانشکده فنی و مهندسی، دانشگاه شاهد، تهران، ایران.

Email: Noroozi.sepehr@gmail.com

احسان موسوی خانقاه^{۲*}

۲. دانشیار، گروه مهندسی کامپیوتر، دانشکده فنی و مهندسی، دانشگاه شاهد، تهران، ایران،
(نویسنده مسئول).

Email: emousavi@shahed.ac.ir

چکیده

هدف: گسترش کاربردهای هوش مصنوعی در محیط‌های تصمیم‌گیری سازمانی، ضرورت توجه به شفافیت، تفسیرپذیری و اعتماد کاربران به خروجی‌های الگوریتمی را افزایش داده است. در این میان، هوش مصنوعی توضیح‌پذیر به عنوان رویکردی برای قابل‌فهم ساختن منطق تصمیم‌گیری مدل‌های پیچیده مطرح شده و داشبوردهای تصمیم‌گیری به یکی از مهم‌ترین بسترهای ارائه این توضیحات تبدیل شده‌اند. با وجود اهمیت و نوظهور بودن این حوزه، ساختار دانشی و روندهای موضوعی پژوهش‌های مرتبط با هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری به صورت نظام‌مند بررسی نشده است. از این رو، هدف این پژوهش، تحلیل ساختار فکری و تکامل موضوعی این حوزه علمی است.

روش‌شناسی: این پژوهش با بهره‌گیری از رویکردی علم‌سنجی و با استفاده از فنون پیشرفته کتاب‌سنجی و تحلیل شبکه انجام شد. داده‌های پژوهش از پایگاه‌های Scopus و Web of science گردآوری شد. پس از اجرای فرایند غربالگری نظام‌مند بر اساس دستورالعمل پریزما (PRISMA)، مجموعه نهایی شامل ۲۶۹ مقاله پژوهشی، با استفاده از بسته Bibliometrix در محیط Biblioshiny مورد تحلیل‌های علم‌سنجی قرار گرفت. همچنین، ضمن به کارگیری نقشه تکامل موضوعی برای بررسی روند تحول پژوهش‌ها، برای تحلیل هم‌رخدادی واژگان و ترسیم نقشه‌های شبکه‌ای، چگالی و همپوشانی زمانی، از نرم‌افزار VOSviewer استفاده شد.

یافته‌ها: نتایج نشان داد که هسته مفهومی این حوزه در پیوند میان سه مفهوم «هوش مصنوعی توضیح‌پذیر»، «تصمیم‌گیری» و «یادگیری عمیق» شکل گرفته است. تحلیل خوشه‌بندی مفهومی، وجود چند خوشه اصلی شامل ابعاد انسانی-اخلاقی، فنی-روش‌شناسانه و کاربردی را نشان داد. همچنین، نتایج حاکی از آن است که مسیر تحول پژوهش‌ها از روش‌های کلاسیک یادگیری ماشین به سمت موضوعاتی مانند بصری‌سازی، اعتماد و سیستم‌های تصمیم‌یار هوشمند حرکت کرده است.

نتیجه‌گیری: یافته‌ها نشان می‌دهد که این حوزه در حال گذار از پارادایم «فناوری‌محور» به پارادایم «انسان‌محور و تصمیم‌محور» است. در این چارچوب، توضیح‌پذیری پلی میان مدل‌های هوش مصنوعی و فهم کاربران محسوب می‌شود و توسعه آینده داشبوردهای تصمیم‌گیری مستلزم تلفیق یادگیری عمیق، سازوکارهای توضیح‌پذیری و طراحی رابط کاربر انسان‌محور خواهد بود.

واژگان کلیدی: هوش مصنوعی توضیح‌پذیر، داشبوردهای تصمیم‌گیری، سیستم‌های پشتیبان تصمیم، تحلیل علم‌سنجی، یادگیری عمیق، بصری‌سازی، تصمیم‌گیری، جبهه‌های پژوهشی، روندهای نوظهور، جهت‌گیری‌های آینده پژوهش

صفحه ۲۶۸-۲۳۵

دریافت: ۱۴۰۴/۱۱/۲۴

بازنگری نهایی: ۱۴۰۵/۰۳/۲۲

پذیرش: ۱۴۰۵/۰۳/۳۰

زودآیند: ۱۴۰۵/۰۳/۳۱

انتشار: ۱۴۰۵/۰۷/۰۹



مقدمه و بیان مسئله

هوش مصنوعی در سال‌های اخیر به یکی از اثرگذارترین فناوری‌های عمومی منفعت‌محور تبدیل شده و در حوزه‌هایی مانند سلامت، آموزش، صنعت، مدیریت و تصمیم‌گیری سازمانی نقش فزاینده‌ای یافته است (Tveita & Hustad, 2025). گزارش‌های معتبر بین‌المللی نیز نشان می‌دهند که اثرگذاری این فناوری بر بهره‌وری، کیفیت کار و شیوه تصمیم‌سازی انسان‌ها روزبه‌روز پررنگ‌تر می‌شود. با این حال، گسترش کاربردهای هوش مصنوعی، به طور هم‌زمان با افزایش نگرانی‌ها درباره شفافیت، اعتماد، سوگیریو قابلیت فهم خروجی مدل‌ها همراه بوده است؛ به همین دلیل، ضرورت حرکت به‌سوی هوش مصنوعی توضیح‌پذیر^۱ بیش از پیش مورد توجه قرار گرفته است (Barredo et al., 2020).

در این میان، طراحی رابط کاربری برای هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، به‌عنوان یکی از مهم‌ترین مسیرهای تبدیل خروجی‌های پیچیده مدل‌های هوش مصنوعی به اطلاعات قابل‌استفاده برای انسان مطرح شده است. این حوزه بر آن است که توضیحات مدل، نه صرفاً از نظر فنی، بلکه از منظر کاربر، زمینه کاربرد، سطح دانش مخاطب و نیاز تصمیم‌گیرنده طراحی شوند. در واقع، یک داشبورد تصمیم‌گیری مبتنی بر هوش مصنوعی توضیح‌پذیر باید بتواند میان «دقت الگوریتمی» و «فهم انسانی» پیوند برقرار کند؛ به این معنا که توضیحات ارائه‌شده نه‌تنها معتبر باشند، بلکه برای کاربر نیز قابل‌درک، عملی و متناسب با وظیفه تصمیم‌گیری باشند (Miller, 2019). پژوهش‌های اخیر نیز نشان می‌دهند که موفقیت هوش مصنوعی توضیح‌پذیر به‌شدت به طراحی انسان‌محور، ارزیابی تجربه کاربر و یکپارچگی توضیحات با جریان واقعی کار وابسته است (Kim et al., 2024).

این حوزه را می‌توان یک حوزه نوظهور و در حال تثبیت دانست. شواهد کتاب‌سنجی نشان می‌دهد که ادبیات XAI رشد بسیار سریعی داشته و بخش عمده این رشد از سال‌های پایانی دهه ۲۰۱۰ م به بعد شکل گرفته است؛ به‌طوری‌که برخی مطالعات کتاب‌سنجی با بررسی هزاران مقاله هوش مصنوعی توضیح‌پذیر، جهش کمی این حوزه و شکل‌گیری خوشه‌های پژوهشی جدید را گزارش کرده‌اند (Alonso et al., 2018; Ammar et al., 2026; Chen et al., 2023; Rejeb et al., 2026; Russo & Vistocco, 2026). در کنار آن، مطالعات مروری دیگر نیز نشان می‌دهند که هنوز اجماع روشنی درباره استانداردهای طراحی، شیوه‌های ارزیابی و الگوهای مؤثر برای توضیح‌پذیری وجود ندارد (Ali et al., 2023)؛ برای مثال، در یک مرور نظام‌مند از مطالعات کاربرمحور هوش مصنوعی توضیح‌پذیر، تنها بخش محدودی از پژوهش‌ها از چارچوب‌های ارزیابی مشترک استفاده کرده‌اند که خود نشان‌دهنده پراکندگی و ناپختگی نسبی این حوزه است (Kim et al., 2024). همچنین پژوهش‌های جدید در زمینه تصمیم‌یارهای بالینی^۲ و داشبوردهای مبتنی بر هوش مصنوعی نشان داده‌اند که کاربران، به‌ویژه متخصصان، توضیحات کوتاه، بصری، ساده و متناسب با زمینه کار را ترجیح می‌دهند و در عین حال، بسیاری از نمایش‌های رایج مانند برخی نمودارهای تبیینی پیچیده را به‌راحتی قابل‌فهم نمی‌دانند (Fouad et al., 2026; Hunsicker et al., 2026).

در حوزه‌های مرتبط نیز ابعاد مطالعاتی متنوعی دیده می‌شود که نشان می‌دهد طراحی رابط کاربری برای هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری صرفاً یک مسئله فنی نیست، بلکه یک موضوع میان‌رشته‌ای است که به تعامل انسان و کامپیوتر، بصری‌سازی داده، بار شناختی، تجربه کاربری و زمینه کاربرد وابسته است. برای نمونه،

1 . Explainable Artificial Intelligence (XAI)

2 . Clinical Decision Aids

مطالعات اخیر درباره هوش مصنوعی توضیح‌پذیر در سیستم‌های پشتیبان تصمیم‌گیری بالینی بر نیاز به توضیحات عملی، قابل‌هضم و ادغام‌شده با روند کاری تأکید کرده‌اند (Rezaeian et al., 2026; Shiddik, 2026). در همین راستا، پژوهش‌های دیگر نیز به نقش عناصر بصری، تعامل‌پذیری، دسترس‌پذیری و ارزیابی کاربر در طراحی داشبوردهای تصمیم‌گیری مبتنی بر هوش مصنوعی پرداخته‌اند (Mahanta et al., 2026; Prasad et al., 2026). این تنوع موضوعی نشان می‌دهد که حوزه مورد نظر شما در حال شکل‌گیری به‌عنوان یک قلمرو دانشی مستقل است، اما هنوز نقشه روشن و منسجمی از ساختار درونی آن ارائه نشده است.

در چنین شرایطی، علم‌سنجی می‌تواند ابزار مناسبی برای آشکارسازی ساختار دانشی این حوزه باشد. علم‌سنجی با تحلیل تولیدات علمی، استنادها، هم‌رخدادی کلیدواژه‌ها، شبکه نویسندگان، کشورها، دانشگاه‌ها و مجلات، امکان ترسیم نقشه‌ای از رشد، تمرکزهای موضوعی، ساختارهای راهبردی دانش و روندهای تحول یک حوزه پژوهشی را فراهم می‌کند. به بیان دیگر، مطالعه ساختار دانشی می‌تواند به پژوهشگران نشان دهد که مرزهای دانشی حوزه تا به کجا پیش رفته است، کدام جریان‌های فکری و موضوعی در حال شکل‌دادن به آن هستند، بازیگران اصلی چه کسانی‌اند و شکاف‌های پژوهشی در کجا قرار دارند (Haghani, 2023; Ivancheva, 2008; Van Leeuwen, 2006). مطالعات علم‌سنجی نشان داده است که این حوزه می‌تواند برای شناسایی خوشه‌های پژوهشی، موضوعات داغ، کشورها و نویسندگان اثرگذار و نیز جهت‌گیری‌های آینده حوزه هوش مصنوعی توضیح‌پذیر نیز بسیار مفید و ضروری باشد. حتی تحلیل شناسایی نویسندگان و نشریات برجسته که از مهم‌ترین قوانین کلاسیک حوزه علم‌سنجی محسوب می‌شود، با ظهور ابزارهای نوین بیش از پیش مورد توجه پژوهشگران برای ارزیابی ساختار دانش قرار گرفته است (Bradford, 1985; Lotka, 1926; Lotka & Bradford, 1926; Torbati & Noroozi Chakoli, 2013). براین‌اساس، مسئله اصلی این پژوهش آن است که با توجه به اهمیت روزافزون هوش مصنوعی و جایگاه طراحی رابط کاربری برای حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، ساختار دانشی این حوزه چگونه شکل گرفته است؟ این که چه زیرشاخه‌ها و حوزه‌های مرتبط و راهبردی در این قلمرو وجود دارد، این حوزه نوظهور تاکنون بر چه زمینه‌های موضوعی و پژوهشی تمرکز داشته است و قوانین لوتکا و بردفورد چگونه بر روابط نویسندگان و مجلات این حوزه حاکم است، از مسائل اساسی مطرح در این پژوهش به شمار می‌رود. بی‌تردید، پاسخ به این پرسش‌ها می‌تواند تصویری روشن از وضعیت فعلی و مسیرهای آتی پژوهش در این حوزه ارائه دهد و به پژوهشگران کمک کند تا با شناخت بهتر از جبهه‌های پژوهش، تصمیم‌های دقیق‌تری برای ادامه پژوهش، همکاری علمی و انتخاب مسیرهای آینده اتخاذ کنند.

پرسش‌های پژوهش

- باتوجه‌به مطالب مطرح شده در بخش پیشین، این پژوهش در پی پاسخگویی به پرسش‌های زیر است:
۱. الگوی رشد و توزیع زمانی تولیدات علمی در حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری در سطح بین‌المللی چگونه است؟
 ۲. الگوی توزیع انتشارات در مجلات حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، چگونه از قانون برادفورد پیروی می‌کند و چه نوع ساختاری را در تمرکز و پراکندگی مجلات نشان می‌دهد؟
 ۳. الگوی بهره‌وری نویسندگان حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، چگونه سازمان یافته

تحلیل ساختار دانشی، تکامل موضوعی و جهت‌گیری‌های آینده و نوظهور در پژوهش‌های ...

است و بر اساس قانون لوتکا، رهبران علمی تا چه میزان جهت‌گیری و تمرکز بدنه دانش این حوزه را شکل می‌دهند؟

۴. ساختار فکری و راهبردی خوشه‌های اصلی دانش در حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری از چه الگویی تبعیت می‌کند و مؤلفه‌های متمایز و کلیدی هر خوشه کدام‌اند؟

۵. با رویکرد آینده‌پژوهانه، موضوعات هسته اصلی و نوظهور پژوهش در حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، کدام‌اند و روند تکامل ساختار دانشی آن‌ها در طول زمان چگونه قابل تفسیر است؟

چارچوب نظری

هوش مصنوعی توضیح‌پذیر (XAI) به‌عنوان یک حوزه پژوهشی مهم و راهبردی با هدف افزایش شفافیت، قابلیت اعتماد و تفسیرپذیری سیستم‌های هوشمند ظهور کرده است. در این زمینه، گراف‌های دانش^۱ با فراهم کردن یک چارچوب ساختاریافته برای نمایش و سازمان‌دهی روابط پیچیده میان موجودیت‌های داده، نقشی اساسی در بهبود درک مفهومی داده‌ها ایفا می‌کنند. از سال ۲۰۲۰ م. به بعد، قابلیت توضیح‌پذیری^۲ به‌عنوان یکی از عوامل کلیدی برای پذیرش سیستم‌های هوش مصنوعی، کاربرد گسترده‌ای یافته است. دلایل مجهز کردن سیستم‌های هوشمند به قابلیت توضیح‌پذیری، فقط به حقوق کاربران یا پذیرش فناوری محدود نمی‌شود؛ بلکه بهره‌گیری از این قابلیت برای طراحان و توسعه‌دهندگان نیز ضروری است تا بتوانند استحکام سیستم را افزایش دهند و امکان عیب‌یابی را برای جلوگیری از سوگیری، بی‌عدالتی و تبعیض فراهم کنند. همچنین، این ویژگی می‌تواند اعتماد کاربران را نسبت به اینکه چرا و چگونه تصمیم‌ها گرفته می‌شوند، افزایش دهد (Confalonieri et al., 2021).

باوجوداین، هنوز توافق روشنی درباره اینکه «توضیح» دقیقاً چیست یا یک «توضیح خوب» چه ویژگی‌هایی دارد وجود ندارد؛ بطوری که برداشت‌های مختلفی از توضیح در گونه‌های متفاوت سیستم‌های هوش مصنوعی و رشته‌های گوناگون بررسی شده‌اند. نخستین مفاهیم توضیح‌پذیری در هوش مصنوعی همراه با سیستم‌های خبره پس از اواسط دهه ۱۹۸۰ م. کمرنگ شدند (Buchanan & Shortliffe, 1984)، اما با موفقیت‌های اخیر در فناوری یادگیری ماشین دوباره مورد توجه قرار گرفتند (Guidotti et al., 2018).

آنجلوف (Angelov et al., 2021) در سال ۲۰۲۱ م. بر این باور بود که پژوهش‌های حوزه هوش مصنوعی توضیح‌پذیر هنوز عمدتاً به تحلیل حساسیت، انتشار لایه‌به‌لایه، اهمیت ویژگی‌ها و انتساب، توضیح‌های شبه‌محلی به‌وسیله LIME، توضیح‌های افزایشی شیلی مبتنی بر نظریه بازی (SHAP)، مکان‌یابی مبتنی بر گرادینان و Grad-CAM، یا مدل‌های جانشین محدود است. اما واقعیت این است که پژوهش‌های این حوزه از سال ۲۰۲۲ م. به بعد به کلی دگرگون و متحول شده و به عرصه‌های گوناگونی وارد شده و به ویژه در مسیر تأثیرگذاری بر داشبوردهای تصمیم‌گیری ورود پیدا کرده است (Akkem et al., 2026; Bornmann & Leydesdorff, 2014; Buñay, 2026; Guisñan et al., 2026; Nalela et al., 2026; Pereira & Maciel, 2026).

در چنین شرایطی است که علم‌سنجی می‌تواند برای شناسایی زمینه‌های پژوهشی گذشته و حال و پیش‌بینی‌کننده مسیرهای پژوهشی آینده به کمک بیاید. علم‌سنجی با برخورداری از روش‌ها، فنون و ابزارهای گوناگون قادر است مسیری که حوزه‌های علمی پیموده‌اند و به تکامل رسیده‌اند را رصد کند، حوزه‌های نوظهور آن‌ها را شناسایی کند و

1 . Knowledge Graphs

2 . Explainability

راهی روشن پیش روی پژوهشگران و علاقمندان قرار دهد (Bornmann & Leydesdorff, 2014). در صورت تکیه بر نتایج علم سنجی در مسیر سیاست‌گذاری و برنامه‌ریزی پژوهشی، برنامه‌ها و نقشه راه آینده در حوزه‌های موضوعی، از جمله در حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، از پشتوانه علمی و عملی قابل اتکا برخوردار خواهد شد و احتمال خطا در تدوین راهبردهای عملیاتی، به شدت کاهش می‌یابد.

پیشینه پژوهش

تاکنون پژوهشی از نوع علم‌سنجی در خصوص هوش مصنوعی توضیح‌پذیر (XAI) در داشبوردهای تصمیم‌گیری، یافت نشده است. با وجود این، پژوهش‌های متعددی وجود دارد که در آن‌ها، با تکیه بر فنون و روش‌های علم‌سنجی و کتاب‌سنجی، روندهای موضوعی حوزه‌های علمی گوناگون، از جمله حوزه هوش مصنوعی توضیح‌پذیر، بدون در نظر گرفتن ارتباط آن با داشبوردهای تصمیم‌گیری مورد مطالعه قرار گرفته است. از این مطالعات به عنوان پیشینه پژوهش، الهام گرفته شده است.

در همین راستا، کوتسوپاس و نوسیوس (Koutsoupas & Nosios, 2026)، در پژوهشی، یک تحلیل کتاب‌سنجی نظام‌مند از مقالات مجلات علمی داوری شده را که بین سال‌های ۱۹۷۵ تا ۲۰۲۶م. در پایگاه Scopus حوزه هوش مصنوعی توضیح‌پذیر برای پژوهش‌های علوم اجتماعی و علوم انسانی منتشر شده بود، ارائه دادند. آن‌ها با استفاده از بسته bibliometrix در زبان R و با بررسی روند انتشار مقالات، الگوهای واژگانی، ساختارهای موضوعی و الگوهای استنادی، تلاش کردند نشان دهند ساختاربندی و ادغام هوش مصنوعی توضیح‌پذیر در پژوهش‌های علوم اجتماعی و انسانی چگونه بوده است. در نهایت نشان دادند که پژوهش‌های هوش مصنوعی توضیح‌پذیر در علوم اجتماعی و علوم انسانی از جایگاه مستحکمی در زیرساخت‌های گسترده‌تر یادگیری ماشین برخوردارند. همچنین، نشان دادند که این پژوهش‌ها در پذیرش روش‌ها و تمرکزهای موضوعی، الگوهای متمایزی از خود نشان می‌دهند.

در پژوهشی دیگر، روسو و ویستوکو (Russo & Vistocco, 2026; Talmoudi & Choukir, 2026)، با انجام پژوهش کتاب‌سنجی جامع در خصوص روند توسعه تحقیقات هوش مصنوعی توضیح‌پذیر در بازه زمانی ۱۹۹۳ تا ۲۰۲۴م. با استفاده از روش تحلیل شبکه، تحلیل تطابق چندگانه^۱ و تکنیک‌های هم‌استنادی^۲، به شناسایی مشارکت‌کنندگان کلیدی، خوشه‌های اصلی پژوهشی، روندهای موضوعی و تکامل روش‌شناسی‌ها در این حوزه پرداختند. آن‌ها نشان دادند که پژوهش‌های XAI بر دو محور اصلی «توسعه روش‌های فنی برای بهبود تفسیرپذیری مدل‌ها» و «کاربرد این روش‌ها در حوزه‌های متنوعی مانند سلامت، مدیریت ریسک، علم اقلیم» استوار بوده است. همچنین، آن‌ها شبکه تاریخ‌نگارانه^۳ روند تحول XAI را از مفاهیم بنیادی تا کاربردهای تخصصی به تصویر کشیده و بر رشد میان‌رشته‌ای این حوزه تأکید می‌کند و در نهایت با هدف هدایت پیشرفت‌های آینده و تقویت همکاری بیشتر در این حوزه مهم، بینش‌های ارزشمندی درباره مسیر تحول پژوهش‌های XAI ارائه دادند.

تلمود و چوکیا (Talmoudi & Choukir, 2026) در تحلیلی کتاب‌سنجی مبتنی بر پروتکل پریزما از نرم‌افزار R و بسته bibliometrix و با استفاده از الگوریتم خوشه‌بندی لووین^۴، به تحلیل هم‌رخدادی واژگان و ترسیم نقشه ساختار مفهومی مقالات منتشرشده در حوزه سیستم‌های هشدار زودهنگام پرداختند و نشان دادند، موضوع‌های هوش

1. Multiple Correspondence Analysis
2. Co-Citation
3. Historiographic Network
4. Louvain

تحلیل ساختار دانشی، تکامل موضوعی و جهت‌گیری‌های آینده و نوظهور در پژوهش‌های ...

مصنوعی مولد^۱ و سازه‌های روان‌شناختی از جمله موضوعات محرک و موضوع هوش مصنوعی توضیح‌پذیر به عنوان موضوع نوظهور محسوب می‌شوند. آن‌ها همچنین پیشنهاد دادند که مدیران آموزش عالی باید در حین رسیدگی به حکمرانی اخلاقی هوش مصنوعی، XAI را نیز با سیستم‌های خود ادغام کنند.

دائوویسون (Daovisan, 2026)، در مطالعه‌ای دیگر به بررسی هوش مصنوعی توضیح‌پذیر ۵.۰ (XAI 5.0) در چارچوب فناوری‌های نوظهور برای نوآوری مسئولانه^۲ پرداخت. این مدل‌سازی آینده‌نگرانه فناوری نشان داد که مدل چندجمله‌ای درجه سوم^۳ بهترین برازش را با داده‌ها دارد و با ضرایب تعیین بالا همراه است.

مارچیانو و همکاران (Marciano et al., 2026) در پژوهشی درباره هوش مصنوعی توضیح‌پذیر در حوزه سلامت، به مطالعه روند سالانه انتشار مقالات، نویسندگان اثرگذار، مؤسسات برجسته، توزیع جغرافیایی و الگوهای کلیدواژه‌ها پرداختند و سه خوشه مفهومی اصلی «کاربردهای بالینی و تصویربرداری تشخیصی»، «مبانی فنی شامل XAI و یادگیری عمیق»، و «حکمرانی الگوریتمی»^۴ را شناسایی کردند. آن‌ها این پژوهش را به عنوان نخستین تحلیل کتاب‌سنجی معرفی می‌کنند که به‌طور صریح به تلاقی توضیح‌پذیری، اخلاق و مقررات‌گذاری در زمینه هوش مصنوعی پزشکی می‌پردازد و علاوه بر ارائه یک نمای کلی کمی و یک ترکیب کیفی از این حوزه میان‌رشته‌ای در حال تحول فراهم می‌کند.

با استناد به پیشینه مرور شده، می‌توان دریافت که هوش مصنوعی توضیح‌پذیر از یک مفهوم انتزاعی در علوم رایانه به یک ضرورت راهبردی در حوزه‌های میان‌رشته‌ای همچون سلامت، علوم اجتماعی و حکمرانی اخلاقی بدل شده است. درحالی‌که مطالعات پیشین بر شناسایی خوشه‌های فنی و روند انتشار در حوزه‌های کلی تمرکز داشته‌اند، اما خلأ یک تحلیل جامع علم‌سنجی که به‌طور اختصاصی بر «داشبوردهای تصمیم‌گیری» به عنوان نقطه تلاقی تحلیل بصری و هوش مصنوعی متمرکز باشد، کاملاً مشهود است. پژوهش حاضر با بهره‌گیری از متدولوژی‌های پیشرفته علم‌سنجی و مدل‌سازی آینده‌پژوهانه که در مطالعاتی همچون دائوویسون (Daovisan, 2026) و مارچیانو و همکاران (Marciano et al., 2026) مورد تأکید قرار گرفته، درصدد است تا این شکاف پژوهشی را پوشش داده و نقشه راهی برای ادغام نظام‌مند توضیح‌پذیری در اکوسیستم‌های تصمیم‌گیری ترسیم نماید.

روش‌شناسی پژوهش

این پژوهش با رویکردی کمی و با استفاده از فنون پیشرفته کتاب‌سنجی و تحلیل شبکه، رویکردی علم‌سنجانه را برای ترسیم چشم‌انداز فکری، ساختار دانشی و الگوی حاکم بر پژوهش‌های حوزه «هوش مصنوعی توضیح‌پذیر (XAI) در داشبوردهای تصمیم‌گیری» اتخاذ می‌کند. با توجه به پراکندگی داده‌ها در پایگاه‌های مختلف، فرایند اجرای پژوهش در چهار گام اصلی عملیاتی شد:

۱. شناسایی و استخراج داده‌ها

جهت بازیابی رکوردهای مرتبط با کلیدواژه‌های اصلی و مترادف‌های حوزه «هوش مصنوعی توضیح‌پذیر» و

1. Generative AI
2. Responsible Innovation
3. Cubic Polynomial
4. Algorithmic Governance

«دانشوردهای تصمیم‌گیری^۱، رشته جستجوهای^۲ تخصصی و جامعی مطابق فرمول‌های ۱ و ۲ و با استفاده از عملگرهای منطقی (AND/OR) در تاریخ ۱۸ تا ۲۰ خرداد ۱۴۰۵ در دو پایگاه استنادی Scopus و Web of Science صورت پذیرفت:

فرمول ۱. فرمول جستجو در Scopus

((TITLE-ABS-KEY ("decision support")) OR (TITLE-ABS-KEY ("decision-making")) OR (TITLE-ABS-KEY ("expert system*")) OR (TITLE-ABS-KEY ("decision dashboard*")) OR (TITLE-ABS-KEY ("decision aid*"))) AND (((TITLE-ABS-KEY ("user interface")) OR (TITLE-ABS-KEY ("interface design")) OR (TITLE-ABS-KEY ("ui")) OR (TITLE-ABS-KEY ("dashboard*")) OR (TITLE-ABS-KEY ("visual analytics")) OR (TITLE-ABS-KEY ("human-computer interaction")) OR (TITLE-ABS-KEY ("hci")) OR (TITLE-ABS-KEY ("interaction design")) OR (TITLE-ABS-KEY ("explainable ui")))) AND ((TITLE-ABS-KEY ("explainable machine learning")) OR (TITLE-ABS-KEY (interpretable machine learning)) OR (TITLE-ABS-KEY (xai)) OR (TITLE-ABS-KEY (explainable artificial intelligence)) OR (TITLE-ABS-KEY (explainable ai)))))

فرمول ۲. فرمول جستجو در WoS

TS= (("explainable AI" OR "explainable artificial intelligence" OR "XAI" OR "interpretable machine learning" OR "explainable machine learning") AND ("user interface" OR "interface design" OR "UI" OR "dashboard*" OR "visual analytics" OR "human-computer interaction" OR "HCI" OR "interaction design" OR "explainable UI") AND ("decision support" OR "decision-making" OR "decision dashboard*" OR "expert system*" OR "decision aid*"))

در این مرحله، در مجموع تعداد ۹۴۲ رکورد اولیه، شامل ۷۳۲ رکورد از Scopus و ۲۱۰ رکورد از WoS بازیابی شد. خروجی داده‌های Scopus در قالب‌های CSV و RIS و خروجی داده‌های WoS در قالب‌های BibTex و RIS ذخیره شد تا در مراحل بعد، امکان استفاده در نرم‌افزارهای تحلیل علم‌سنجی فراهم شود:

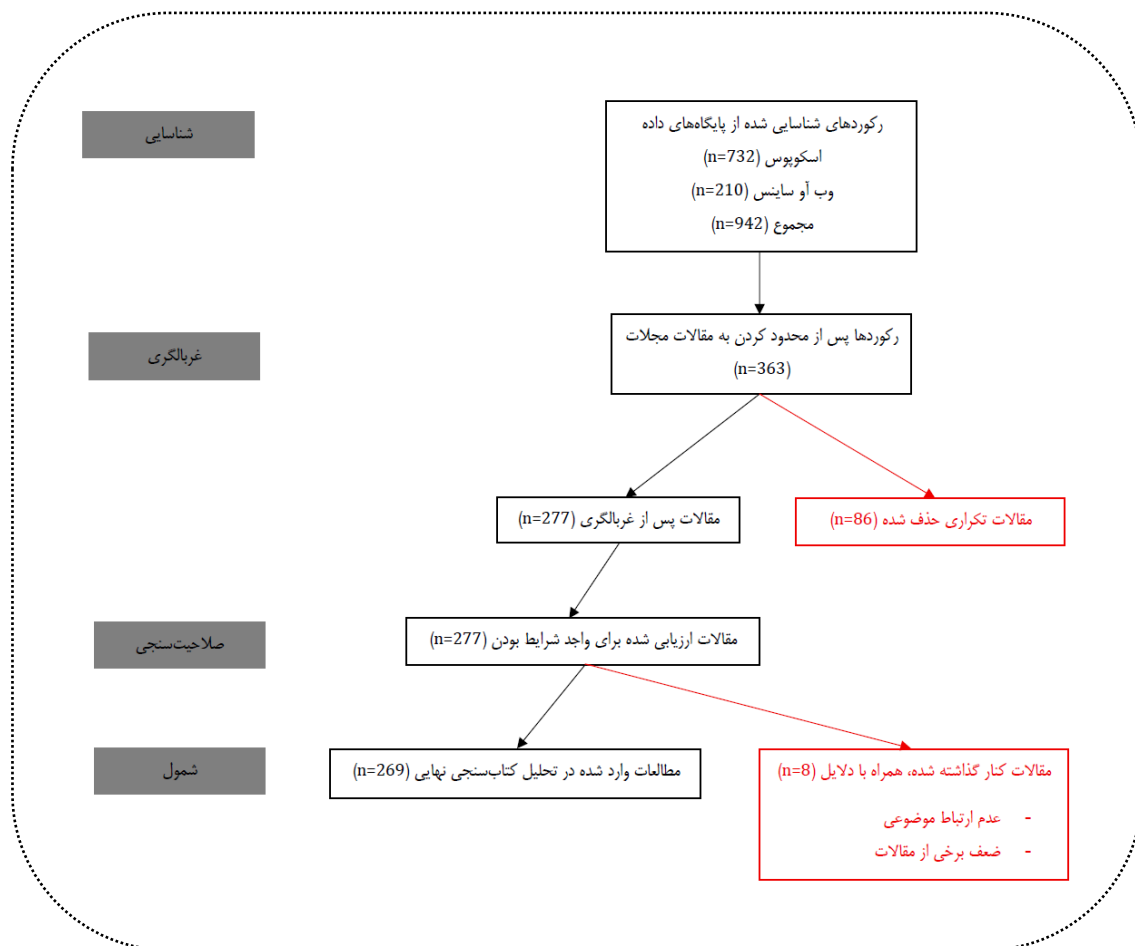
۲. پالایش و استانداردسازی^۳

باتوجه به تفاوت در ساختار متادیتا و قالب‌های استخراجی پایگاه‌های Scopus و WoS، فرآیند غربالگری^۴ و استانداردسازی داده‌ها با هدف دستیابی به یک بانک اطلاعاتی واحد و معتبر طی سه گام اصلی و بر اساس پروتکل پریزما^۵ مطابق تصویر (۱) در چند مرحله انجام شد. نخست، جهت ارتقای روایی نتایج، جامعه پژوهش تنها به «مقالات اصیل پژوهشی منتشر شده در مجلات بین‌المللی و داوری شده معتبر» محدود گردید که منجر به شناسایی ۲۵۱ رکورد از Scopus و ۱۱۲ رکورد از WoS شد. محدود کردن نتایج به مجلات داوری شده از آن جهت حائز اهمیت است که این نوع مجلات، مقالات را پس از فرایند سخت گیرانه داوری علمی مورد پذیرش قرار می دهند و منتشر می کنند (نوروزی چاکلی و همکاران، ۱۴۰۳). در گام دوم، فرآیند ادغام و همسان‌سازی فیلدها در محیط Bibliometrix انجام گرفت که طی آن ۸۶ مورد هم‌پوشانی و رکورد تکراری شناسایی و حذف شدند. در نهایت، با انجام غربالگری کیفی بر اساس چکیده و محتوای تخصصی، ۸ مورد دیگر به دلیل انطباق نداشتن دقیق با اهداف

1. Decision Dashboard
2. Search Strings
3. Data Pre-Processing
4. Screening
5. PRISMA

تحلیل ساختار دانشی، تکامل موضوعی و جهت‌گیری‌های آینده و نوظهور در پژوهش‌های ...

موضوعی پژوهش کنار گذاشته شدند. بدین ترتیب، مجموعه داده نهایی مشتمل بر ۲۶۹ مقاله منتخب برای انجام تحلیل‌های علم‌سنجی و ترسیم نقشه‌های دانش در نرم‌افزارهای Biblioshiny و Vosviewer آماده‌سازی شد.



تصویر ۱. مراحل پژوهش مطابق پروتکل پریزما

۳. کنترل واژگان و تدوین اصطلاح‌نامه^۱

یکی از چالش‌های اصلی در تحلیل‌های واژگانی، وجود مترادف‌ها یا املای متفاوت واژگان بود که می‌توانست نتایج تحلیل خوشه‌بندی را با خطا مواجه کند. برای رفع این چالش، پیش از ورود داده‌ها به بیبیلوشاینی، نسبت به یکدست‌سازی واژگان اقدام شد و همزمان با ورود داده‌ها به Vosviewer، یک اصطلاح‌نامه اختصاصی در قالب فایل اکسل تدوین شد. در این فایل:

- مفاهیم مشابه (مانند "XAI" و "Explainable Artificial Intelligence" یا "User Interface" و "UI") تحت یک برچسب واحد استانداردسازی شدند.
- واژگان عمومی و بی‌ارتباط^۲ با اهداف پژوهش شناسایی و از فرایند تحلیل شبکه حذف شدند تا دقت خوشه‌بندی‌ها افزایش یابد.

1 . Thesaurus Construction

2 . Stop Words

۴. تحلیل‌های علم‌سنجی و تجسم داده‌ها

پس از پالایش و آماده‌سازی داده‌ها، مجموعه نهایی شامل ۲۶۹ مقاله برای انجام تحلیل‌های علم‌سنجی وارد محیط تحلیل شد. در این مرحله از بسته Bibliometrix در محیط Biblioshiny و همچنین نرم‌افزار Vosviewer برای تحلیل شبکه‌ها و تجسم ساختار دانشی استفاده شد. ابتدا داده‌های نهایی در محیط Biblioshiny مورد تحلیل‌های توصیفی و ساختاری قرار گرفت. این تحلیل‌ها شامل بررسی روند رشد سالانه تولیدات علمی، شناسایی نویسندگان و مجلات پربازده و تحلیل توزیع انتشارات علمی بود. در این راستا برای ارزیابی الگوهای بهره‌وری علمی از قانون لوتکا جهت بررسی توزیع تولیدات علمی نویسندگان و از قانون برادفورد برای شناسایی سهم مجلات هسته در حوزه پژوهش استفاده شد. این تحلیل‌ها امکان درک ساختار انتشار دانش و تمرکز مجلات علمی در حوزه «هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری» را فراهم می‌کند.

در گام بعد، برای تحلیل ساختار مفهومی و روابط میان مفاهیم پژوهشی، داده‌ها به نرم‌افزار Vosviewer منتقل شدند. در این نرم‌افزار، با استفاده از روش هم‌رخدادی واژگان^۱، شبکه ارتباطی میان کلیدواژه‌های پژوهش استخراج شد. به منظور افزایش دقت خوشه‌بندی مفاهیم و جلوگیری از پراکندگی اصطلاحات مشابه، فایل اصطلاح‌نامه (Thesaurus) که در مرحله قبل تهیه شده بود در فرایند تحلیل اعمال شد تا واژگان مترادف یا معادل‌های مختلف مفاهیم (مانند XAI، Explainable AI و Explainable Artificial Intelligence) یکپارچه‌سازی شوند.

خروجی‌های حاصل از VOSviewer در قالب سه نوع نقشه تحلیلی ارائه شدند:

- نقشه شبکه‌ای^۲ برای نمایش خوشه‌های مفهومی و روابط میان کلیدواژه‌ها؛
 - نقشه چگالی (بلوغ)^۳ برای شناسایی حوزه‌های پرتراکم و پر مطالعه در ادبیات پژوهش؛
 - نقشه پوششی (همپوشانی زمانی)^۴ برای تحلیل روند زمانی شکل‌گیری و تحول موضوعات پژوهشی.
- این نقشه‌ها به شناسایی ساختار فکری حوزه پژوهش، خوشه‌های دانشی غالب، و مسیرهای پژوهشی نوظهور کمک می‌کنند.

۵. تحلیل ساختار مفهومی و شناسایی روندهای پژوهشی

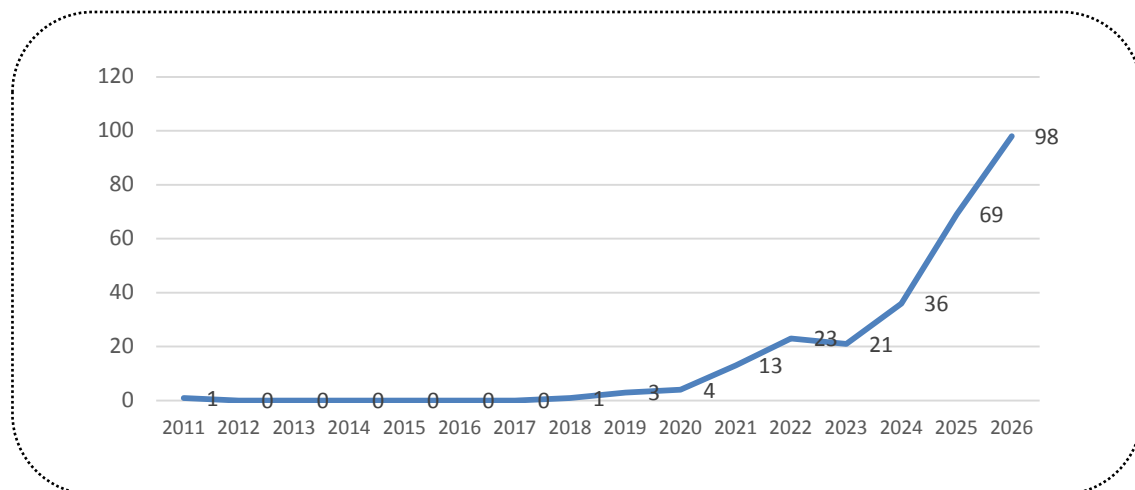
به منظور درک عمیق‌تر ساختار مفهومی حوزه مورد مطالعه، از روش‌های تحلیل شبکه واژگان برای شناسایی خوشه‌های موضوعی استفاده شد. خوشه‌های استخراج شده نمایانگر جریان‌های اصلی پژوهشی در حوزه «هوش مصنوعی توضیح‌پذیر در سیستم‌ها و داشبوردهای تصمیم‌یار» هستند. تحلیل این خوشه‌ها امکان شناسایی موضوعات مرکزی، پیوندهای میان حوزه‌های پژوهشی و روندهای در حال ظهور را فراهم می‌کند. علاوه بر این، با استفاده از نقشه‌های هم‌رخدادی، تحلیل تراکم و چگالی (بلوغ)، تحلیل پوششی (تحلیل همپوشانی زمانی) و تحلیل روند زمانی^۵ کلیدواژه‌ها، موضوعات نوظهور و مسیرهای آینده پژوهش در این حوزه شناسایی شدند. این تحلیل‌ها مبنایی برای ارائه چارچوبی از تحولات دانشی و جهت‌گیری‌های آینده پژوهش فراهم می‌سازد.

1. Co-Occurrence Analysis
2. Network Visualization
3. Density Visualization
4. Overlay Visualization
5. Temporal Aanalysis

یافته های پژوهش

پاسخ به پرسش نخست پژوهش. الگوی رشد و توزیع زمانی ظهور حوزه هوش مصنوعی توضیح پذیر در داشبوردهای تصمیم گیری در سطح بین المللی چگونه است؟

مجموعه اطلاعات ارائه شده در تصویر (۲) و جدول (۱)، به خوبی بر نوظهور بودن حوزه هوش مصنوعی توضیح پذیر در داشبوردهای تصمیم گیری صحنه می گذارد. تصویر (۲) به روشنی نشان می دهد با وجود این که نخستین مقاله بین المللی مرتبط با این حوزه در مجلات پایگاه های WoS و Scopus، در سال ۲۰۱۱ م. منتشر شده است؛ باید سال واقعی آغاز رشد مقاله های این حوزه را از ۲۰۱۸ م. در نظر گرفت؛ زیرا روند رشد مستمر انتشار مقالات این حوزه در مجلات بین المللی از آن سال آغاز شده و تاکنون با رشدی سریع و روزافزون ادامه یافته است. علاوه بر این، جدول ۱ به روشنی حاکی از آن است که این ۲۶۹ مقاله منتشر شده در مجلات بین المللی که تنها دارای ۱۲۹۲ نویسنده بوده، با نرخ رشد سالانه قریب به ۳۶ درصدی، به طور متوسط برای هر مقاله ۲۱ استناد دریافت کرده اند.



تصویر ۲. توزیع زمانی رشد مقالات حوزه هوش مصنوعی توضیح پذیر در داشبوردهای تصمیم گیری در مجلات بین المللی

جدول ۱. توصیف اجمالی وضعیت مقالات منتشر شده حوزه هوش مصنوعی توضیح پذیر در

داشبوردهای تصمیم گیری در مجلات بین المللی

نتایج	توصیف
۲۰۱۱:۲۰۲۶	بازه زمانی
۱۸۹	تعداد مجلات
۲۶۹	مقالات
۳۵.۷۵	نرخ رشد سالانه (%)
۱.۵۹	سن متوسط هر مقاله
۲۱.۱۷	میانگین استناد به هر مقاله
۱۹۹۲	مجموع تعداد کلیدواژه ها ^۱

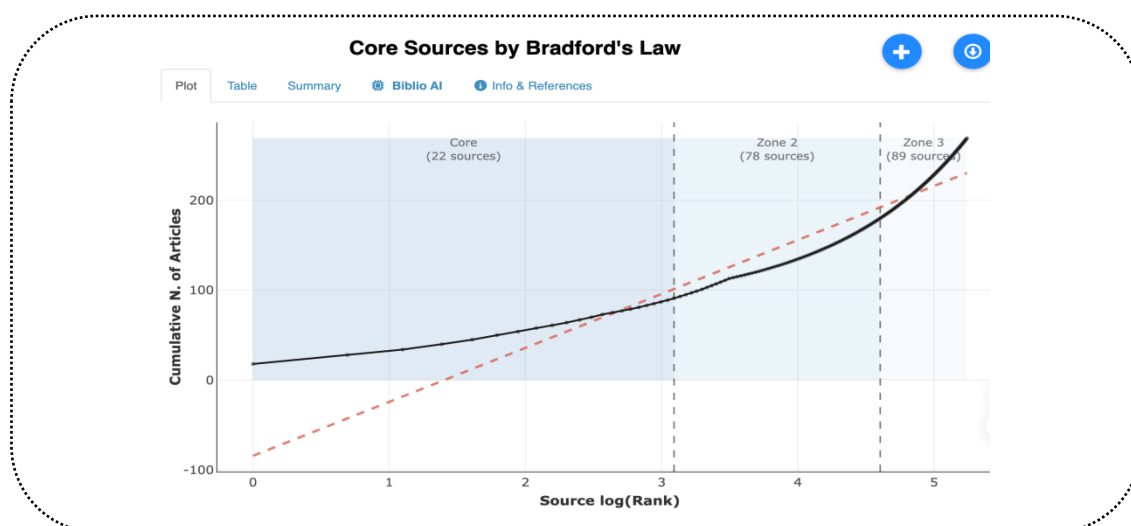
1 . Keywords Plus (ID)

ادامه جدول ۱. توصیف اجمالی وضعیت مقالات منتشر شده حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری در مجلات بین‌المللی

توصیف	نتایج
کلیدواژه‌های نویسندگان ^۱	۱۱۴۰
تعداد نویسندگان	۱۲۹۲
تعداد نویسندگان مقالات تک نویسنده	۹
تعداد هم‌نویسندگی برای هر مقاله	۵.۳
سهم همکاری بین‌المللی (%)	۲۰.۴۵

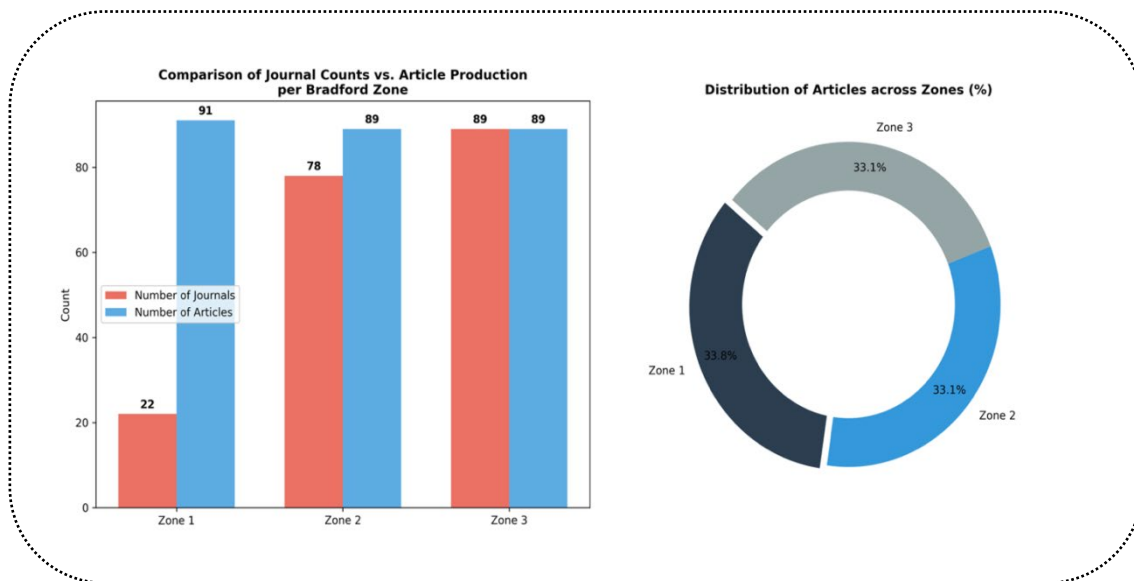
پاسخ به پرسش دوم پژوهش. الگوی توزیع انتشارات در مجلات حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، چگونه از قانون برادفورد پیروی می‌کند و چه نوع ساختاری را در تمرکز و پراکندگی مجلات نشان می‌دهد؟

به منظور شناسایی مجلات هسته و تحلیل ساختار انتشارات در حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، از قانون برادفورد^۲ استفاده شد. این قانون با تقسیم مجلات فعال در هر حوزه موضوعی به سه منطقه^۳، ساختار تمرکز و پراکندگی دانش در آن حوزه را مشخص می‌کند و نشان می‌دهد چگونه، باوجود تقریباً برابر بودن تعداد مقالات در هر منطقه، تعداد مجلات به صورت تصاعدی افزایش می‌یابد (Bradford, 1985). مطابق با منحنی توزیع تجمعی مقالات در برابر لگاریتم رتبه مجلات تصویر (۳)، داده‌های این پژوهش مدل کلاسیک برادفورد (S-Curve) را تأیید می‌کنند. در این پژوهش، تعداد ۲۲ مجله در منطقه هسته (Zone 1) قرار گرفته‌اند که ۳۳.۸٪ از کل حجم انتشارات (۹۱ مقاله) را به خود اختصاص داده‌اند. این امر نشان‌دهنده تمرکز بالای موضوعی در هسته‌ای کوچک و تخصصی از مجلات است.



تصویر ۳. منحنی توزیع تجمعی مقالات در برابر لگاریتم رتبه مجلات بین‌المللی حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری

1. Author's keywords (DE)
2. Bradford's Law
3. Zone



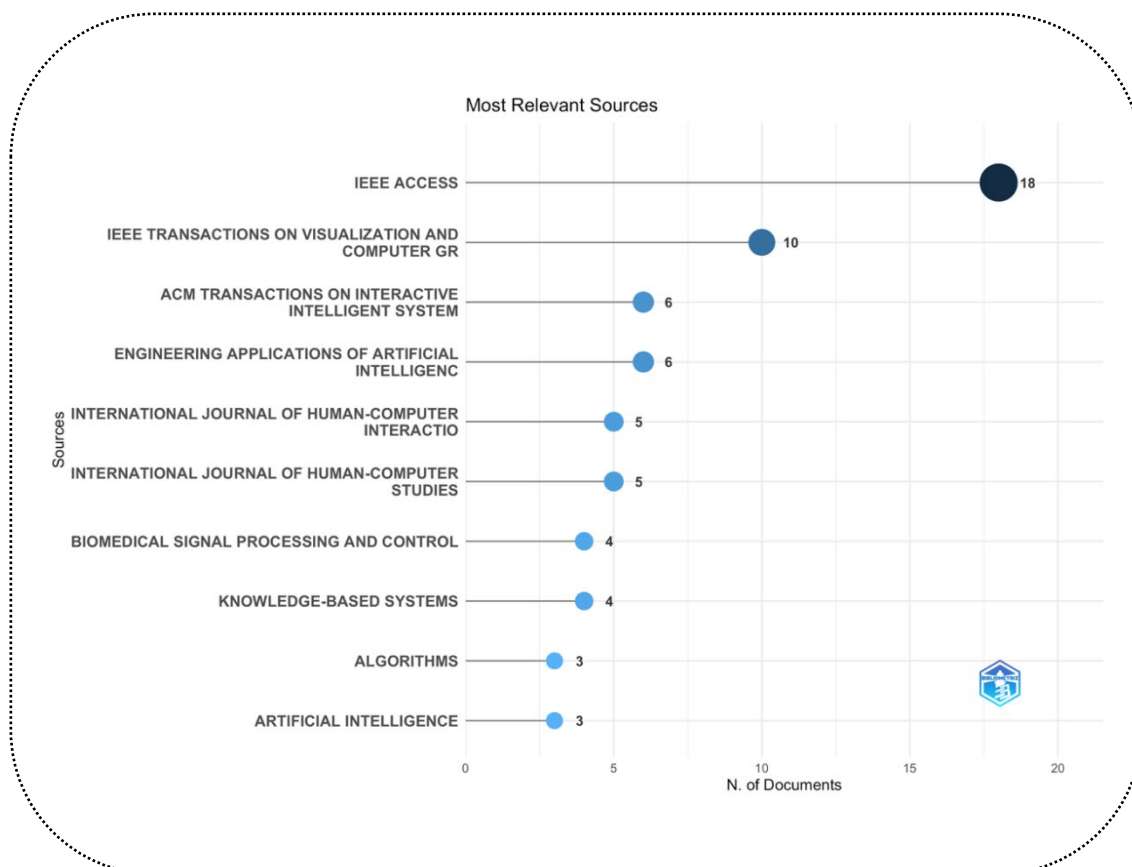
تصویر ۴. مقایسه قانون برادفورد در مجلات بین‌المللی حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری

در مقابل، همان‌طور که در تصویر (۴) مشاهده می‌شود، برای دستیابی به یک‌سوم پایانی مقالات (منطقه حاشیه یا Zone 3)، تعداد مجلات به شدت افزایش یافته و به ۸۹ مجله می‌رسد. در این منطقه، نسبت مجله به مقاله تقریباً ۱ به ۱ است که نشان‌دهنده پراکندگی بسیار زیاد موضوع هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری در مجلات عمومی‌تر و حوزه‌های میان‌رشته‌ای است. به این ترتیب، ساختار توزیع در این حوزه نشان‌دهنده یک «هسته بالغ» است. وجود ۲۲ مجله کلیدی به عنوان هسته علمی، بیانگر آن است که با وجود نوظهور بودن مفاهیمی مانند هوش مصنوعی توضیح‌پذیر (XAI)، این حوزه توانسته است جایگاه تخصصی خود را در مجلات معتبر تثبیت کند. ضریب پراکندگی در منطقه سوم نیز حاکی از پتانسیل بالای این موضوع برای جذب مخاطب در گرایش‌های مختلف علمی (از علوم کامپیوتر تا مدیریت و پزشکی) است و پژوهشگران این حوزه می‌توانند با تمرکز بر مجلات منطقه هسته، به معتبرترین و مرتبط‌ترین جریان‌های دانش دسترسی پیدا کنند. تحلیل ساختار توزیع مقالات در سه منطقه به شرح زیر است:

- منطقه ۱ (هسته): شامل تنها ۸ مجله است که مجموعاً ۵۸ مقاله (حدود ۲۲٪ از کل دانش موضوعی) را منتشر کرده‌اند. این تمرکز بالا در تعداد اندکی از مجلات، نشان‌دهنده شکل‌گیری یک هسته تخصصی و بلوغ نسبی در این حوزه پژوهشی است.
 - منطقه ۲ (مرتبط): شامل ۴۴ مجله است که ۸۹ مقاله را پوشش داده‌اند.
 - منطقه ۳ (حاشیه‌ای): شامل ۱۳۷ مجله تک مقاله‌ای یا با تعداد مقالات بسیار اندک است که نشان‌دهنده ماهیت بین‌رشته‌ای موضوع و نفوذ مفاهیم هوش مصنوعی توضیح‌پذیر به حوزه‌های متنوع علمی (از پزشکی تا مدیریت) است.
- باتوجه به تصویر (۵)، مجلات هسته^۱ (Zone 1) که بیشترین نقش را در تولید دانش این حوزه داشته‌اند به شرح زیر تحلیل می‌شوند:

1 . Core Journals

۱. مجله IEEE Access (رتبه ۱- انتشار ۱۸ مقاله): به عنوان پیشروترین مجله، نشان‌دهنده تمایل پژوهشگران به انتشار در مجلات با دسترسی آزاد^۱ و تمرکز بر جنبه‌های فنی و مهندسی سیستم‌های توضیح‌پذیر است.
۲. مجله IEEE Transactions on Visualization and Computer Graphics (رتبه ۲- انتشار ۱۰ مقاله): حضور این مجله در رده دوم بسیار حائز اهمیت است؛ زیرا تأیید می‌کند که «تصویرسازی داده‌ها»^۲ رکن اصلی در طراحی رابط‌های کاربری هوش مصنوعی توضیح‌پذیر برای داشبوردهاست.
۳. مجلات ACM Transactions on Interactive Intelligent Systems و International Journal of Human Computer Interaction: قرارگیری این مجلات در رده‌های برتر (با ۶ و ۵ مقاله) نشان‌دهنده غلبه رویکرد «تعامل انسان و کامپیوتر» (HCI) در این حوزه است. این موضوع ثابت می‌کند که داشبوردهای هوش مصنوعی توضیح‌پذیر صرفاً یک ابزار محاسباتی نیستند، بلکه بر بهبود تجربه کاربری^۳ و تعامل هوشمند تمرکز دارند.
۴. علاوه بر این، حضور مجله‌ای مانند Biomedical Signal Processing and Control در میان ۸ مجله هسته اصلی، نشان‌دهنده این است که داشبوردهای تصمیم‌گیری مبتنی بر هوش مصنوعی توضیح‌پذیر به طور ویژه‌ای در حوزه پزشکی و سلامت کاربرد عملیاتی پیدا کرده‌اند؛ چرا که در این حوزه، توضیح‌پذیری مدل‌ها با جان انسان‌ها در ارتباط است.



تصویر ۵. مجلات هسته و تأثیرگذار در حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری

1. Open Access
2. Data Visualization
3. UX

تحلیل ساختار دانشی، تکامل موضوعی و جهت‌گیری‌های آینده و نوظهور در پژوهش‌های ...

از طرفی، این الگو بیانگر آن است که پژوهش‌های مربوط به طراحی رابط کاربری برای هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، دارای یک هسته انتشاراتی مشخص و در عین حال ماهیت بین‌رشته‌ای گسترده‌ای هستند.

پاسخ به پرسش سوم پژوهش. الگوی بهره‌وری نویسندگان حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، چگونه سازمان یافته است و بر اساس قانون لوتکا، رهبران علمی تا چه میزان جهت‌گیری و تمرکز بدنه دانش این حوزه را شکل می‌دهند؟

به زبان ساده، قانون لوتکا می‌گوید در هر حوزه علمی، تعداد کمی از نویسندگان بسیار پرکار هستند و اکثریت نویسندگان تنها یک یا دو مقاله منتشر کرده‌اند (Bradford, 1985). اطلاعات مندرج در جدول (۲) حکایت از غلبه نویسندگان تک‌مقاله‌ای (نویسندگان گذرا)^۱ در این حوزه دارد؛ چرا که بیش از ۹۲.۳٪ از نویسندگان (۱۱۹۲ نفر) تنها یک مقاله در حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری منتشر کرده‌اند. این عدد بسیار بالا نشان‌دهنده حوزه «جذاب و نوظهور» است که پژوهشگران زیادی از رشته‌های مختلف برای یک بار وارد آن شده‌اند، اما هنوز به تخصص اصلی و مستمر اکثر آن‌ها تبدیل نشده است.

جدول ۲. توزیع نویسندگان مقالات منتشر شده حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری در مجلات بین‌المللی

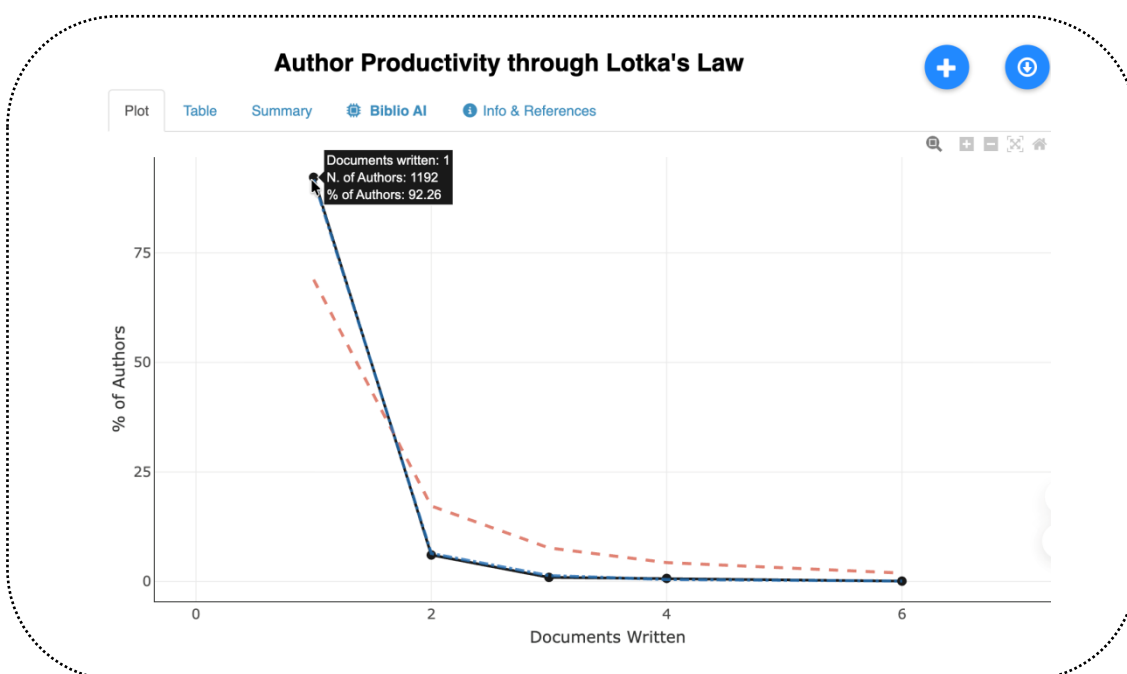
تعداد مقالات نوشته شده	تعداد نویسندگان	نسبت نویسندگان
۱	۱۱۹۲	۰.۹۲
۲	۷۸	۰.۰۶
۳	۱۲	۰.۰۰۹
۴	۹	۰.۰۰۷
۶	۱	۰.۰۰۰۸

در خصوص حضور تأثیرگذارترین نویسندگان، نخبگان حوزه^۲ و نویسندگان پرکار^۳ در این حوزه باید اذعان داشت که تنها ۱ نویسنده (Zhang Y.) وجود دارد که به عنوان فعال‌ترین پژوهشگر این حوزه، ۶ مقاله منتشر کرده است (۰.۱٪ از کل) و تمرکز پژوهش‌های وی، عمدتاً بر «شفافیت در سیستم‌های هوشمند» و «تعامل انسان و کامپیوتر»^۴ است که هسته اصلی موضوع هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری را شکل می‌دهد. همچنین، تنها ۹ نویسنده دارای ۴ مقاله هستند. این افراد معمولاً چارچوب‌های اولیه^۵ برای توضیح‌پذیری در رابط‌های کاربری را تدوین کرده‌اند و همان «قطب‌های علمی» یا رهبران فکری این حوزه هستند که احتمالاً در خوشه‌هایی که در ادامه مقاله ارائه خواهد شد، نقش کلیدی دارند.

به لحاظ نرخ مشارکت نویسندگان، کاهش شدید سهم از ۹۲٪ (برای ۱ مقاله) به ۶٪ (برای ۲ مقاله) در تصویر (۶)

1. Transient Authors
2. Most Productive Authors
3. Core Authors
4. HCI
5. Frameworks

نشان‌دهنده آن است که وفاداری نویسندگان به موضوع خاص رابط کاربری برای هوش مصنوعی توضیح‌پذیر، هنوز در حال شکل‌گیری است و از این نظر با حوزه کلاسیک و قدیمی تفاوت دارد. و این وضعیت نشان‌دهنده نوظهور بودن این گرایش تخصصی است. مقایسه این سهم با قانون لوتکا می‌تواند قابل توجه باشد. در حالی که طبق قانون لوتکا در علم‌سنجی، انتظار می‌رود حدود ۶۰٪ نویسندگان تک‌مقاله‌ای باشند، در داده‌های این پژوهش این میزان ۹۲٪ است. این تفاوت معنادار نشان می‌دهد که موضوع رابط کاربری برای هوش مصنوعی توضیح‌پذیر، یک حوزه به شدت توزیع‌شده، میان‌رشته‌ای و جذاب است؛ به این معنا که متخصصان علوم کامپیوتر، روان‌شناسی، علوم شناختی، پزشکی و مدیریت هر کدام از زاویه دید خود یک بار به این موضوع پرداخته‌اند. به بیانی دیگر، اگرچه بدنه دانش این حوزه توسط تعداد محدودی از متخصصان تعامل انسان و کامپیوتر رهبری می‌شود، اما نفوذ این موضوع در سایر علوم بسیار گسترده است. این موضوع برای داشبوردهای تصمیم‌گیری یک مزیت محسوب می‌شود، زیرا نشان‌دهنده پتانسیل بالای این ابزار برای پذیرش در جوامع علمی گوناگون است.

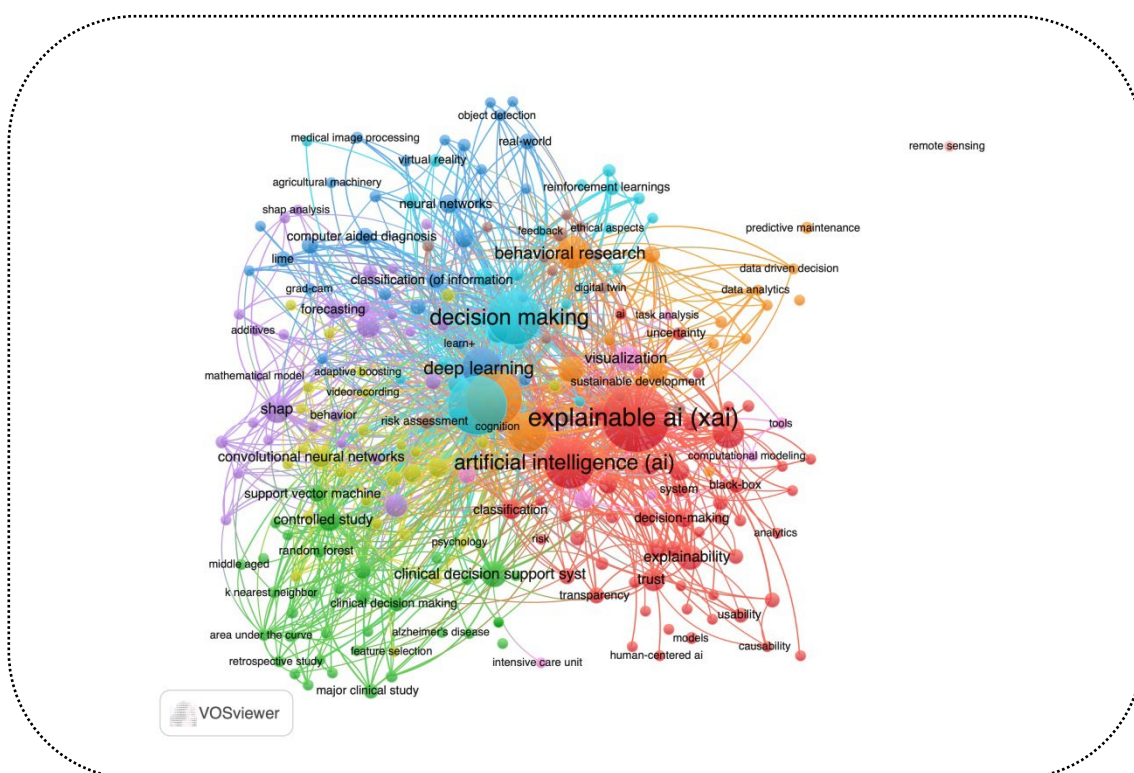


تصویر ۶. الگوی مقایسه‌ای بهره‌وری نویسندگان در حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری با توزیع نظری لوتکا

پاسخ به پرسش چهارم پژوهش. ساختار فکری و راهبردی خوشه‌های اصلی دانش در حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری از چه الگویی تبعیت می‌کند و مؤلفه‌های متمایز و کلیدی هر خوشه کدام‌اند؟

واکاوی ساختار فکری حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری با استفاده از تحلیل شبکه هم‌رخدادی واژگان (تصویر ۷) و نقشه‌های چگالی و زمانی، نشان می‌دهد که این حوزه بر یک اکوسیستم دانش چندبعدی و به هم پیوسته استوار است. این ساختار بر پایه ۱۰ خوشه متمایز استوار است که طیف وسیعی از مباحث، از زیرساخت‌های ریاضی تا ملاحظات اخلاقی و کاربردهای بالینی را در بر می‌گیرند.

تحلیل ساختار دانشی، تکامل موضوعی و جهت‌گیری‌های آینده و نوظهور در پژوهش‌های ...



تصویر ۷. نقشه شبکه‌ای ساختار فکری و خوشه‌های اصلی حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری به منظور تبیین بیشتر، مجموع خوشه‌های معرفی شده در تصویر (۷)، در جدول (۳) دسته‌بندی شده است.

جدول ۳. طبقه‌بندی ساختار فکری و خوشه‌های اصلی حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری مندرج در تصویر ۱

شماره خوشه	عنوان خوشه (Cluster Theme)	تعداد واژگان تمرکز مفهومی	واژگان محوری (Top Keywords)	نقش در نقشه ذهنی
1	حاکمیت اخلاقی و مبانی تعامل انسان و هوش مصنوعی	Trust, Ethics, Human-AI Collaboratio, Transparency	Explainable AI, Trust, Transparency, Ethics, Human-Centered AI, Fairness, Interpretability, Trust Calibration	این خوشه هسته نظری شبکه را تشکیل می‌دهد و نشان می‌دهد ادبیات XAI از رویکرد فنی صرف به سمت رویکردهای انسان‌محور و مدیریت اعتماد حرکت کرده است.
2	کاربردهای بالینی و سیستم‌های پشتیبان تصمیم پزشکی	Clinical Decision Support, Diagnosis, Healthcare	Clinical Decision Support, Medical Diagnosis, Patient Safety, Prognosis, Healthcare AI	این خوشه نشان می‌دهد کاربردهای پزشکی یکی از مهم‌ترین حوزه‌های استفاده از XAI هستند و توضیح‌پذیری در این حوزه برای افزایش ایمنی و اعتماد ضروری است.
3	تکنیک‌های یادگیری ماشین و زیرساخت‌های محاسباتی	Random Forest, SVM, Cybersecurity	Machine Learning, Classification, Feature Selection, Cybersecurity, Prediction	این خوشه لایه فنی شبکه را نشان می‌دهد و بیانگر تمرکز پژوهش‌ها بر توسعه و استفاده از مدل‌های یادگیری ماشین در مسائل پیش‌بینی و امنیت است.

ادامه جدول ۳. طبقه‌بندی ساختار فکری و خوشه‌های اصلی حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری مندرج در تصویر ۱

شماره خوشه	عنوان خوشه (Cluster Theme)	تعداد واژگان	تمرکز مفهومی	واژگان محوری (Top Keywords)	نقش در نقشه ذهنی
4	مدل‌سازی پیش‌بینی‌کننده و یادگیری عمیق (DL)	۲۷	Neural Networks, Deep Learning	Deep Learning, Neural Networks, Pattern Recognition, Representation Learning	این خوشه نشان می‌دهد رشد مدل‌های یادگیری عمیق موجب افزایش نیاز به روش‌های توضیح‌پذیری شده است تا رفتار مدل‌های پیچیده قابل درک باشد.
5	متدولوژی‌های ارزیابی و تحلیل اهمیت ویژگی‌ها (SHAP)	۲۷	Feature Attribution, Sensitivity Analysis	SHAP, Model Evaluation, Calibration, Local Explanation, Global Explanation	این خوشه بیانگر ابزارها و روش‌های اصلی تولید توضیح در مدل‌های هوش مصنوعی است و SHAP به‌عنوان یکی از رایج‌ترین روش‌ها در ادبیات معرفی می‌شود.
6	سیستم‌های خبره و پردازش زبان طبیعی (NLP)	۲۳	Expert Systems, NLP, Human Factors	Expert Systems, Natural Language Processing, Social Effects, Decision Support	این خوشه پیوند میان رویکردهای کلاسیک سیستم‌های خبره و رویکردهای جدید XAI را نشان می‌دهد و ابعاد انسانی و اجتماعی هوش مصنوعی را برجسته می‌کند.
7	تحلیل کلان‌داده‌ها و تصمیم‌گیری داده‌محور صنعتی	۱۷	Big Data, Data Analytics, DSS	Big Data, Data-Driven Decision Making, Industrial Analytics, DSS	این خوشه کاربرد XAI در محیط‌های سازمانی و صنعتی را نشان می‌دهد که در آن تحلیل داده‌های بزرگ برای پشتیبانی از تصمیم‌گیری استفاده می‌شود.
8	سلامت دیجیتال و اینترنت اشیا (IoT)	۱۵	Wearables, Monitoring, Telemedicine	IoT, Digital Health, Remote Monitoring, Smart Systems	این خوشه نشان‌دهنده ارتباط میان XAI و داده‌های حسگری در محیط‌های سلامت دیجیتال و سیستم‌های هوشمند است.
9	سیستم‌های خودمختار و تحلیل وظایف شناختی	۹	Autonomous Vehicles, Cognition, Risk	Autonomous Systems, Human Cognition, Risk Assessment, Task Analysis	این خوشه بر نقش توضیح‌پذیری در سیستم‌های خودمختار تمرکز دارد تا کاربران بتوانند رفتار سیستم را در شرایط مختلف درک کنند.
10	سنجش از دور و حوزه‌های جانبی	۱	Remote Sensing	Remote Sensing	این خوشه نشان‌دهنده کاربردهای پراکنده و تخصصی XAI در حوزه‌هایی مانند سنجش از دور است که هنوز تراکم پژوهشی کمتری دارند.

تحلیل ساختار دانشی، تکامل موضوعی و جهت‌گیری‌های آینده و نوظهور در پژوهش‌های ...

دسته‌بندی خوشه‌ها در جدول (۳) نشان می‌دهد که پژوهش در حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، فقط به توسعه مدل‌های پیش‌بینی محدود نیست، بلکه به‌طور فزاینده‌ای بر فهم‌پذیری مدل، تعامل انسان و کامپیوتر و اعتبار تصمیمات تولیدشده توسط سامانه‌ها تمرکز دارد. در مجموع، ساختار فکری این حوزه نشان می‌دهد که هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری از یک موضوع صرفاً الگوریتمی به یک حوزه انسان‌محور، کاربردی و مبتنی بر اعتماد تبدیل شده است. مطابق جدول (۳)، بدنه اصلی خوشه‌های دانش دهگانه در این حوزه را می‌توان به پنج کلان‌دسته طبقه‌بندی کرد:

هسته انسانی - اخلاقی (خوشه ۱): این خوشه با تمرکز بر مفاهیمی همچون Trust و Ethics، نشان می‌دهد که حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری از یک ابزار فنی صرف به یک واسطه «اعتمادساز» تبدیل شده است که هدف آن نهادینه کردن اعتماد کاربر و رعایت اصول اخلاقی در هوش مصنوعی انسان‌محور است.

- **هسته فنی - روش‌شناسانه (خوشه‌های ۳، ۴، ۵ و ۶):** این بخش، موتور محرک و زیربنای محاسباتی حوزه است. در اینجا تقابل میان پیچیدگی مدل‌های Deep Learning و متدهای تبیین‌گر نظیر SHAP و LIME جریان دارد. همچنین پیوند میان سیستم‌های خبره کلاسیک و پردازش زبان طبیعی (NLP) در این هسته، نشان‌دهنده تلاش برای ایجاد زبان مشترک میان ماشین و انسان است.

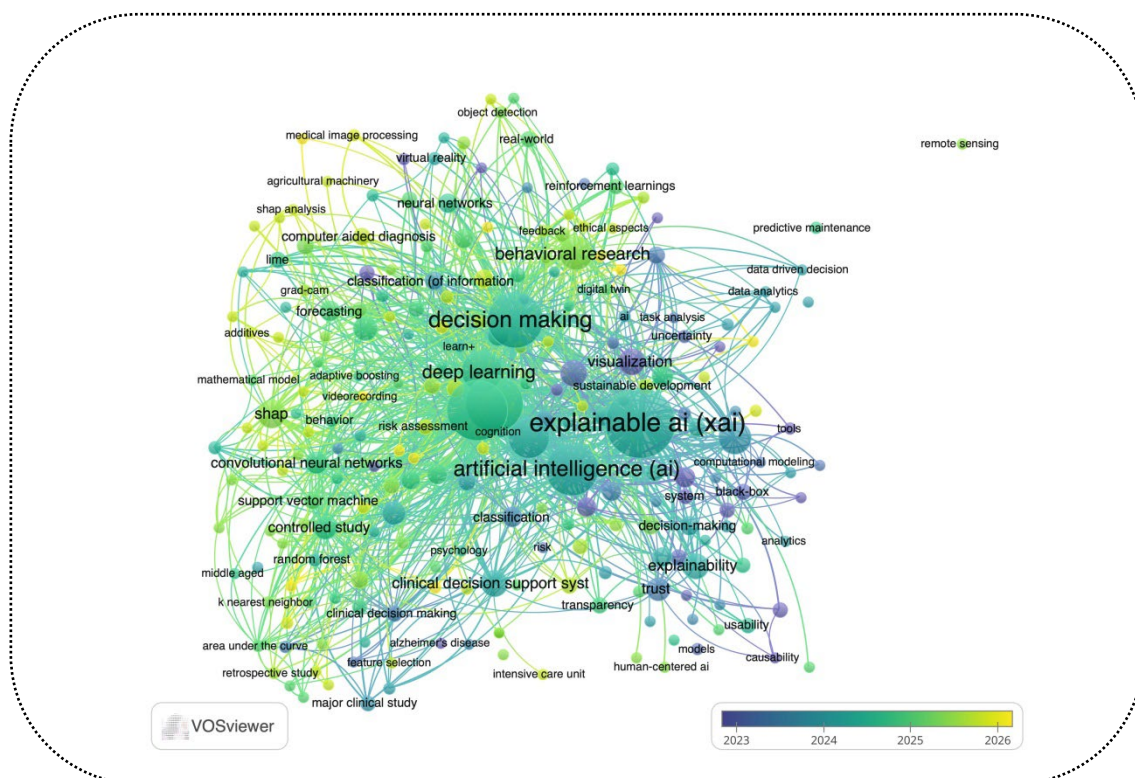
- **هسته کاربردی - عملیاتی (خوشه‌های ۲، ۷ و ۸):** این هسته بر پیاده‌سازی در دنیای واقعی تمرکز دارد. حوزه‌های سلامت (Clinical DSS)، صنعت و اینترنت اشیا (IoT)، به عنوان اصلی‌ترین بسترهای عملیاتی شناسایی شده‌اند که در آن‌ها توضیح‌پذیری عاملی حیاتی برای پذیرش سیستم‌های پشتیبان تصمیم محسوب می‌شود.

- **هسته شناختی و سیستم‌های خودمختار (خوشه ۹):** این بخش از شبکه که بر مفاهیمی چون Cognition و Task Analysis تمرکز دارد، لایه عمیق‌تری از تحلیل است. در این لایه، پژوهشگران به دنبال درک فرآیندهای ذهنی کاربر در مواجهه با سیستم‌های خودمختار هستند تا توضیحاتی ارائه دهند که با مدل‌های شناختی انسان همخوانی داشته باشد.

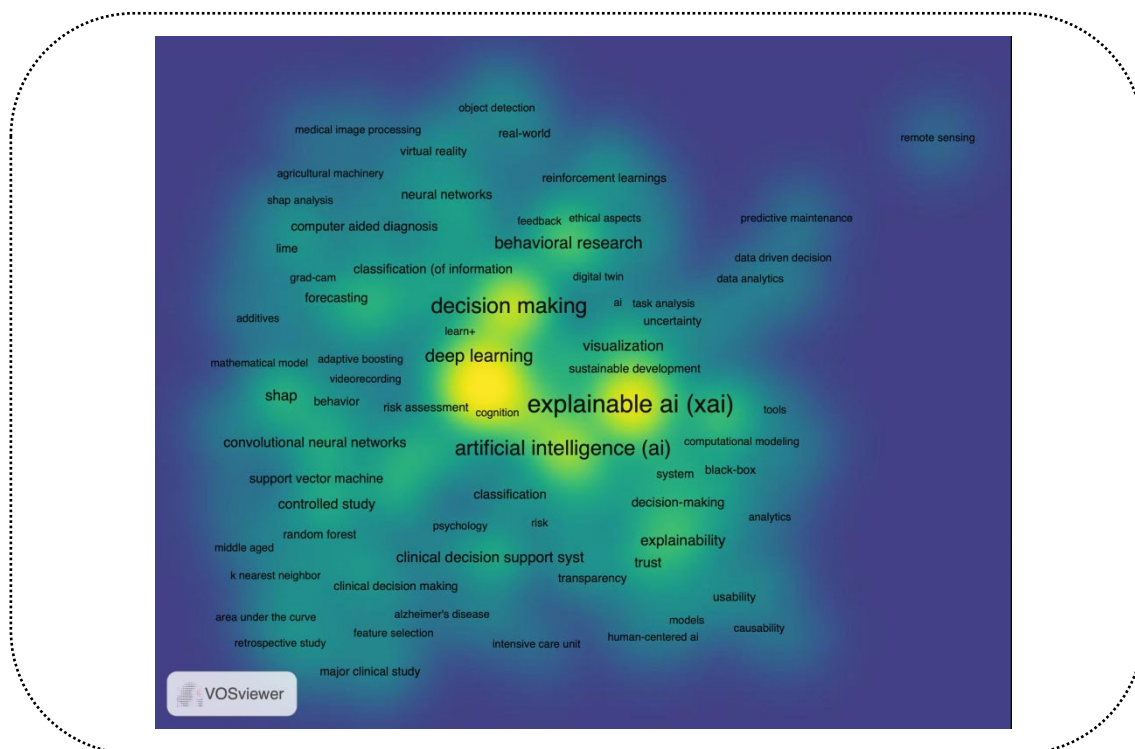
- **هسته موضوعات نوظهور و حوزه‌های حاشیه‌ای (خوشه ۱۰):** حضور موضوعاتی مانند Remote Sensing سنجش از دور در این بخش، گویای نفوذ تدریجی مفاهیم توضیح‌پذیری به حوزه‌های تخصصی و جدید است. اگرچه این خوشه در حال حاضر تراکم کمتری دارد، اما پتانسیل رشد و تبدیل شدن به خوشه‌های کاربردی در آینده را داراست.

در کنار نقشه شبکه ارائه شده در تصویر (۷)، نقشه چگالی واژگان (تصویر ۸) نشان می‌دهد که بیشترین تراکم پژوهشی و «نقطه داغ»^۱ حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، در مثلث میان Decision Making، Explainable AI و Deep Learning قرار دارد. شدت رنگ در این نواحی تأیید می‌کند که دغدغه اصلی پژوهشگران، «توضیح‌پذیر کردن مدل‌های یادگیری عمیق برای ارتقای کیفیت تصمیم‌گیری» است. سایر مفاهیم مانند Clinical DSS در رده دوم چگالی قرار دارند که نشان‌دهنده بلوغ نسبی کاربرد XAI در پزشکی نسبت به سایر صنایع است.

1 . Hotspot



تصویر ۸. نقشه تحلیل چگالی و تراکم (بلوغ) پژوهشی ساختار فکری و خوشه های اصلی حوزه هوش مصنوعی توضیح پذیر در داشبوردهای تصمیم گیری



تصویر ۹. نقشه تحلیل بلوغ روند و تکامل زمانی ساختار فکری و خوشه های اصلی حوزه هوش مصنوعی توضیح پذیر در داشبوردهای تصمیم گیری

تحلیل ساختار دانشی، تکامل موضوعی و جهت‌گیری‌های آینده و نوظهور در پژوهش‌های ...

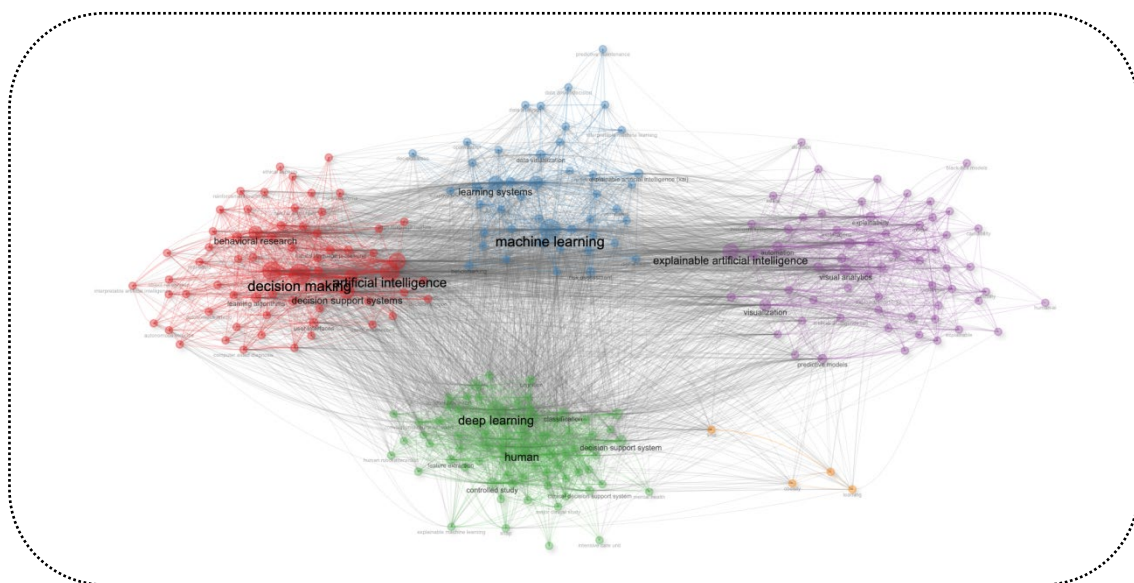
همچنین نقشه هم‌پوشانی زمانی (تصویر ۹) تحول راهبردی موضوعات را نمایان می‌سازد. در تصویر (۹)، موضوعات پیشرو و کلاسیک نظیر Neural Networks و Support Vector Machine با رنگ‌های سرد (بنفش) مربوط به سال‌های آغازین هستند که بر جنبه‌های فنی هوش مصنوعی کلاسیک تمرکز داشتند. علاوه بر این، موضوعات معاصر و در حال ظهور با رنگ‌های گرم (زرد) مشخص شده‌اند. به این ترتیب، مفاهیمی که به رنگ زرد متمایل هستند (مربوط به بازه ۲۰۲۵-۲۰۲۶م.) شامل Human-Centered AI, Causability, Trust Calibration و روش‌های پیشرفته ارزیابی در محیط‌های Real-world هستند. این روند نشان می‌دهد که کانون توجه پژوهش‌ها از «دقت مدل» به سمت «تأثیر بر کاربر» و «تفسیر علی» حرکت کرده است.

به این ترتیب، با توجه به تصاویر (۷ تا ۹) وس و ویوور و جدول (۳) مرتبط با آن‌ها، در مجموع ساختار فکری هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری نشان‌دهنده الگوی «تکامل از مدل محوری به انسان‌محوری» است. وجود خوشه‌های تخصصی مانند خوشه ۵ (ارزیابی با SHAP) در کنار خوشه ۹ (شناخت انسانی)، گویای این واقعیت است که داشبوردهای تصمیم‌گیری مدرن، محصول تلفیق الگوریتم‌های پیچیده ریاضی با مدل‌های ذهنی کاربران هستند. این ساختار نه تنها به دنبال «چگونگی» کارکرد مدل، بلکه به دنبال «چگونگی پذیرش» آن توسط تصمیم‌گیرندگان است.

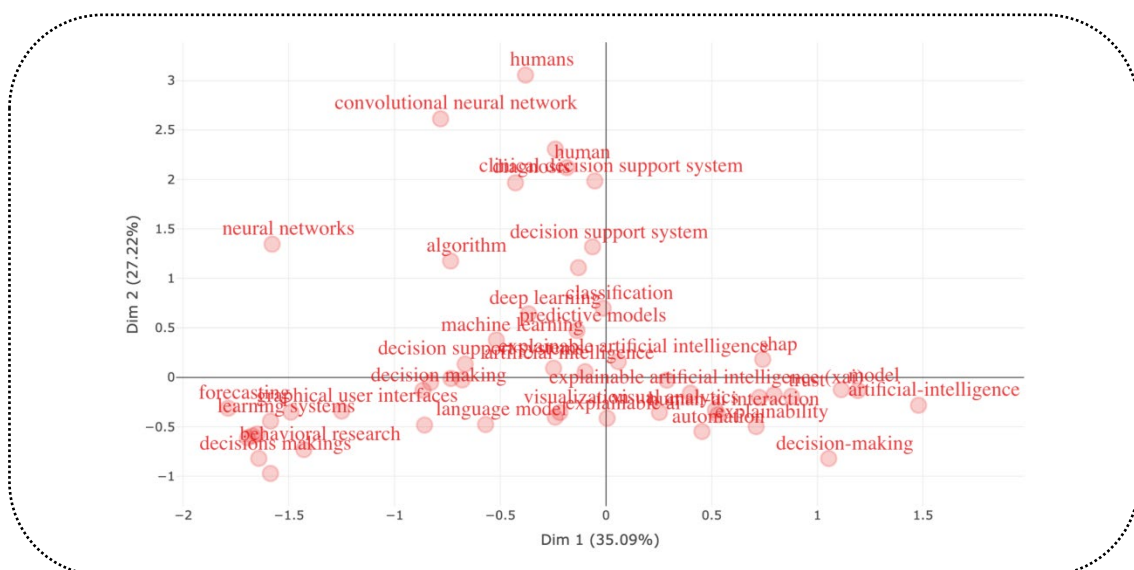
با وجود این، برای تنوع بخشی و عمق بخشی تحلیل ساختار مفهومی و راهبردی موضوعات، تدوین نقشه تماتیک^۱ (تصویر ۱۰) و تحلیل عاملی^۲ بیلبوشاینی (تصویر ۱۱) می‌تواند مفید و ضروری باشد؛ چرا که این دو نوع تحلیل به خوبی نشان می‌دهند که مفاهیم اصلی این حوزه در چه موقعیتی از نظر مرکزیت و توسعه قرار دارند و چه موضوعاتی به عنوان موتورهای دانشی آن عمل می‌کنند. در نقشه تماتیک بیلبوشاینی (تصویر ۱۰) چهار خوشه اصلی قابل مشاهده است که یکی بر مفاهیم فنی مانند Machine learning, Deep learning و Human متمرکز است که نشان می‌دهد مبانی فنی و انسانی به صورت هم‌زمان در پیش‌رانش حوزه نقش دارند. در این بخش، یادگیری ماشین و یادگیری عمیق همچنان از نظر مرکزیت و پویایی در سطح بالایی قرار دارند، اما به‌تنبهایی کافی نیستند و در پیوند با عامل انسانی معنا پیدا می‌کنند. خوشه دیگر با تمرکز بر موضوعاتی نظیر Artificial intelligence, Explainable AI, Behavioral research و Decision making، در موقعیتی قرار گرفته‌اند که بیانگر اهمیت راهبردی آن‌ها در ساختار پژوهش است. این خوشه نشان می‌دهد که پژوهش این حوزه به سمت طراحی تعامل مؤثر انسان و کامپیوتر و استفاده از هوش مصنوعی توضیح‌پذیر برای تقویت تصمیم‌گیری انسانی حرکت کرده است. خوشه‌ای دیگر که می‌توان از آن با عنوان خوشه موضوعات نوظهور یا رو به افول یاد کرد، بر موضوعاتی همچون Explainability, Visualization و Visual analytics تمرکز دارد. حضور این موضوعات در این خوشه بیانگر آن است که ابزارهای بصری و روش‌های تبیینی هنوز در حال تثبیت و بلوغ هستند و با وجود اهمیت زیاد، در مقایسه با هسته‌های مرکزی، هنوز در مرحله توسعه روش‌شناختی و کاربردی قرار دارند. علاوه بر این، خوشه کوچک‌تری نیز بر Forecasting و Machine learning تمرکز دارد. در مجموع، این الگو نشان می‌دهد که ساختار مفهومی هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری به صورت یک شبکه چندلایه سازمان یافته است: لایه نخست شامل مبانی مدل‌سازی و یادگیری ماشین، لایه دوم شامل توضیح‌پذیری و تبیین تصمیم، و لایه سوم شامل کاربردهای تصمیم‌یار و انسانی.

1 . Biblioshiny Thematic Map/ Network View

2 . Factorial Analysis



تصویر ۱۰. نقشه تمانیک ساختار فکری و خوشه‌های اصلی حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری ازسوی‌دیگر، نقشه تحلیل عاملی بیلیوشاینی (تصویر ۱۱) نیز همین الگو را تأیید می‌کند. در این تحلیل، واژگانی مانند Decision, Trust, Decision-making, Explainable AI, Explainability, Artificial-intelligence support system در نزدیکی هم قرار گرفته‌اند و نشان می‌دهند که پیوند میان توضیح‌پذیری، اعتماد، و پشتیبانی تصمیم از مهم‌ترین بنیان‌های مفهومی حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری است. در مقابل، واژگانی مانند User interfaces, Forecasting, Neural networks, Convolutional neural network, Clinical decision support systems و Natural language processing در پیرامون این هسته قرار گرفته‌اند و نقش حوزه‌های پشتیبان یا کاربردی را ایفا می‌کنند.



تصویر ۱۱. نقشه تحلیل عاملی ساختار فکری و خوشه‌های اصلی حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری

تحلیل ساختار دانشی، تکامل موضوعی و جهت‌گیری‌های آینده و نوظهور در پژوهش‌های ...

بنابراین، ساختار مفهومی حوزه نشان می‌دهد که توضیح‌پذیری، اعتماد، تعامل انسان و کامپیوتر و تصمیم‌گیری محورهای راهبردی دانش در این زمینه هستند و ابزارهای بصری و تحلیلی، در حال گذار از موضوعات نوظهور به مؤلفه‌های تثبیت‌شده پژوهش محسوب می‌شوند.

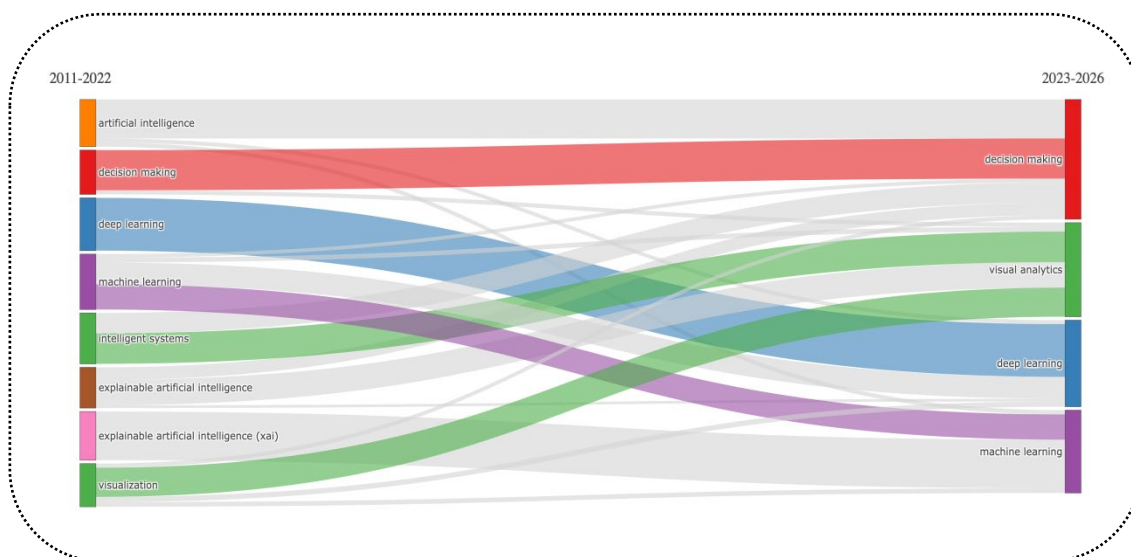
پاسخ به پرسش پنجم پژوهش. با رویکرد آینده‌پژوهانه، موضوعات هسته اصلی و نوظهور پژوهش در حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، کدام‌اند و روند تکامل ساختار دانشی آن‌ها در طول زمان چگونه قابل تفسیر است؟

برای شناسایی موضوعات هسته اصلی و بررسی روند تکامل موضوعی^۱ و آینده‌پژوهی حوزه، مقالات در دو دوره زمانی ۲۰۱۱-۲۰۲۲ م. و ۲۰۲۳-۲۰۲۶ م. تقسیم شد و با استفاده از تحلیل تراکم و تحلیل همپوشانی زمانی، روند تکامل موضوعی مورد مقایسه قرار گرفت و نتایج در تصاویر (۱۲)، (۱۳) و (۱۴) ارائه شد. این تحلیل امکان تشخیص موضوعات پرتکرار، خوشه‌های داغ پژوهشی و نیز مسیر تحول ادبیات در طول زمان را فراهم می‌کنند. تحلیل روند تکامل موضوعی مندرج در تصویر (۱۲)، نمایانگر یک گذار ساختاری و پارادایمیک در بدنه دانش این حوزه است. یافته‌های حاصل از این نمودار حاکی از آن است که مفهوم Decision making در هر دو دوره (۲۰۱۱-۲۰۲۲ م.) و (۲۰۲۳-۲۰۲۶ م.) به‌عنوان قدرتمندترین و پایدارترین جریان (ستون فقرات پژوهش) باقی مانده است. ثبات این جریان نشان‌دهنده آن است که هدف غایی تمامی نوآوری‌های فنی در حوزه هوش مصنوعی توضیح‌پذیر، همواره پشتیبانی مستمر از فرآیند تصمیم‌گیری بوده و این اصالت موضوعی در طول زمان حفظ شده است. با این حال، جذاب‌ترین یافته در این نقشه تکاملی (تصویر ۱۰)، نحوه جابه‌جایی و بازآرایی سایر خوشه‌هاست؛ به طوری که بخش‌های بزرگی از جریان‌های Artificial Intelligence و Intelligent Systems در دوره اول، در دوره دوم به خوشه‌های Visual analytics و Decision making ریخته شده‌اند. این تغییر به معنای عبور از «هوش مصنوعی انتزاعی» به سمت «هوش کاربردی و بصری» است؛ جایی که هوش مصنوعی نه به عنوان یک جعبه سیاه، بلکه در قالب ابزارهای تحلیل بصری برای انسان معنا می‌یابد. علاوه بر این، برخلاف تصورات اولیه، جریان Deep Learning که در دوره اول نیز حضور داشت، در دوره دوم با جذب بخشی از بدنه Machine Learning و سایر سیستم‌های هوشمند، ضخامت و اهمیت بیشتری یافته است. این واقعیت نشان‌دهنده آن است که یادگیری عمیق از یک رویکرد فرعی، به زیربنای فنی غیرقابل انکار در سیستم‌های توضیح‌پذیر مدرن تبدیل شده است. نکته حائز اهمیت دیگر، نحوه اتصال خوشه‌های Explainable Artificial Intelligence از دوره اول به خوشه‌های Visual Analytics و Machine Learning در دوره دوم است. این پیوند تأیید می‌کند که «توضیح‌پذیری» دیگر یک لایه الحاقی نیست، بلکه در تار و پود روش‌های یادگیری ماشین و ابزارهای تحلیل بصری ادغام شده است. بنابراین، مطالعه روند تکامل موضوعی حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، نشان‌دهنده یک چرخش راهبردی^۲ از «ساخت سیستم‌های هوشمند کلاسیک» به سمت «توسعه اکوسیستم‌های یادگیری عمیق بصری» است که مستقیماً در خدمت ارتقای کیفیت تصمیم‌گیری انسانی قرار دارند.

برش اختصاصی دوره زمانی ۲۰۱۱-۲۰۲۲ م. در تصویر (۱۳) این موارد را تأیید می‌کند. در این دوره، موضوعات تحلیل

1 . Thematic Evolution

2 . Strategic Shift



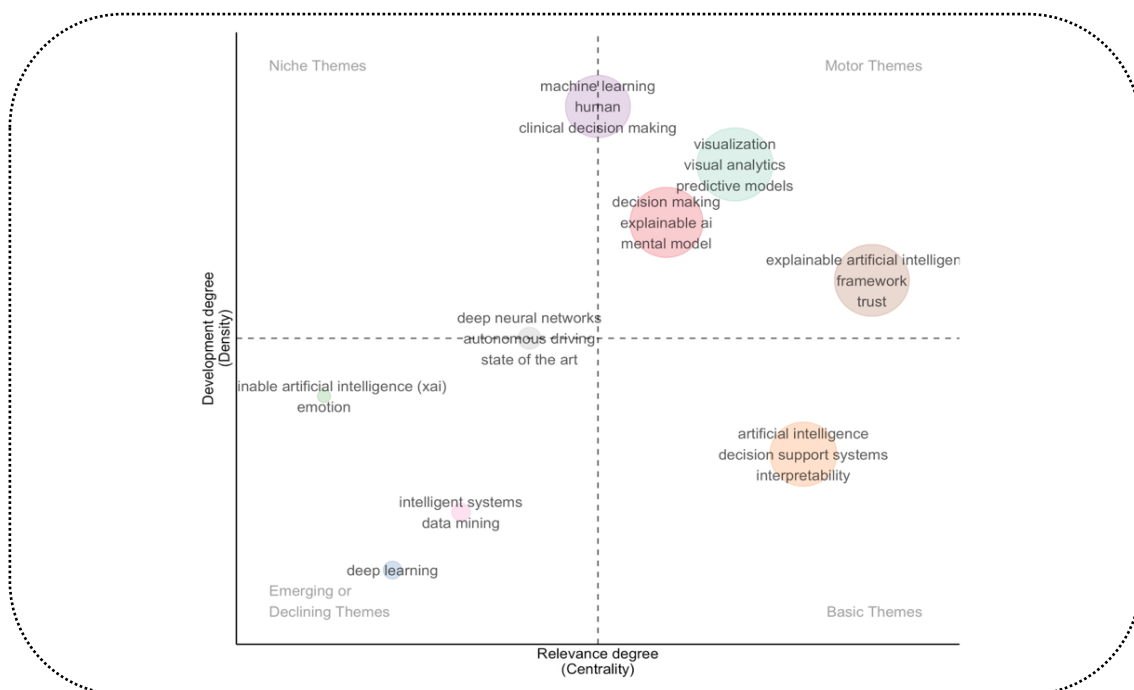
تصویر ۱۲. تکامل ساختار دانش در حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، در دو بازه قبل و بعد از ۲۰۲۲م.

بصری و مصورسازی^۱ به همراه موضوع مدل‌های پیش‌بینی‌گر^۲ در ربع موضوعات محرک و توسعه‌یافته^۳ محسوب می‌شوند. این نشان می‌دهد که در این دوره، «تحلیل بصری» و «مدل‌های پیش‌بینی‌گر» با توسعه یافتگی بالا و همچنین مرکزیت بالا، هسته مرکزی و توسعه‌یافته‌ترین بخش حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری بوده‌اند. به علاوه، خوشه Explainable AI همراه با مفاهیم Trust (اعتماد) و Framework در لبه ربع موضوعات محرک و پایه قرار گرفته است. این نشان‌دهنده تلاش پژوهشگران برای ایجاد «چارچوب‌های نظام‌مند» جهت جلب اعتماد کاربر به هوش مصنوعی است. در عین حال، خوشه Human, Machine Learning و Clinical Decision Making با ویژگی توسعه یافتگی بالا و مرکزیت پایین، در ناحیه موضوعات تخصصی و حاشیه‌ای^۴ قرار دارد که نشان می‌دهد پیش از سال ۲۰۲۲م، کاربرد هوش مصنوعی توضیح‌پذیر در «تصمیم‌گیری‌های بالینی و پزشکی» یک حوزه بسیار تخصصی (Niche) بوده که روابط درونی قوی داشته، اما هنوز به کل بدنه دانش متصل نشده بود. علاوه بر این، خوشه پایه و اساسی با ویژگی توسعه یافتگی پایین و همچنین مرکزیت پایین، شامل Artificial Intelligence, Decision Support Systems و Interpretability است که در این ربع قرار دارد. موضوع‌ها الفبای حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری محسوب می‌شوند. افزون بر این، Interpretability (تفسیرپذیری) برخلاف Explainability، به عنوان مفهومی پایه^۵ و بنیادین شناخته می‌شود. نکته بسیار مهم دیگر، موضوع Deep Learning است که در این بازه (تا سال ۲۰۲۲م)، در دورترین نقطه ربع نوظهور قرار دارد. این امر تأیید می‌کند که در این دوره، یادگیری عمیق هنوز جایگاه تثبیت‌شده‌ای در داشبوردهای تصمیم‌گیری مبتنی بر هوش مصنوعی توضیح‌پذیر نداشته و به عنوان رویکردی جدید، در حال ورود به این حوزه بوده

1. Visualization
2. Predictive Models
3. Motor Themes
4. Niche Themes
5. Basic Themes

تحلیل ساختار دانشی، تکامل موضوعی و جهت‌گیری‌های آینده و نوظهور در پژوهش‌های ...

است. از همین رو، در بازه پیش از سال ۲۰۲۲ م. موضوع Deep learning در حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری هنوز دارای توسعه یافتگی پائین و مرکزیت پائین بوده و به همین دلیل، در ناحیه موضوعات نوظهور قرار گرفته بود.



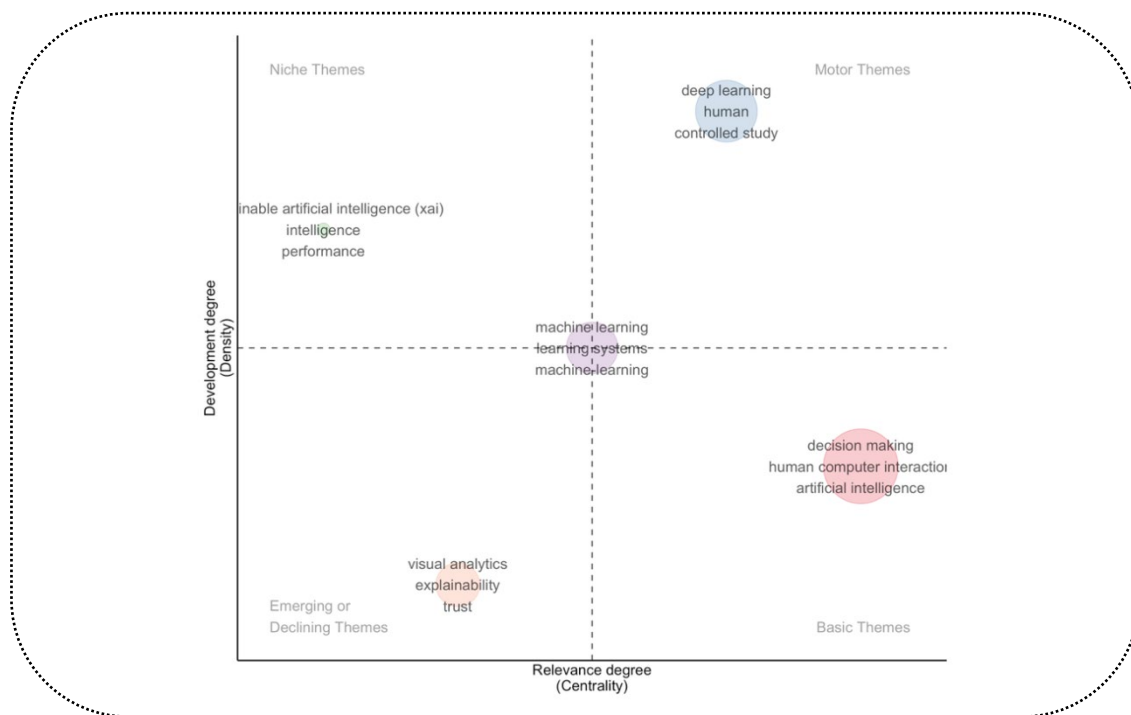
تصویر ۱۳. روند راهبردی تکامل ساختار دانش در حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، در دوره زمانی ۲۰۲۲-۲۰۱۱ م. (دوران شکل‌گیری و تمرکز بصری)

این در حالی است که مطابق جزئیات مندرج در تصویر (۱۴)، در دوره دوم (۲۰۲۳-۲۰۲۶ م.) یک جابه‌جایی قدرت^۱ در مفاهیم مشاهده می‌شود. مهم‌ترین تغییر این دوره نسبت به دوره قبل، صعود خیره‌کننده Deep Learning به ناحیه موضوعات محرک و توسعه‌یافته است. به عبارتی، ظهور یک قدرت جدید در این حوزه است. به این ترتیب، در سال‌های اخیر، یادگیری عمیق دیگر یک موضوع حاشیه‌ای نیست، بلکه به موتور اصلی پیش‌برنده هوش مصنوعی توضیح‌پذیر در داشبوردها تبدیل شده است؛ به طوری که پژوهش‌ها اکنون بر لایه‌های عمیق مدل‌ها تمرکز دارند. همچنین، علاوه بر Deep Learning، شامل موضوعات Human و Controlled Study نیز در خوشه محرک و توسعه‌یافته دیده می‌شود. این می‌تواند یکی از مهم‌ترین یافته‌های این پژوهش قلمداد شود؛ چرا که یادگیری عمیق که در دوره قبل «نوظهور» بود، اکنون به محرک اصلی (Motor Theme) تبدیل شده است. حضور عبارت Controlled Study در کنار آن نشان می‌دهد که پژوهش‌ها از مباحث تئوریک فراتر رفته و به سمت «مطالعات کنترل‌شده تجربی» برای سنجش تأثیر این مدل‌ها بر انسان حرکت کرده‌اند. به علاوه، در ناحیه موضوعاتی پایه‌ای، جایگاه خوشه بزرگی شامل Decision Making، Human Computer Interaction (HCI) و Artificial Intelligence تثبیت شده است. این نشان می‌دهد که تعامل انسان و کامپیوتر اکنون به «ستون فقرات» و دانش پایه این حوزه تبدیل شده است؛ به طوری که هر پژوهشی در این دوره، الزماً باید به جنبه‌های انسانی نیز توجه کند. علاوه بر

1 . Power Shift

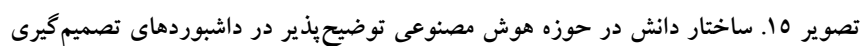
این، قرارگیری خوشه موضوعی (XAI) Explainable Artificial Intelligence و Performance در ربع موضوعات تخصصی (Niche) نشان می‌دهد که مباحث مربوط به «سنجش عملکرد» و جنبه‌های فنی خود «هوش توضیح‌پذیر»، به یک حوزه کاملاً تخصصی و آکادمیک تبدیل شده است. این موضوع نشان‌دهنده آن است که پژوهشگران در حال کار بر روی لبه‌های فنی سیستم (مانند دقت توضیح‌ها و معیارهای ارزیابی هوشمندی) هستند، که اگرچه بسیار توسعه‌یافته است، اما هنوز به صورت یک استاندارد عمومی در تمامی داشبوردهای تصمیم‌گیری (که در سمت راست تصویر ۱۴ هستند) فراگیر نشده است. در موضوع‌های نوظهور یا رو به افول، موضوعاتی همچون Visual Analytics، Trust و Explainability به چشم می‌خورد. باید توجه داشت که حضور این موضوعات در این ناحیه به معنای فراموش شدن آن‌ها نیست؛ بلکه نشان‌دهنده این است که این مفاهیم از حالت «موضوعات مستقل» خارج شده و در دل سیستم‌های یادگیری عمیق حل شده‌اند. به عبارتی، توضیح‌پذیری بصری اکنون یک «پیش‌فرض» است، نه یک موضوع مستقل.

به این ترتیب، در مجموع باید اذعان داشت که در دوره اول (تصویر ۱۳)، تمرکز بر «چگونه نمایش دادن» (Visual Analytics) بود و یادگیری عمیق هنوز ضعیف بود. اما در دوره دوم (تصویر ۱۴)، تمرکز بر «هوشمندسازی عمیق» (Deep Learning) است و تعامل انسان و کامپیوتر (HCI) به زیربنای اصلی کار تبدیل شده است.

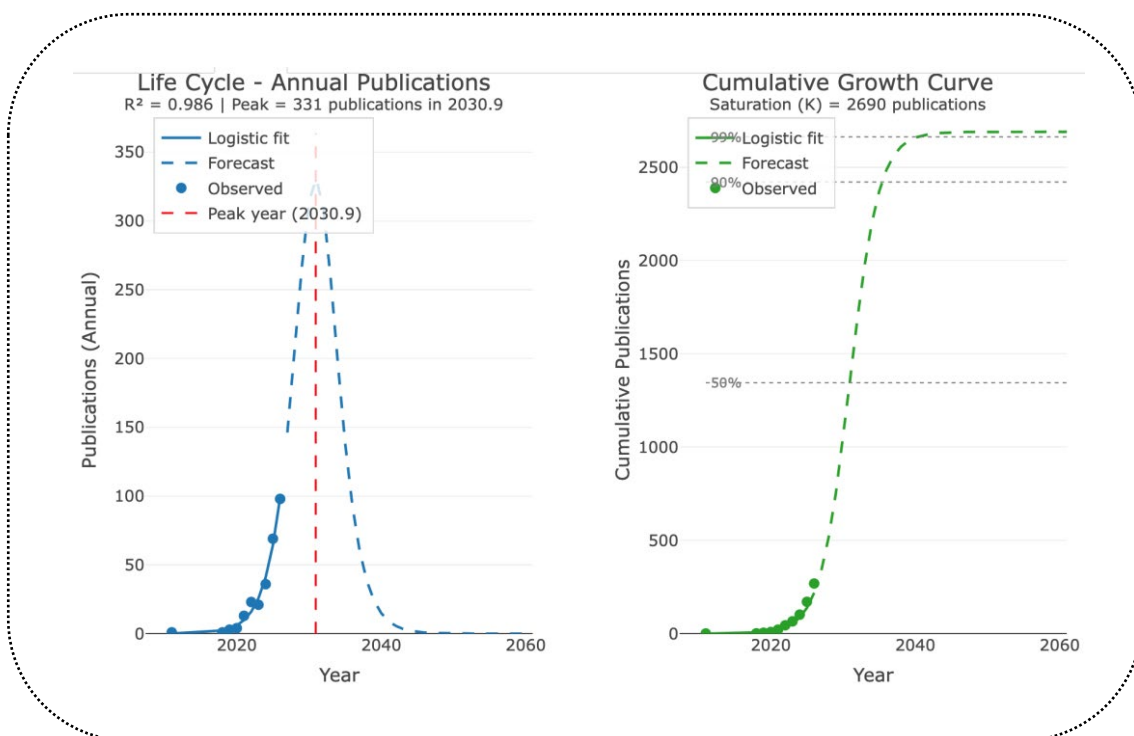


تصویر ۱۴. روند راهبردی تکامل ساختار دانش در حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، در دوره زمانی ۲۰۲۶-۲۰۱۳ م. (دوران شکل‌گیری و تمرکز بصری)

جمع‌بندی مجموع موضوعات مورد بحث در تصاویر ۱۰ تا ۱۲، در تصویر (۱۵) به نمایش درآمده است. چنان‌که ملاحظه می‌شود، موضوع‌های Human computer interaction، Machine learning، Decision making و Artificial intelligence در کانون توجهات و نقش‌آفرینی در حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری قرار دارند.



- 1 . Logistic Fit
- 2 . Life Cycle
- 3 . Peak
- 4 . Cumulative Growth
- 5 . Fast Ggrowth Phase



تصویر ۱۶. مدل سازی پیش بینی چرخه حیات و منحنی رشد تولیدات علمی در حوزه هوش مصنوعی
 توضیح پذیر در داشبوردهای تصمیم گیری

بحث و نتیجه گیری

یافته‌های پژوهش نشان داد که حوزه هوش مصنوعی توضیح پذیر در داشبوردهای تصمیم گیری، از قلمرویی صرفاً فنی و الگوریتم محور، به ساختاری فکری، چندلایه و میان رشته‌ای در حال گذار است؛ ساختاری که در آن، مفاهیمی چون اعتماد، شفافیت، تفسیرپذیری، تعامل انسان - کامپیوتر و کاربردپذیری در کنار روش‌های یادگیری عمیق و مدل‌های پیش‌بینی‌گر قرار گرفته‌اند. تحلیل هم‌رخدادی واژگان و خوشه‌بندی مفاهیم نشان داد که هسته اصلی دانش این حوزه، در پیوند میان «تصمیم‌گیری»، «هوش مصنوعی توضیح‌پذیر» و «یادگیری عمیق» شکل گرفته است؛ به گونه‌ای که این سه گانه مفهومی، مهم‌ترین کانون تولید دانش در نقشه ساختاری پژوهش محسوب می‌شود. این نتیجه نشان می‌دهد که حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، دیگر صرفاً یک افزونه تبیینی برای مدل‌های یادگیری ماشین نیست، بلکه به عنوان یک مؤلفه بنیادین در طراحی سامانه‌های تصمیم‌یار هوشمند جایگاهی تثبیت شده است.

در تبیین این یافته می‌توان گفت که تمرکز پررنگ بر «تصمیم‌گیری» بیانگر آن است که مسئله اصلی پژوهش‌های این حوزه، تنها افزایش دقت مدل‌ها نیست، بلکه ارتقای کیفیت تصمیم در شرایط پیچیده و پُرخطر نیز است. این امر با نتایج مطالعات پیشین، از جمله پژوهش‌های کوتسوپاس و نوسیوس (Koutsoupias & Nosios, 2026) و مارچیانو و همکاران (Marciano et al., 2026) هم‌راستا است؛ زیرا آن مطالعات نیز نشان داده‌اند که ارزش هوش مصنوعی توضیح‌پذیر، زمانی به طور واقعی آشکار می‌شود که تبیین‌ها در خدمت اعتماد، پذیرش کاربران و تصمیم‌گیری عملی قرار گیرند. بنابراین، یافته‌های پژوهش حاضر تأیید می‌کند که مسیر تحول حوزه هوش مصنوعی توضیح‌پذیر در

تحلیل ساختار دانشی، تکامل موضوعی و جهت‌گیری‌های آینده و نوظهور در پژوهش‌های ...

داشبوردهای تصمیم‌گیری از سطح «فناوری توضیح» به سطح «زیرساخت تصمیم‌سازی» در حال تکامل است. از سوی دیگر، نتایج تحلیل ساختار فکری نشان داد که حوزه مورد بررسی تنها بر مسائل فنی استوار نیست، بلکه در چند خوشه محتوایی متمایز، از جمله خوشه‌های انسانی - اخلاقی، فنی - روش‌شناسانه، کاربردی - عملیاتی، شناختی - خودمختار و موضوعات تخصصی و حاشیه‌ای، قابل طبقه‌بندی است. این الگو نشان می‌دهد که پژوهش‌های حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، به‌طور هم‌زمان به دو منطق پاسخ می‌دهد: نخست، منطق مهندسی و کارایی الگوریتمی؛ و دوم، منطق انسانی و اخلاقی فهم، اعتماد و مسئولیت‌پذیری. این دو وجه شناسایی شده در پژوهش‌های حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری، با یافته‌های پیشینه نیز همسو است؛ زیرا بسیاری از پژوهش‌های این حوزه تأکید کرده‌اند که هرچه نقش سیستم‌های هوشمند در تصمیم‌سازی پررنگ‌تر می‌شود، نیاز به تبیین‌پذیری، پاسخ‌گویی و کنترل انسانی نیز افزایش می‌یابد. در نتیجه، حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری را باید نه فقط یک راهکار فنی، بلکه یک سازوکار نهادی برای کاهش شکاف میان عملکرد مدل و ادراک کاربر دانست.

تحلیل نقشه چگالی و خوشه‌های پرتراکم نیز نشان داد که بیشترین تمرکز دانش در مجاورت مفاهیم «Decision Making, Explainable AI» و «Deep Learning» قرار دارد. این الگو بیانگر آن است که یادگیری عمیق به موتور اصلی پیشرفت این حوزه تبدیل شده است؛ اما به‌تنهایی کافی نیست و باید با سازوکارهای توضیح‌پذیری و طراحی رابط کاربر ادغام شود. در همین راستا، مطالعات پیشین نیز نشان داده‌اند که با پیچیده‌تر شدن مدل‌های یادگیری عمیق، مسئله توضیح‌پذیری از یک نیاز جانبی به ضرورتی مرکزی تبدیل شده است. بنابراین، یافته‌های پژوهش حاضر این ادعا را تقویت می‌کند که آینده داشبوردهای تصمیم‌گیری، نه در جایگزینی کامل مدل‌های پیچیده، بلکه در تلفیق دقت تحلیلی با شفافیت ادراکی رقم خواهد خورد.

از منظر تحول موضوعی نیز، نتایج حاکی از آن بود که این حوزه از موضوعات کلاسیک، مانند «شبکه‌های عصبی» و «ماشین بردار پشتیبان» به تدریج به سمت موضوعات جدیدتر و کاربردی‌تر همچون «یادگیری عمیق بصری»، «بصری‌سازی»، «اعتماد» و «سیستم‌های تصمیم‌یار هوشمند» حرکت کرده است. این جابه‌جایی نشان می‌دهد که تمرکز پژوهش‌ها از بهبود صرف عملکرد مدل‌ها، به سمت طراحی اکوسیستم‌هایی سوق یافته است که در آن کاربر، زمینه تصمیم و تجربه تعاملی نقش تعیین‌کننده‌ای دارند. این تحول با پژوهش‌های پیشین نیز قابل تطبیق است؛ زیرا پژوهش‌های انجام شده در حوزه هوش مصنوعی توضیح‌پذیر همواره بر این نکته تأکید داشته‌اند که موفقیت سامانه‌های توضیح‌پذیر تنها در صورتی تحقق می‌یابد که تبیین‌ها متناسب با نیاز کاربر، بافت کاربرد، و سطح دانش تصمیم‌گیرنده طراحی شوند.

یافته مهم دیگر پژوهش حاضر آن است که «اعتماد» و «بصری‌سازی» از مفاهیم پیرامونی به عناصر درونی سیستم‌های هوش مصنوعی توضیح‌پذیر تبدیل شده‌اند. این موضوع نشان می‌دهد که مطالعات اخیر دیگر صرفاً به پرسش «مدل چه می‌گوید؟» محدود نیستند، بلکه به این پرسش نیز پاسخ می‌دهند که «کاربر چگونه می‌فهمد، می‌پذیرد و بر اساس آن اقدام می‌کند؟» نیز پاسخ می‌دهند. از این منظر، نتایج پژوهش حاضر با مسیر نظری مطالعات پیشین همخوان است که توضیح‌پذیری را پلی میان هوش مصنوعی و فهم انسانی می‌دانند. به بیان دیگر، تبیین در داشبوردهای تصمیم‌گیری نه یک خروجی جانبی، بلکه بخشی از فرآیند شناختی تصمیم‌سازی است.

همچنین، برآورد مدل رشد لجستیک نشان داد که این حوزه هنوز در مرحله رشد سریع قرار دارد و با وجود

افزایش شتاب انتشار، به سقف اشباع علمی خود نرسیده است. پیش‌بینی سال اوج انتشار و ظرفیت نهایی نیز بیانگر آن است که ادبیات این حوزه در سال‌های پیش رو همچنان ظرفیت توسعه قابل‌توجهی دارد. این نتیجه از حیث سیاست‌گذاری پژوهشی اهمیت دارد؛ زیرا نشان می‌دهد که هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری نه یک موضوع تثبیت‌شده و اشباع‌شده، بلکه یک حوزه پویا و در حال بلوغ است که هنوز امکان شکل‌گیری مباحث نظری و کاربردی جدید در آن وجود دارد.

در مجموع، این پژوهش نشان داد که دانش مربوط به هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری در حال حرکت از یک پارادایم «فناوری - محور» به یک پارادایم «انسان - محور و تصمیم - محور» است. بنابراین، مهم‌ترین دستاورد این مطالعه آن است که ساختار فکری و موضوعی این حوزه را نه به صورت خطی، بلکه به صورت شبکه‌ای و چندلایه ترسیم کرده و نشان داده است که آینده آن در گرو تلفیق دقت الگوریتمی، شفافیت تبیینی، و تناسب با نیاز کاربر خواهد بود. براین‌اساس، پژوهش حاضر ضمن تأیید و امتداد یافته‌های مطالعات پیشین، شکاف میان توسعه فنی مدل‌ها و نیازهای واقعی تصمیم‌گیرندگان را برجسته می‌سازد و نشان می‌دهد که مسیر آتی حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری باید بیش از پیش به سمت طراحی توضیح‌های قابل‌فهم، زمینه‌مند و کاربردی هدایت شود.

پیشنهادهای اجرایی پژوهش

مبتنی بر یافته‌های این پژوهش، پیشنهادهای اجرایی زیر برای سیاست‌گذاران، برنامه ریزان و مدیران پژوهشی کشور قابل ارائه است:

- توسعه رویکردهای انسان‌محور در طراحی داشبوردهای مبتنی بر هوش مصنوعی توضیح‌پذیر: باتوجه‌به اینکه یافته‌های پژوهش نشان دادند مفاهیمی مانند اعتماد، شفافیت و تعامل انسان - ماشین به عناصر مرکزی این حوزه تبدیل شده‌اند، مدیران و طراحان سیستم‌های تصمیم‌یار لازم است در طراحی داشبوردها رویکرد انسان‌محور را تقویت کرده و توضیحات مدل‌های هوش مصنوعی را به گونه‌ای ارائه کنند که برای کاربران غیرمتخصص نیز قابل درک و قابل استفاده در فرآیند تصمیم‌گیری باشد.
- ادغام سازوکارهای بصری‌سازی توضیح‌پذیر در سامانه‌های تصمیم‌یار: نتایج تحلیل شبکه مفهومی نشان داد که پیوند میان «تصمیم‌گیری»، «یادگیری عمیق» و «بصری‌سازی» از کانون‌های اصلی پژوهش است. براین‌اساس، پیشنهاد می‌شود سازمان‌ها در توسعه داشبوردهای مدیریتی از ابزارهای بصری‌سازی پیشرفته برای نمایش منطق تصمیم‌گیری مدل‌ها استفاده کنند تا امکان تفسیر بهتر نتایج برای تصمیم‌گیران فراهم شود.
- حمایت از پژوهش‌های میان‌رشته‌ای در حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری و طراحی رابط کاربر: ساختار فکری حوزه نشان داد که این حوزه در مرز میان علوم داده، تعامل انسان - کامپیوتر و علوم تصمیم قرار دارد. بنابراین، سیاست‌گذاران پژوهشی می‌توانند با ایجاد برنامه‌های پژوهشی میان‌رشته‌ای و حمایت از همکاری میان متخصصان هوش مصنوعی، طراحی رابط کاربر و علوم مدیریت، به توسعه کاربردهای مؤثرتر هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری کمک کنند.
- تدوین چارچوب‌های استاندارد برای ارزیابی اعتماد و شفافیت در سیستم‌های تصمیم‌یار: باتوجه‌به برجسته بودن مفاهیمی مانند اعتماد و توضیح‌پذیری در خوشه‌های اصلی پژوهش، پیشنهاد می‌شود نهادهای سیاست‌گذار و

تحلیل ساختار دانشی، تکامل موضوعی و جهت‌گیری‌های آینده و نوظهور در پژوهش‌های ...

سازمان‌های توسعه‌دهنده سیستم‌های هوشمند، شاخص‌ها و استانداردهایی برای ارزیابی میزان شفافیت، قابلیت تفسیر و قابلیت اعتماد در داشبوردهای تصمیم‌گیری مبتنی بر هوش مصنوعی تدوین کنند.

- سرمایه‌گذاری هدفمند در حوزه‌های نوظهور و تخصصی پژوهش: نتایج نقشه ساختار مفهومی نشان داد که برخی موضوعات در قالب «موضوعات تخصصی و حاشیه‌ای» در حال شکل‌گیری هستند. برنامه‌ریزان پژوهشی می‌توانند با حمایت هدفمند از این حوزه‌های نوظهور، زمینه توسعه کاربردهای جدید حوزه هوش مصنوعی توضیح‌پذیر در داشبوردهای تصمیم‌گیری در حوزه‌هایی مانند سیستم‌های تصمیم‌یار تخصصی و کاربردهای داده‌محور را فراهم کنند.

پیشنهاد برای پژوهش‌های آتی

- باتوجه به یافته‌های این پژوهش، پیشنهادهای زیر برای پژوهش‌های آتی ارائه می‌شود:
- طراحی و اعتبارسنجی چارچوب‌های ارزیابی تجربه کاربر در داشبوردهای هوش مصنوعی توضیح‌پذیر؛
- ارزیابی تأثیر انواع روش‌های توضیح‌پذیری یادگیری عمیق بر کیفیت تصمیم‌گیری مدیران؛
- مطالعه تطبیقی تحول موضوعی هوش مصنوعی توضیح‌پذیر در حوزه‌های کاربردی مختلف (مانند سلامت، مالی، مدیریت شهری)؛
- تحلیل عمیق خوشه‌های تخصصی و حاشیه‌ای برای شناسایی روندهای نوظهور؛
- توسعه مدل‌های یکپارچه برای تلفیق بصری‌سازی، تبیین‌پذیری و تعامل کاربر در داشبوردهای هوشمند.

تقدیر و تشکر (Acknowledgement)

این مقاله برگرفته از پروژه کارشناسی نویسنده اول با عنوان بررسی مکانیزم‌های پایه ای هوش مصنوعی در سیستم‌های داشبوردهای تجاری است که تحت راهنمایی نویسنده دوم به انجام رسیده است. نویسندگان لازم می‌دانند نهایت تشکر و قدردانی خود را از داوران محترم که نظراتشان سهم به‌سزایی در بهبود بخشیدن کیفیت داشت، به عمل آورند.

تعارض منافع (Conflict of Interest)

نویسندگان اعلام می‌دارند که در خصوص انتشار این مقاله تضاد منافع وجود ندارد. علاوه بر این، موضوعات اخلاقی، از جمله سرقت ادبی، رضایت آگاهانه، سوءرفتار، جعل داده‌ها، انتشار، ارسال مجدد، مکرر و همچنین سیاست مجله در قبال استفاده از هوش مصنوعی از سوی نویسندگان رعایت شده است.

یادداشت دفتر مجله (Publisher's/ Editorial Note)

با توجه به وجود نسبت خویشاوندی میان یکی از نویسندگان و سردبیر مجله، و به‌منظور رعایت موازین کمیته بین‌المللی اخلاق نشر (COPE)، سردبیر از تمامی مراحل ارزیابی و داوری این مقاله عزل گشته و داوری اثر، ارجاع به داوران و تصمیم‌گیری نهایی در خصوص پذیرش یا رد مقاله، به صورت کامل و مستقل توسط عضو دیگری از هیئت تحریریه که هیچ وابستگی خانوادگی یا شخصی با نویسندگان ندارد انجام شده است. بدین سبب، بی‌طرفی داوری و تصمیم‌گیری این مقاله تضمین می‌شود.

فهرست منابع

- نوروزی چاکلی، ع.، نوروزی چاکلی، س.، و نوروزی، ح. (۱۴۰۳). شناسایی ابعاد، چالش‌ها، معیارها، شاخص‌ها و الزامات مطرح در چرخه داوری علمی ارزیابانه عملیاتی و ارائه چارچوبی برای داوری آثار علمی در کشور [گزارش طرح پژوهشی]. وزارت علوم، تحقیقات و فناوری، معاونت پژوهش و فناوری. بازیابی شده در اردیبهشت ۲۵، ۱۴۰۵، از <https://doi.org/10.6084/m9.figshare.32964269>
- Akkem, Y., Biswas, S. K., & Aruna, V. (2026). Deciphering the black box: Interactive crop recommendation system using Explainable AI with visualisation dashboards. *Journal of Experimental & Theoretical Artificial Intelligence*, 38(2), 219–259. <https://doi.org/10.1080/0952813X.2025.2595024>
- Ali, S., Abuhmed, T., El-Sappagh, S., Muhammad, K., Alonso-Moral, J. M., Confalonieri, R., Guidotti, R., Del Ser, J., Díaz-Rodríguez, N., & Herrera, F. (2023). Explainable Artificial Intelligence (XAI): What we know and what is left to attain Trustworthy Artificial Intelligence. *Information Fusion*, 99. <https://doi.org/10.1016/j.inffus.2023.101805>
- Alonso, J. M., Castiello, C., & Mencar, C. (2018). A bibliometric analysis of the explainable artificial intelligence research field. In J. Medina, M. Ojeda-Aciego, J. L. Verdegay, D. A. Pelta, I. P. Cabrera, B. Bouchon-Meunier, & R. R. Yager (Eds.), *Information Processing and Management of Uncertainty in Knowledge-Based Systems. Theory and Foundations. IPMU 2018. Communications in Computer and Information Science* (Vol. 853, pp. 3–15). Springer. https://doi.org/10.1007/978-3-319-91473-2_1
- Ammar, O. H., Rejeb, A., & Rejeb, K. (2026). Explainable Artificial Intelligence in finance: A bibliometric review. *Finance Research Open*, 100131. <https://doi.org/10.1016/j.finr.2026.100131>
- Angelov, P. P., Soares, E. A., Jiang, R., Arnold, N. I., & Atkinson, P. M. (2021). Explainable artificial intelligence: an analytical review. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 11(5), e1424. <https://doi.org/10.1002/widm.1424>
- Barredo, A. A., Del Ser, J., Gil-Lopez, S., Díaz-Rodríguez, N., Bennetot, A., Chatila, R., Tabik, S., Garcia, S., Molina, D., & Herrera, F. (2020). Explainable Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information Fusion*, 58, 82–115. <https://doi.org/10.1016/j.inffus.2019.12.012>
- Bornmann, L., & Leydesdorff, L. (2014). Scientometrics in a changing research landscape. *EMBO Reports*, 15(12), 1228–1232. <https://doi.org/10.15252/embr.201439608>
- Bradford, S. C. (1985). Sources of information on specific subjects. *Journal of Information Science*, 10(4), 173–180. <https://doi.org/10.1177/016555158501000406>
- Buchanan, B. G., & Shortliffe, E. H. (1984). *Rule based expert systems: the Mycin experiments of the Stanford Heuristic Programming Project (the Addison-Wesley series in artificial intelligence)*. Addison-Wesley Longman Publishing Co., Inc. <https://scispace.com/pdf/rule-based-expert-systems-the-mycin-experiments-of-the-1ohsm4se0g.pdf>

- Buñay-Guisñan, P., Lara, J. A., Cano, A., Cerezo, R., & Romero, C. (2026). Towards accessible AI for addressing students' academic dropout: An auto machine learning and explainable artificial intelligence approach. *Universal Access in the Information Society*, 25(1). <https://doi.org/10.1007/s10209-025-01278-4>
- Chen, X.-Q., Ma, C.-Q., Ren, Y.-S., Lei, Y.-T., Huynh, N. Q. A., & Narayan, S. (2023). Explainable artificial intelligence in finance: A bibliometric review. *Finance Research Letters*, 56, 104145. <https://ideas.repec.org/a/eee/finlet/v56y2023ics1544612323005172.html>
- Confalonieri, R., Coba, L., Wagner, B., & Besold, T. R. (2021). A historical perspective of explainable artificial intelligence. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 11(1), e1391. <https://doi.org/10.1002/widm.1424>
- Daovisan, H. (2026). Explainable artificial intelligence 5.0 for emerging technologies of responsible innovation. *Discover Artificial Intelligence*, 6, 595. <https://doi.org/10.1007/s44163-026-01250-y>
- Fouad, S., Hakobyan, L., Ihongbe, I. E., Kavakli-Thorne, M., Atkins, S., & Bhatia, B. (2026). Human-centered user interface design for explainable AI in chest radiology: A multi-phase co-design approach. *IEEE Access*, 14, 12498-12513. https://publications.aston.ac.uk/id/eprint/48600/3/S_Fouad_et_al_Human-Centered_User_Interface_Design_for_Explainable_AI_in_Chest_Radiology_A_Multi-Phase_Co-Design_Approach.pdf
- Guidotti, R., Monreale, A., Ruggieri, S., Turini, F., Giannotti, F., & Pedreschi, D. (2018). A survey of methods for explaining black box models. *ACM Computing Surveys (CSUR)*, 51(5), 1–42. <https://doi.org/10.1145/3236009>
- Haghani, M. (2023). What makes an informative and publication-worthy scientometric analysis of literature: A guide for authors, reviewers and editors. *Transportation Research Interdisciplinary Perspectives*, 22, 100956. <https://doi.org/10.1016/j.trip.2023.100956>
- Hunsicker, T., Schulz, A., Leist, R. A., Kiefer, S., Boden, K. T., Rothaus, K., König, C. J., & Langer, M. (2026). Efficiency Pitfalls of Explainable AI in Clinical Diagnostic and Treatment Human-AI Workflows. *Human Factors*, 00187208261443764. <https://doi.org/10.1177/00187208261443764>
- Ivancheva, L. (2008). Scientometrics today: A methodological overview. *Collnet Journal of Scientometrics and Information Management*, 2(2), 47–56. <https://doi.org/10.1080/09737766.2008.10700853>
- Kim, M., Kim, S., Kim, J., Song, T. J., & Kim, Y. (2024). Do stakeholder needs differ? - Designing stakeholder-tailored Explainable Artificial Intelligence (XAI) interfaces. *International Journal of Human-Computer Studies*, 181, 103160. <https://doi.org/10.1016/j.ijhcs.2023.103160>
- Koutsoupias, N., & Nosios, M. (2026). Explainable Artificial Intelligence for Social Sciences and Humanities: A systematic bibliometric analysis. *Engineering Proceedings*, 124(1), 113. <https://doi.org/10.3390/engproc2026124113>

- Lotka, A. J. (1926). The frequency distribution of scientific productivity. *Journal of the Washington Academy of Sciences*, 16(12), 317–323. <https://www.jstor.org/stable/24529203>
- Lotka, A. J., & Bradford, S. C. (1926). The frequency distribution of scientific productivity. *Journal of Information Science*, 10(4), 317–323. <https://www.jstor.org/stable/24529203>
- Mahanta, P., Bhattacharya, M., & Habermeier, J. (2026). Balancing AI complexity with usability: a study of AI-based UX patterns in business applications. *Information Technology & People*. <https://doi.org/10.1108/ITP-07-2025-1115>
- Marciano, J. M. V., Machado, V. P., & de Araújo, A. H. M. (2026). Bibliometric Review of the Ethical and Legal Perspectives of Explainable Artificial Intelligence in Health [Preprint]. *Research Square*. <https://doi.org/10.21203/rs.3.rs-9448591/v1>
- Miller, T. (2019). But why? Understanding explainable artificial intelligence. *XRDS: Crossroads, The ACM Magazine for Students*, 25(3), 20–25. <https://doi.org/10.1145/3313107>
- Nalela, P., Rao, D. P., & Rao, P. V. (2026). Explainable AI for predicting mortality risk in metastatic cancer: Retrospective cohort study using the memorial sloan Mettering-Metastatic Dataset. *JMIR Cancer*, 12, e74196. <https://doi.org/10.2196/74196>
- Noroozi Chakoli, A., Noroozi Chakoli, S., & Norouzi, H. (2024). *Identifying the Dimensions, Challenges, Criteria, Indicators, and Requirements of the Scholarly Peer Review Process and Proposing a Framework for the Evaluation of Scientific Works in Iran* [Research project]. Ministry of Science, Research and Technology of Iran, Deputy of Research and Technology. Retrieved May 15, 2026, from <https://doi.org/10.6084/m9.figshare.32964269> [In Persian].
- Pereira, A., & Maciel, A. M. A. (2026). Integration of an explainable dashboard to enhance autoML transparency. In A. Rocha, F. G. Penalvo, C. J. Costa, & R. Goncalves (Eds.), *Proceedings of 20th Iberian Conference on Information Systems and Technologies, CISTI 2025, Vol. 2* (Vol. 1717, Numbers 20th Iberian Conference on Information Systems and Technologies-CISTI, pp. 725–736). https://doi.org/10.1007/978-3-032-10721-3_62
- Prasad, P. W. C., Sayeed, M. S., Nguyen, D.-M., Hutabarat, D. P., & Mohiuddin, G. M. (2026). Explainable AI: Enhancing decision-making in the detection of cyber threats. *Frontiers in Computer Science*, 8, 1762332. <https://doi.org/10.3389/fcomp.2026.1762332>
- Rejeb, A., Rejeb, K., & Treiblmaier, H. (2026). Explainable artificial intelligence in finance: a bibliometric and topic modeling analysis using BERTopic. *Quality & Quantity*, 1–25. <https://doi.org/10.1007/s11135-026-02837-4>
- Rezaeian, O., Bayrak, A. E., & Asan, O. (2026). Explainability and AI confidence in clinical decision support systems: Effects on trust, diagnostic performance, and cognitive load in breast cancer care. *International Journal of Human–Computer Interaction*, 42(6), 4477–4497. <https://doi.org/10.1080/10447318.2025.2539458>
- Russo, M., & Vistocco, D. (2026). Explainable Artificial Intelligence through the lens of bibliometric citation analysis. *Applied Stochastic Models in Business and Industry*, 42(3), e70091. <https://doi.org/10.1002/asmb.70091>

- Shiddik, M. A. B. (2026). Explainable Artificial Intelligence in healthcare: current landscape, challenges, and future directions. *Health Science Reports*, 9(3), e72172. <https://doi.org/10.1002/hsr2.72172>
- Talmoudi, R., & Choukir, J. (2026). From predictive analytics to Explainable AI in higher education: A bibliometric mapping. *Qubahan Academic Journal*, 6(2), 437–461. <https://doi.org/10.48161/qaj.v6n2a2559>
- Torbati, A. S., & Noroozi Chakoli, A. (2013). Empirical examination of Lotka's law for applied mathematics. *Life Science Journal*, 10(SUPPL. 5). https://www.lifesciencesite.com/lmj/life1005s/104_17439life1005s_601_607.pdf
- Tveita, L. J., & Hustad, E. (2025). Benefits and challenges of Artificial Intelligence in public sector: A literature review. *Procedia Computer Science*, 256, 222–229. <https://doi.org/10.1016/j.procs.2025.02.115>
- Van Leeuwen, T. (2006). The application of bibliometric analyses in the evaluation of social science research. Who benefits from it, and why it is still feasible. *Scientometrics*, 66(1), 133–154. <https://doi.org/10.1007/s11192-006-0010-7>