



Research Article

Nonlinear Modeling Of The Interaction Of Ethical, Cultural, And Structural Factors On The Intention To Disclose Sensitive Knowledge Using Response Surface Methodology: A Human Versus AI Comparison

AliAsghar Abdeslahi^{1*}  Hojat RezaeiArjman²  Shervin Vahidian Qazvini³ 

1. Ph.D. Candidate in Public Administration, Department of Management, Faculty of Management and Economics, Lorestan University, Khorramabad, Iran. E-mail: aliasghar.abdeslahi@gmail.com

2. Department of Business Administration, Faculty of Management, Islamic Azad University, Arak, Iran.
E-mail: hrezaeiarmand56@gmail.com

3. Master's Student in Applied Physics, Department of Applied Physics, Faculty of Physics and Astronomy, University of Bologna, Bologna, Italy. E-mail: Shervin.Vahidian@studio.unibo.it

Received: 20 May 2025; Revised: 3 August 2025 ; Accepted: 23 September 2025; Published: 30 September 2025

Abstract

Purpose: Public sector organizations today operate under intensifying demands to uphold ethical standards, mitigate corruption, and ensure public trust. Whistleblowing—the deliberate, principled disclosure of sensitive internal information such as financial irregularities, misconduct, or policy violations—serves as a vital corrective tool. However, the decision to blow the whistle is not driven by a single stimulus; instead, it emerges from a complex nexus of factors. An individual's moral inclination, shaped by personal values and integrity, provides the foundational readiness to report wrongdoing. Simultaneously, the prevailing ethical knowledge-sharing culture within an organization either encourages or stifles such disclosures by signaling whether employees will be supported or ostracized. Finally, structural barriers—bureaucratic hurdles, opaque reporting procedures, fear of retaliation, and lack of legal safeguards—directly impede the act of disclosure. While prior research has examined these dimensions in isolation and largely through linear models, the present study addresses a critical gap by modeling their joint, potentially non-linear influences on whistleblowing intention. Moreover, the advent of advanced language models, such as Grok-3, affords a unique opportunity to compare algorithmic predictions against actual human behavior. Accordingly, this research aims to (1) develop a comprehensive, non-linear quantitative model capturing the combined effects of moral inclination, ethical culture, and structural barriers on whistleblowing intention among public service employees in Lorestan Province, Iran, and (2) evaluate the comparative predictive accuracy and limitations of human survey data versus Grok-3's outputs under identical vignette scenarios.

Methodology: An applied, descriptive-experimental field study was undertaken, employing Central Composite Design (CCD) in conjunction with Response Surface Methodology (RSM) to explore main, interaction, and quadratic effects. From a population of 700 service-sector employees, a stratified random sample of 385 was determined via Cochran's formula, ensuring sufficient statistical power. Data collection used a scenario-based questionnaire featuring fifteen unique vignettes, each systematically varying three independent variables—individual moral inclination (A), ethical knowledge-sharing culture (C), and structural barriers (B)—across low, medium, and high levels. Participants rated their likelihood to disclose sensitive knowledge on a seven-point Likert scale (1 = “definitely would not” to 7 = “definitely would”). Content validity was confirmed through expert review by five scholars in knowledge management and organizational behavior, and internal consistency reliability was established with Cronbach's alpha = 0.88 in a pilot test of 30 employees. In parallel, the same vignette scenarios were input into Grok-3 to generate AI-based intention scores. Four regression frameworks—purely linear, two-way interaction, quadratic (second degree), and cubic (third degree)—were fitted for both human and AI datasets using Design-Expert 13. Model selection criteria included coefficient of determination (R^2), adjusted R^2 , ANOVA F-tests, lack-of-fit tests, and sum of squared errors (SSE). The quadratic model demonstrated superior explanatory power ($R^2 = 0.95$ human; $R^2 = 0.97$ AI) with non-significant lack-of-fit and was selected for detailed analysis. A composite utility function was then applied to pinpoint the optimal factor combination that maximizes whistleblowing intention.

Findings: The human-based quadratic model explained 95% of variance in whistleblowing intention, with highly significant coefficients. Moral inclination (A) exhibited the most pronounced positive linear effect (coefficient = 0.96,

$p < 0.01$), highlighting its central motivational role. Its accompanying negative quadratic term ($-0.76, p < 0.05$) revealed a saturation threshold beyond which additional moral motivation produced diminishing increases in disclosure intent—an effect seldom captured in linear analyses. Ethical knowledge-sharing culture (C) registered a robust positive linear coefficient of 0.95 ($p < 0.01$) across all levels, with no evidence of saturation, signifying its consistent enabling function. Structural barriers (B) exerted a significant negative linear effect ($-0.59, p < 0.01$), indicating that each incremental barrier unit steadily suppresses willingness to report. Under the optimal conditions ($A = 7, C = 7, B = 1$), predicted human intention reached 5.55 on the seven-point scale with a composite utility of 0.87. Grok-3's quadratic model paralleled these trends but with distinct magnitudes: A's linear coefficient was 1.64 ($p < 0.01$), $B = -0.60$ ($p < 0.01$), and $C = 0.80$ ($p < 0.01$). Notably, AI identified significant interactions: $A \times B$ ($-0.375, p < 0.05$) suggested that high moral drive combined with strong barriers markedly dampens intention, while $A \times C$ ($0.225, p < 0.05$) signified synergistic gains when moral drive aligns with a supportive culture. The AI quadratic A^2 term ($-0.43, p < 0.05$) reaffirmed saturation in moral motivation. Grok-3's optimal predicted intention soared to 7.00, approximately 26% higher than human respondents, reflecting AI's lack of psychosocial risk aversion and highlighting the complexity of real-world disclosure decisions.

Research limitations/implications: The cross-sectional vignette approach, though experimentally robust, cannot fully emulate the emotional stakes, group dynamics, and organizational politics of authentic whistleblowing. The geographic focus on a single province reduces the generalizability to distinct cultural, regulatory, or institutional contexts. Self-report measures risk social desirability bias and cognitive fatigue, especially over multiple scenario evaluations. Grok-3's opaque proprietary architecture precludes in-depth understanding of how it weights and processes scenario information. Future research should incorporate longitudinal designs tracking actual disclosure behaviors, broaden samples across diverse settings, integrate objective whistleblowing records, and explore moderating influences such as employees' perceived organizational support and individual risk tolerance.

Practical implications: This study furnishes a multi-lever, evidence-based roadmap for public sector management. First, cultivate intrinsic moral motivation via scenario-based ethics training, peer support networks, and visible ethical leadership—but calibrate intensity to avoid motivational saturation. Second, strengthen an ethical knowledge-sharing culture by embedding transparency in organizational mission statements, celebrating exemplary whistleblowers, and maintaining open communication channels for ethical concerns. Third, dismantle structural barriers through simplified, confidential reporting procedures, robust anti-retaliation policies, and clear legal safeguards. While AI models like Grok-3 can support scenario planning and policy simulations, they must complement—rather than replace—human judgment, particularly where psychosocial and cultural nuances dictate employee risk calculations.

Originality/value: This research represents the first Iranian application of an integrated CCD-RSM experimental design combined with a cutting-edge large language model to analyze whistleblowing intentions. Departing from conventional linear regression, it uncovers novel saturation dynamics in moral motivation and maps intricate factor interactions. The introduction of a utility-optimization metric offers actionable guidelines for calibrating moral, cultural, and structural interventions. By juxtaposing human survey data with AI predictions, the study elucidates both the promise and the limitations of AI in ethically sensitive organizational contexts, thereby advancing methodological frontiers in the study of organizational ethics and public governance.

Keywords: Artificial Intelligence (AI), Knowledge Management, Organizational Behavior, Organizational Structure, Response Surface Methodology, Sensitive Knowledge Disclosure

Cite this article: AliAsghar Abdeslahi1, Hojat RezaeiArjman, Shervin Vahidian Qazvini. (2025). Nonlinear Modeling Of The Interaction Of Ethical, Cultural, And Structural Factors On The Intention To Disclose Sensitive Knowledge Using Response Surface Methodology: A Human Versus AI Comparison. Strategic Management of Organizational Knowledge, 8 (3), 89-113. <https://doi.org/10.47176/SMOK.2025.1929>

© 2024 The Authors. Strategic Management of Organizational Knowledge published by Imam Hussein University. This is an open-access article under the CC-BY 4.0 license. (<https://creativecommons.org/licenses/by/4.0/>)

Funding

None.

Author contributions

AliAsghar Abdeslahi: Conceptualization, Project Administration, Methodology, Formal Analysis, Data Curation, Statistical Computation, Writing – Original Draft.

Hojat RezaeiArjman: Literature Review, Instrument Design, Investigation, Writing – Review & Editing.

Shervin Vahidian Qazvini: Selection, experimentation, and interpretation of results derived from artificial intelligence (AI).

Conflicts of interest

The author declares no conflict of interest.

Acknowledgments

None.



مدل سازی غیر خطی تعامل عوامل اخلاقی، فرهنگی و ساختاری بر نیت افشای دانش حساس با روش سطح پاسخ: مقایسه انسان و هوش مصنوعی

علی اصغر عبدشاهی^{۱*}، حجت رضایی ارجمند^۲، شروین وحیدیان قزوینی^۳

۱. دانشجوی دکتری مدیریت دولتی، گروه مدیریت، دانشکده مدیریت و اقتصاد، دانشگاه لرستان، خرم‌آباد، ایران. E-mail: aliasghar.abdeshahi@gmail.com

۲. گروه مدیریت بازرگانی، دانشکده مدیریت، دانشگاه آزاد اسلامی، اراک، ایران. E-mail: hrezaeiarjmand56@gmail.com

۳. دانشجوی کارشناسی ارشد فیزیک کاربردی، گروه فیزیک کاربردی، دانشکده فیزیک و کیهان‌شناسی، دانشگاه بلونیا، بلونیا، ایتالیا. Email: Shervin.Vahidian@studio.unibo.it

تاریخ دریافت: ۳۰ اردیبهشت ۱۴۰۴؛ تاریخ بازنگری: ۱۲ مرداد ۱۴۰۴؛ تاریخ پذیرش: ۱ مهر ۱۴۰۴؛ تاریخ انتشار: ۸ مهر ۱۴۰۴

چکیده

هدف: با توجه به افزایش اهمیت شفافیت، پاسخ‌گویی و مقابله با فساد در سازمان‌های دولتی، افشای آگاهانه دانش حساس به‌عنوان یک رفتار داوطلبانه و اخلاقی‌مدار در بطن خط‌مشی‌گذاری عمومی اهمیت یافته است. این پژوهش با هدف مدل‌سازی تأثیر توأمان و غیرخطی سه متغیر «تمایل اخلاقی فردی»، «فرهنگ اشتراک دانش اخلاقی» و «موانع ساختاری اشتراک دانش» بر «نیت افشای دانش حساس» و مقایسه الگوی انسانی با مدل هوش مصنوعی گروک-۳ در بین کارکنان ادارات خدماتی استان لرستان انجام شد.

روش پژوهش: تحقیق حاضر از نوع کاربردی-توصیفی و آزمایشی است و با بهره‌گیری از طراحی مرکب مرکزی و روش سطح پاسخ انجام گرفت. جامعه آماری شامل ۷۰۰ نفر از کارکنان ادارات خدماتی استان لرستان و حجم نمونه ۳۸۵ نفر بود که با روش نمونه‌گیری تصادفی طبقه‌ای انتخاب شدند. ابزار گردآوری داده‌ها پرسش‌نامه سناریو محور با ۱۵ ترکیب سه‌سطحی بر اساس و مقیاس لیکرت ۷ درجه‌ای برای سنجش نیت افشای دانش بود. روایی محتوایی با نظر پنج متخصص و پایایی با ضریب آلفای کرونباخ ۰/۸۶ تأیید شد. داده‌ها با استفاده از نرم‌افزار دیزاین اکسپرت نسخه ۱۳ و برازش مدل درجه دوم تحلیل شدند.

یافته‌ها: نتایج نشان داد هر سه متغیر مستقل در سطح ۱ درصد اثر معناداری بر نیت افشا دارند. ضریب مربعی «تمایل اخلاقی فردی» منفی و معنادار بود که نشان‌دهنده اشباع اثر آن در سطوح بالا است، در حالی که «فرهنگ اشتراک دانش اخلاقی» و «موانع ساختاری» اثری خطی و مستقیم داشتند. مدل هوش مصنوعی نیز روند مشابهی را نشان داد ولی ضریب خطی «تمایل اخلاقی» را بیش از دو برابر تخمین زد. در شرایط بهینه (تمایل=۷، فرهنگ=۷، موانع=۱)، نیت افشا در مدل انسانی برابر ۵/۵۵ و در مدل هوش مصنوعی برابر ۷ محاسبه شد که حدود ۲۶ درصد تفاوت را نشان می‌دهد.

نتیجه‌گیری: نیت افشای دانش حساس متأثر از ترکیب اخلاق فردی، زمینه فرهنگی و موانع ساختاری است و در شرایط بهینه نیز به‌طور کامل تحقق نمی‌یابد. یافته‌ها بیانگر آن است که هوش مصنوعی در شناسایی روندهای کلی عملکرد قابل قبولی دارد، اما نمی‌تواند پیچیدگی‌های روان‌شناختی و اجتماعی کارکنان را به‌طور کامل بازنمایی کند.

اصالت/ارزش: این پژوهش نخستین مطالعه ایرانی است که با استفاده از ترکیب روش سطح پاسخ و مدل زبانی گروک-۳ به تحلیل غیرخطی و مقایسه انسان-هوش مصنوعی در حوزه افشای دانش حساس می‌پردازد. استفاده هم‌زمان از سناریوهای طراحی‌شده، مدلسازی سطح پاسخ و الگوریتم‌های زبانی پیشرفته، مرزهای رایج روش‌شناسی در تحلیل‌های رفتاری را توسعه می‌دهد و زمینه تدوین خط‌مشی‌های هوشمندانه‌تری را برای ارتقاء شفافیت سازمانی فراهم می‌سازد.

کلیدواژه‌ها: افشای دانش حساس، رفتار سازمانی، ساختار سازمانی، روش سطح پاسخ، مدیریت دانش، هوش مصنوعی.

مقدمه و بیان مسئله

با گسترش روزافزون فناوری‌های دیجیتال و افزایش پیچیدگی فرایندهای اداری و دولتی، افشای آگاهانه دانش حساس توسط کارکنان به‌عنوان یکی از سازوکارهای کلیدی برای کاهش تقارن اطلاعاتی، تقویت شفافیت نهادی و افزایش مسئولیت‌پذیری شناخته شده است (Kang, 2023; Zia et al., 2021). این پدیده که در ادبیات مدیریت دانش با عنوان «افشاگری دانشی»^۱ مطرح می‌شود، با تسهیل جریان هدفمند و امن اطلاعات حساس میان سطوح و بخش‌های مختلف سازمان، به بهبود کیفیت تصمیم‌گیری، ارتقای پاسخگویی و افزایش کارایی سازمان کمک می‌کند (Nicholls et al., 2021).

یکی از اصلی‌ترین پیش‌بینی‌کنندگان نیت افشاگری، انگیزه اخلاقی درونی یا «تمایل خدمت عمومی»^۲ کارکنان است؛ به‌طوری که هرچه کارکنان تعهد اخلاقی و حس مسئولیت بالاتری نسبت به منافع عمومی داشته باشند، احتمال افشای ناهنجاری‌ها افزایش می‌یابد (Potipiroon & Wongpreedee, 2020). علاوه بر این، انتظارات کارکنان در خصوص دریافت پاداش‌های حمایتی، تشویق‌های سازمانی یا برخورداری از حمایت‌های قانونی کافی و موثر از افشاگران می‌تواند این انگیزه اخلاقی را در مسیر عمل تقویت کرده یا در صورت نبود سازوکارهای حمایتی کارآمد، آن را تضعیف کند (Potipiroon, 2024).

با وجود انگیزه‌های اخلاقی قوی، کارکنان به دلیل «ترس از تلافی»^۳، ابهام در سازوکارهای رسمی گزارش‌دهی و عدم اعتماد به حمایت‌های نهادی، از به‌اشتراک‌گذاری دانش حساس خودداری می‌کنند (Cheliatsidou et al., 2023). بررسی‌های نظام‌مند حوزه مدیریت دانش و رفتار سازمانی نشان می‌دهد که نبود سیاست‌های حمایتی جامع، احساس ناامنی روانی، و فرهنگ سازمانی غیراخلاقی از مهم‌ترین عوامل مؤثر در ایجاد سکوت سازمانی و خودداری از افشاگری به شمار می‌روند (Nicholls et al., 2021). «فرهنگ اشتراک دانش اخلاقی»^۴ به معنای وجود اقلیمی سازمانی است که اشتراک دانش حساس را تشویق و خطا را فرصتی برای یادگیری تلقی می‌کند. شواهد نشان می‌دهد سازمان‌هایی که از رسانه‌های اجتماعی داخلی برای تعامل و انتقال دانش استفاده می‌کنند، فضایی امن برای افشای اطلاعات حساس فراهم می‌آورند (Razmerita et al., 2016; Zia et al., 2021).

در این پژوهش، علاوه بر پاسخ‌های انسانی گردآوری‌شده از پرسشنامه وینیته‌محور^۵، پاسخ‌های مدل زبانی «گروک-۳»^۶ نیز مورد مقایسه قرار می‌گیرد. گروک-۳ با قابلیت استدلال پیشرفته و دسترسی به موتور جستجوی «دیپ‌سرچ»^۷، سناریوهای اخلاقی پیچیده را مشابه انسان پردازش می‌کند (XAI, 2025; Reuters, 2025; Financial Times, 2025). هوش مصنوعی^۸ با ایجاد تحولات بنیادین در فرآیندهای تصمیم‌گیری راهبردی و خودکارسازی فرایندها، نقش مهمی در شکل‌دهی آینده مدیریت در سطح جهانی ایفا خواهد کرد (Nazari Zadeh et al., 2023).

با گسترش فساد اداری، ضعف نظام‌های گزارش‌دهی، و کاهش سرمایه اجتماعی در ادارات خدماتی استان لرستان، افشای دانش حساس توسط کارکنان به ضرورتی راهبردی تبدیل شده است. بررسی‌های میدانی و مطالعات پیشین (Karimi, Ghaderi & Moradinejad, 2020)، نشان می‌دهند که شهروندان لرستان فساد را بالا ارزیابی کرده و ساختارهای موجود پاسخ‌گوی نیاز به شفافیت نیستند. بااین‌حال، ترس از تلافی، نبود رویه‌های امن گزارش‌دهی، و فرهنگ سازمانی غیراخلاقی مانع افشاگری می‌شوند. این پژوهش، با بهره‌گیری از طراحی مرکب مرکزی و روش سطح^۹ پاسخ، تأثیر توأمان تمایل اخلاقی فردی، فرهنگ اشتراک دانش اخلاقی، و موانع ساختاری را بر نیت افشای دانش حساس بررسی می‌کند. برای نخستین‌بار در ایران، تحلیل انسانی با خروجی‌های مدل هوش مصنوعی «گروک-۳» مقایسه می‌شود؛ رویکردی که به دلیل نبود مطالعات مشابه در ایران، نوآوری روش‌شناختی محسوب می‌شود. این مقایسه نه‌تنها به غنای نظری در فهم رفتارهای اخلاقی سازمانی کمک می‌کند، بلکه با شناسایی دقیق موانع و پیشران‌های افشاگری، به سیاست‌گذاران امکان می‌دهد تا رویه‌های گزارش‌دهی امن، فرهنگ سازمانی اخلاقی، و سیستم‌های خودکار مبتنی بر هوش مصنوعی را طراحی کنند. پرسش اصلی این است که ترکیب و تعامل این سه عامل چگونه نیت افشاگری را شکل می‌دهد و آیا درک انسان و هوش مصنوعی از شرایط اخلاقی همسو یا متفاوت است؟

1. Knowledge Disclosure
2. Desire For Public Service
3. Fear Of Retaliation
4. An Ethical Culture Of Knowledge Sharing
5. Vignette-Based Questionnaire
6. Grok-3
7. Deep Search
8. Artificial Intelligence (AI)
9. Response Surface Methodology (RSM)

ادبیات نظری

تغییرات فناورانه و تحولات مدیریت دانش

از آغاز قرن بیست و یکم، تغییرات شتابان فناورانه، توسعه زیرساخت‌های دیجیتال و گسترش شبکه‌های ارتباطی در سطح جهانی موجب دگرگونی اساسی در الگوهای خلق، ذخیره‌سازی، پردازش و انتقال دانش در سازمان‌ها شده است، به‌طوری که امروزه مدیریت دانش به‌عنوان یکی از مؤلفه‌های حیاتی رقابت‌پذیری و کارایی سازمان‌ها شناخته می‌شود (Laihonen, Kork, & Sinervo, 2023). مدیریت دانش به‌عنوان حوزه‌ای میان‌رشته‌ای، مجموعه‌ای از سازوکارها و فرایندها را برای شناسایی، ثبت، تسهیم و به‌کارگیری دانش در سطوح فردی، گروهی و سازمانی بررسی می‌کند؛ اهمیتی که این روزها به‌ویژه در سازمان‌های خدماتی عمومی، به دلیل ماهیت دانش‌محور فرایندهای تصمیم‌گیری، ارائه خدمات به شهروندان و پاسخ‌گویی شفاف، برجسته‌تر و پررنگ‌تر شده است (Feng, Huang, Tian, & Wang, 2025).

دانش حساس و مفهوم افشاگری

در این میان، «دانش حساس»^۱ به دانشی اطلاق می‌شود که افشای آن می‌تواند تبعات شخصی، سازمانی یا اجتماعی به‌دنبال داشته باشد؛ مواردی مانند خطاهای مدیریتی یا سوءجریان‌های مالی، که معمولاً کارکنان را در مواجهه با ترس، تردید یا فشار اجتماعی قرار می‌دهد (Nicholls et al., 2021). «افشای دانش» رفتاری آگاهانه و هدفمند است که طی آن کارکنان، اطلاعات حساس و گاه بحرانی را درون سازمان یا برای مراجع بیرونی گزارش می‌کنند تا از بروز یا تداوم تخلفات احتمالی جلوگیری شود یا شرایط بهبود یابد (Potipiroon, 2024)، و «نیت افشاگری»^۲ نشان‌دهنده نیت ذهنی یا تمایل عملی فرد به این اقدام است که تحت تأثیر عوامل اخلاقی، شناختی، فرهنگی و ساختاری قرار دارد (Potipiroon & Wongpreedee, 2020).

افشای دانش حساس در ادبیات مدیریت دانش

در ادبیات مدیریت دانش، افشای دانش حساس را می‌توان نوعی انتقال کنترل‌شده‌ی دانش بحرانی^۳ تلقی کرد که مستلزم اعتماد، امنیت روانی، و وجود زیرساخت‌های سازمانی حامی است. به‌ویژه در سازمان‌های دولتی که دانش نه‌تنها ابزار عمل، بلکه ابزار قدرت نیز تلقی می‌شود، افراد در مواجهه با دانش حساس دچار تعارضی میان الزامات اخلاقی و مخاطرات ساختاری می‌شوند. افشای آگاهانه‌ی دانش بحرانی زمانی رخ می‌دهد که سازمان محیطی ایمن، شفاف و ارزش‌محور برای انتقال دانش فراهم کند؛ در غیر این صورت، دانش ضمنی و حساس در سطح فردی باقی می‌ماند و به سرمایه جمعی تبدیل نمی‌شود (Dalkir, 2011).

ابعاد فردی و ساختاری در رفتار افشاگرانه

تمایل اخلاقی فردی یا انگیزش درونی کارکنان برای حفظ منافع عمومی، از طریق سه مولفه اصلی جلوه می‌کند: رهبری نوع‌دوستانه^۴، رفاقت سازمانی^۵، و عواطف مثبت^۶ ناشی از تعاملات حمایتی که احساس امنیت را در افشاگر بالقوه تقویت می‌کند (He & Wei, 2022). موانع ادراک‌شده^۷، موانعی ذهنی-روان‌شناختی هستند که فرد آن‌ها را هزینه‌آور یا خطرناک می‌پندارد و شامل پنهان‌کاری تدافعی^۸، نقش «احمق فرض کردن»^۹ و موجه‌سازی افشای نکردن اطلاعات است (Wen, Zheng, & Ma, 2021). فرهنگ اشتراک دانش اخلاقی، آن اقلیم سازمانی^{۱۰} است که از طریق سه عنصر در

1. Sensitive Knowledge
2. Intention To Disclose
3. Critical Knowledge
4. Altruistic leadership
5. Organizational Camaraderie
6. Positive Emotions
7. Perceived Barriers
8. Defensive Secrecy
9. To Pretend To Be Stupid
10. Organizational Climate

دسترس‌پذیری اطلاعات^۱، رضایت شغلی کارکنان^۲ و پیکربندی زمینه سازمانی^۳، اشتراک مسئولانه و شفاف دانش حساس را تشویق می‌کند (Fischer & Döring, 2022; Feng et al., 2025).

بسترهای فناورانه و نقش هوش مصنوعی

علاوه بر این متغیرها، قابلیت تحلیل کلان‌داده با بهره‌گیری از زیرساخت‌های پیشرفته پردازشی، امکان اکتشاف الگوهای پنهان و بهره‌برداری از داده‌های بزرگ و پیچیده را فراهم می‌سازد و سازمان‌ها را قادر می‌کند تا در فرآیند تصمیم‌گیری، با شناسایی سریع‌تر نشانه‌های اولیه فساد، اقدام‌های اصلاحی مناسب را طراحی و دانش حساس مرتبط را در زمان مناسب و با حداقل ریسک افشا کنند (Wang, Chen, Wang, Ma, & Pan, 2024). همچنین، فرایندهای شکل‌دهی و مدیریت دانش در بخش عمومی باید در بستر مشارکتی، شبکه‌ای و میان‌سازمانی طراحی شوند تا از طریق تسهیل جریان آزاد اطلاعات و ادغام دانش‌های ضمنی و صریح کارکنان و مدیران، هم‌افزایی دانشی تحقق یابد و ظرفیت یادگیری سازمانی برای مواجهه با چالش‌های پیچیده و چندوجهی ارتقا یابد (Laihonen et al., 2023).

با توجه به پیشرفت‌های فناوری، بهره‌گیری از مدل زبانی «گروک ۳» به‌عنوان ابزاری مکمل در تحلیل نیت افشاگری و الگوهای تصمیم‌گیری اخلاقی مورد توجه قرار گرفته است. این مدل با معماری عامل‌محور و دسترسی به موتور جستجوی لحظه‌ای «دیپ‌سرچ» امکان ترکیب داده‌های متنی، اطلاعات تکمیلی و نتایج جستجوی زنده را فراهم می‌سازد. گروک-۳ با استفاده از شبکه عصبی عمیق زبانی و معماری عامل‌محور، سناریوهای اخلاقی پیچیده را مشابه انسان پردازش می‌کند (XAI, 2025; Reuters, 2025)، با این حال، در این پژوهش تمرکز اصلی بر طراحی یا تحلیل فنی مدل نبوده، بلکه هدف اصلی مقایسه خروجی‌های تحلیلی گروک-۳ با پاسخ‌های انسانی در شرایط اداری یکسان است تا میزان درک مدل از رفتار اخلاقی واقعی کارکنان سنجیده شود. بنابراین، تفاوت‌های مشاهده‌شده در نتایج می‌تواند ناشی از ساختار داخلی مدل یا ویژگی‌های داده‌ای باشد که این مقایسه آن را روشن می‌سازد.

خلا نظری و ضرورت پژوهش حاضر

با وجود تحقیقات گسترده در زمینه هر یک از متغیرها، هنوز بررسی همزمان تعامل غیرخطی «تمایل اخلاقی فردی»، «موانع ادراک‌شده» و «فرهنگ اشتراک دانش اخلاقی» در یک چارچوب واحد و نقش آن‌ها بر افشای دانش^۴ حساس، همراه با مقایسه پاسخ‌های انسانی و خروجی‌های گروک ۳، به‌طور جامع انجام نشده است. این خلا نظری مبنای پژوهش حاضر را تشکیل می‌دهد. در عین حال، با توجه به پررنگ شدن نقش هوش مصنوعی در امور روزانه و تعاملات سازمانی، نیاز به ارزیابی میزان فهم مشترک میان انسان‌ها و مدل‌های پیشرفته‌ای مانند گروک-۳ بیش از پیش احساس می‌شود. این ضرورت از آن جهت اهمیت دارد که هوش مصنوعی، به‌ویژه در زمینه‌هایی مانند تصمیم‌گیری اخلاقی^۵ و مدیریت دانش^۶، باید توانایی درک و شبیه‌سازی رفتارها و نگرش‌های انسانی را داشته باشد تا بتواند به‌طور مؤثری در محیط‌های پیچیده و انسانی‌محور مانند ادارات خدماتی عمل کند. مقایسه پاسخ‌های انسانی با خروجی‌های گروک ۳ در این پژوهش، نه‌تنها به شناسایی شکاف‌های موجود در درک متقابل کمک می‌کند، بلکه می‌تواند راهنمایی برای بهبود طراحی مدل‌های هوش مصنوعی در راستای هم‌افزایی بیشتر با انسان‌ها ارائه دهد.

پیشینه پژوهش

با توجه به اهمیت افشای دانش حساس در ارتقای شفافیت و کاهش فساد اداری، بررسی پیشینه نظری و تجربی این حوزه ضروری است. در این راستا، مرور پیشینه پژوهش با روش تحلیل محتوای هدفمند و جستجوی نظام‌مند در پایگاه‌های معتبر داخلی و خارجی انجام شد. برای این منظور، از کلیدواژه‌هایی چون «افشا دانش حساس»، «اخلاق سازمانی»، «موانع اشتراک دانش» و «اقلیم یادگیری» استفاده گردید. هدف از ارائه جدول (۱)، دسته‌بندی مهم‌ترین مطالعات مرتبط و نمایش ساختار کلی حوزه تحقیق است تا زمینه تحلیل دقیق‌تر نقاط قوت و خلأهای موجود فراهم شود.

1. Information Accessibility
2. Employee Job Satisfaction
3. Configuring The Organizational Context
4. Knowledge Disclosure
5. Ethical Decision-Making
6. Knowledge Management

جدول ۱. پیشینه پژوهش

ردیف	نویسنده/ نویسندگان (سال پژوهش)	عنوان پژوهش	مهم‌ترین یافته‌ها و نتایج مرتبط با پژوهش
۱	Laihonen, Kork & Sinervo (2023)	پیشبرد مدیریت دانش در بخش عمومی: به‌سوی درکی از شکل‌گیری دانش در اداره امور عمومی	تأکید بر اهمیت مدیریت دانش در بخش عمومی و نقش آن در بهبود تصمیم‌گیری و ارائه خدمات، که با موضوع افشای دانش حساس در سازمان‌های دولتی مرتبط است.
۲	Wen, Zheng & Ma (2021)	پیش‌زمینه‌های پنهان‌سازی دانش و تأثیر آن‌ها بر عملکرد سازمانی	شناسایی عوامل مؤثر بر پنهان‌سازی دانش و تأثیر آن بر عملکرد سازمانی، که می‌تواند بر نیت افشای دانش حساس تأثیرگذار باشد.
۳	Feng, Huang, Tian & Wang (2025)	تأثیرات پیکربندی‌شده عوامل زمینه‌ای سازمان بر اشتراک دانش در توسعه محصولات پیچیده	تحلیل تأثیر عوامل زمینه‌ای سازمانی بر اشتراک دانش در توسعه سیستم‌های پیچیده، که می‌تواند بر فرهنگ اشتراک دانش اخلاقی تأثیرگذار باشد.
۴	Wang, Chen, Wang, Ma & Pan (2024)	قابلیت تحلیل کلان‌داده و نوآوری اجتماعی: نقش میانجی اکتشاف و بهره‌برداری از دانش	بررسی نقش قابلیت تحلیل کلان‌داده در نوآوری اجتماعی و میانجی‌گری اکتشاف و بهره‌برداری از دانش، که می‌تواند بر نیت افشای دانش حساس تأثیرگذار باشد.
۵	He & Wei (2022)	پیشگیری از رفتارهای پنهان‌سازی دانش از طریق دوستی در محیط کار و رهبری نوع‌دوستانه: نقش میانجی هیجانات مثبت	یافته‌ها نشان می‌دهد که دوستی در محیط کار و رهبری نوع‌دوستانه با میانجی‌گری احساسات مثبت می‌تواند از پنهان‌سازی دانش جلوگیری کند، که مرتبط با نیت افشای دانش حساس است.
۶	Fischer & Döring (2022)	از به اشتراک گذاشتن اطلاعات متشکرم! چگونه اشتراک‌گذاری دانش و دسترسی به اطلاعات بر رضایت شغلی کارمندان دولتی تأثیر می‌گذارد؟	اشتراک دانش و در دسترس بودن اطلاعات تأثیر مثبتی بر رضایت شغلی کارکنان بخش عمومی دارد، که می‌تواند بر فرهنگ اشتراک دانش اخلاقی تأثیرگذار باشد.
۷	Hakimi (2019)	بررسی تأثیر پشتیبانی فناوری اطلاعات از مدیریت دانش بر عملکرد کسب‌وکار، نقش میانجی قابلیت‌های پویا	نقش پشتیبانی فناوری اطلاعات در تسهیل مدیریت دانش و تقویت قابلیت‌های پویای سازمان، به‌عنوان بستری برای بهبود اشتراک و استفاده از دانش سازمانی مورد تأکید قرار گرفت. ارتباط غیرمستقیم با افشای دانش از طریق تسهیل زیرساخت دانش.
۸	Rashid Alipour, Ansari & Seyed Javadin (2019)	بررسی تأثیر اجرای مدیریت دانش بر عملکرد سازمانی (نمونه پژوهش: شرکت شیر پگاه)	پیاده‌سازی مدیریت دانش موجب افزایش اعتماد سازمانی و بهبود عملکرد شد. فراهم‌سازی زمینه برای شکل‌گیری فرهنگ اشتراک دانش در سازمان از نتایج مهم پژوهش بود.
۹	Azimi, Morshedi & Abbasi (2020)	نقش یادگیری سازمانی و فرایند کسب و کار بر عملکرد دانش محور سازمانها (نمونه پژوهش: اداره دارا یی زنجان و سازمان صنعت و معدن بجنورد)	یادگیری سازمانی نقش مهمی در کاهش مقاومت کارکنان در برابر اشتراک دانش و بهبود جریان دانشی داشت؛ ارتباط با مؤلفه‌های فرهنگی و ادراکی مؤثر بر افشای دانش.
۱۰	Naeiji, Ghorchi Beigi, Ghorbani Farab & Seyed Alijaanlou (2022)	بررسی نقش تعدیل‌گری مدیریت دانش استراتژیک در رابطه‌ی ابعاد سرمایه فکری و عملکرد نوآورانه (نمونه پژوهش: شرکت‌های لبنی)	مدیریت راهبردی دانش باعث تقویت اثر سرمایه فکری بر عملکرد نوآورانه شد. ساختارهای حمایتی برای تسهیل بروز دانش پنهان و مشارکت اخلاقی در به‌اشتراک‌گذاری اطلاعات برجسته شد.

جدول ۱. پیشینه پژوهش

ردیف	نویسنده/ نویسندگان (سال پژوهش)	عنوان پژوهش	مهم‌ترین یافته‌ها و نتایج مرتبط با پژوهش
۱۱	Arabshehi, Kabiri & Behboudi (2022)	تاثیر ارزش دانش مدیران عالی بر شیوه‌های اشتراک دانش، نوآوری باز و عملکرد سازمان	ارزش‌گذاری دانش از سوی مدیران ارشد، موجب تشویق کارکنان به اشتراک‌گذاری اطلاعات حساس و تقویت نوآوری باز شد. اهمیت فرهنگ دانشی اخلاقی محور در سازمان تأکید گردید.
۱۲	Hassani Ahangar & Maqdisi (2013)	مطالعه و معرفی الگوریتم‌های هوش مصنوعی جهت حل مسائل بهینه‌سازی؛ کاربرد، مزایا و معایب	الگوریتم‌های هوش مصنوعی به دلیل توانایی در حل مسائل پیچیده بهینه‌سازی، می‌توانند در تحلیل رفتارهای تصمیم‌گیری انسانی، مانند افشای دانش، نقش پشتیبان داشته باشند.
۱۳	Valivand Zamani & Mortezaazadeh (2024)	بررسی تاثیر استفاده از هوش مصنوعی بر فرایند تصمیم‌گیری مدیران در مدیریت بحران‌های سازمانی	استفاده از هوش مصنوعی باعث افزایش دقت، سرعت و انسجام در تصمیم‌گیری مدیران به‌ویژه در شرایط بحرانی می‌شود؛ این یافته به کاربرست مدل‌های زبانی در تحلیل نیت افشاگری باشد.
۱۴	Torabi & Eghbal (2024)	ارائه مدل حاکمیت هوش مصنوعی برای حکمرانی دولت در جمهوری اسلامی ایران	ارائه الگویی از حکمرانی هوش مصنوعی در دولت ایران، که شامل اصول اخلاقی، شفافیت و پاسخ‌گویی است؛ مرتبط با تبیین نقش سامانه‌های هوش مصنوعی در تصمیم‌گیری دانشی و اخلاقی در سازمان‌های دولتی.
۱۵	Moulton (2020)	افشای اطلاعات حساس (رساله دکتری)	افشای اطلاعات حساس تحت تأثیر شدت پیامدها و نوع رابطه فرد با سازمان، عاملی کلیدی در نیت افشای دانش حساس است.
۱۶	Lu & Sinnott (2020)	راهکارهای امنیتی و حریم خصوصی برای سامانه‌های سلامت هوشمند	ملاحظات امنیتی و حریم خصوصی نقش بازدارنده‌ای در اشتراک اطلاعات حساس دارند که می‌تواند نیت افشاگری را محدود کند.
۱۷	Valentine & Godkin (2019)	شدت اخلاقی، تصمیم‌گیری اخلاقی و نیت افشاگرانه	شدت اخلاقی و تصمیم‌گیری اخلاقی به‌طور مستقیم بر نیت افشاگری تأثیر دارند، که نقش تمایل اخلاقی فردی را تأیید می‌کند.
۱۸	Abdullah Sani, Sallehuddin, Salim, Devi, Mohanadas, Yaakup & Azman (2020)	تأثیر عوامل فردی بر نیت افشاگری: از منظر حسابرسان داخلی آینده	متغیرهای فردی مانند ارزش‌های اخلاقی و شناخت پیامدها تأثیر معناداری بر نیت افشاگری دانش حساس دارند.
۱۹	Chua, Thinakaran & Vasudevan (2023)	موانع اشتراک دانش در سازمان‌ها – یک مرور نظری	موانع اشتراک دانش شامل ترس از پیامدها، عدم اعتماد و ساختارهای سخت‌گیرانه، از عوامل کاهش‌دهنده افشای اطلاعات حساس هستند.
۲۰	Saeed, Khan, Zada, Zada, Vega-Muñoz & Contreras-Barraza (2022)	پیوند رهبری اخلاقی با اشتراک دانش پیروان	رهبری اخلاق‌مدار با ایجاد حس مالکیت روان‌شناختی، موجب افزایش تمایل به اشتراک دانش در کارکنان می‌شود که عاملی مؤثر در نیت افشاگری است.
۲۱	Van Portfliet, Irfan & Kenny (2022)	وقتی کارکنان حرف می‌زنند: جنبه‌های مدیریت منابع انسانی در افشاگری	سیاست‌ها و ساختارهای منابع انسانی حمایتی (مثل کانال‌های امن گزارش تخلف) نقش کلیدی در افزایش تمایل کارکنان به افشای تخلفات و دانش حساس دارد.
۲۲	Shen, Lythreatis, Singh & Cooke (2024)	فرا تحلیل رفتار پنهان‌سازی دانش در سازمان‌ها: پیش‌زمینه‌ها، پیامدها و شرایط محدودکننده	مداخلات منابع انسانی مانند آموزش اخلاقی و نظام‌های پاداش‌دهی می‌توانند تا ۳۰ درصد از رفتار پنهان‌سازی دانش را کاهش دهند و به بهبود شفافیت و تسهیم دانش حساس منجر شوند.

جدول ۱. پیشینه پژوهش

ردیف	نویسنده/ نویسندگان (سال پژوهش)	عنوان پژوهش	مهم‌ترین یافته‌ها و نتایج مرتبط با پژوهش
۲۳	Xu, Huang & Huang (2023)	عدالت اطلاعاتی و رفتارهای پنهان‌سازی دانش کارکنان: نقش میانجی شناسایی سازمانی و نقش تعدیل‌گر حساسیت به عدالت	عوامل روان‌شناختی مانند عدم اعتماد و ترس از تلافی تا ۲۵ درصد باعث افت عملکرد تیمی از طریق افزایش پنهان‌سازی دانش می‌شوند؛ پیشنهادهایی برای مداخلات رفتاری ارائه شده است.
۲۴	Wen & Ma (2021)	پیش‌زمینه‌های پنهان‌سازی دانش و تأثیر آن‌ها بر عملکرد سازمانی	ادراک عدالت اطلاعاتی از سوی کارکنان باعث کاهش میل به پنهان‌سازی دانش می‌شود؛ هویت سازمانی قوی این رابطه را تقویت و حساسیت عدالت آن را تعدیل می‌کند.

بررسی (۲۴) مطالعه داخلی و خارجی جدول (۱) نشان می‌دهد که اگرچه مطالعات بسیاری به ابعاد مرتبط پرداخته‌اند، اما تأثیر ابعاد «تمایل اخلاقی فردی»، «فرهنگ اشتراک دانش اخلاقی» و «موانع ادراک‌شده» بر نیت افشاگری در سازمان‌های دولتی هنوز کمتر مورد توجه قرار گرفته است. همچنین پژوهش‌های اخیر به‌طور گسترده‌تری به نقش فناوری‌های نوین مانند کلان‌داده و هوش مصنوعی در تسهیل اشتراک یا محافظت از دانش حساس پرداخته‌اند.

با این حال، شکاف‌ها و خلأهای پژوهشی زیر قابل شناسایی‌اند:

- اغلب پژوهش‌ها به مطالعه ابعاد «دانش آشکار»^۱ یا «اشتراک دانش»^۲ پرداخته‌اند، درحالی‌که موضوع «نیت افشای دانش حساس» در بستر سازمان‌های دولتی (با شرایط خاص فرهنگی، اخلاقی و بوروکراتیک) کمتر به‌صورت هدفمند بررسی شده است.
- ابعاد دانشی مؤثر بر نیت افشاگری (مانند درک پیامدها، اخلاق حرفه‌ای، اعتماد سازمانی، مالکیت روان‌شناختی و قابلیت یادگیری) عمدتاً به صورت پراکنده بررسی شده‌اند و مدل‌های تلفیقی اندکی در این زمینه وجود دارد.
- پژوهش‌ها به‌ندرت از رویکردهای ترکیبی کمی-کیفی، یا روش‌های نوین مانند سطح پاسخ و مدل‌های زبانی هوش مصنوعی در تحلیل نیت افشاگری بهره گرفته‌اند.
- در ادبیات داخلی، علی‌رغم تمرکز بر مدیریت دانش، به ارتباط بین شاخص‌های حکمرانی خوب، شفافیت اداری و نیت افشای دانش کمتر پرداخته شده است.

مرور مطالعات پیشین نشان می‌دهد که هرچند پژوهش‌های متعددی به بررسی عوامل مؤثر بر اشتراک دانش، افشاگری اخلاقی، رهبری اخلاقی‌مدار، ملاحظات امنیتی^۳، و مدیریت دانش در سازمان‌ها پرداخته‌اند، اما خلأهایی جدی در زمینه درک ابعاد ادراکی و اخلاقی افشای دانش حساس در سازمان‌های دولتی وجود دارد. بیشتر پژوهش‌ها به‌صورت پراکنده به نقش متغیرهای فرهنگی، فردی و فناورانه پرداخته‌اند، در حالی که کمتر مطالعه‌ای با رویکردی تلفیقی به تحلیل نیت افشاگری در بستر سازمان‌های خدماتی و دولتی پرداخته است. همچنین، فقدان بهره‌گیری از مدل‌های زبانی نوین در تحلیل تصمیم‌گیری اخلاقی کارکنان، به‌ویژه در زمینه افشای دانش، محسوس است. در همین راستا، پژوهش حاضر با ترکیب روش سطح پاسخ و بهره‌گیری از مدل زبانی گروک-۳^۴، به بررسی نیت افشاگرانه کارکنان ادارات دولتی پرداخته و تلاش کرده است با مدل‌سازی رفتار انسانی و مقایسه آن با تفسیر هوش مصنوعی از شرایط، گامی نو در فهم لایه‌های پنهان تصمیم‌گیری دانشی بردارد.

همچنین، با توجه به پیشرفت روزافزون هوش مصنوعی، فهم سازوکار درک و تفسیر این فناوری از شرایط واقعی محیطی امری ضروری است. این مطالعه با مقایسه نحوه واکنش و تصمیم‌گیری انسان و مدل‌های زبانی هوش مصنوعی در زمینه افشای دانش حساس، گامی نوین در فهم قابلیت‌های شناختی و ادراکی^۴ هوش مصنوعی و کاربرد آن در تحلیل تصمیم‌های اخلاقی کارکنان سازمان‌های دولتی برداشته است.

وجه نوآوری پژوهش

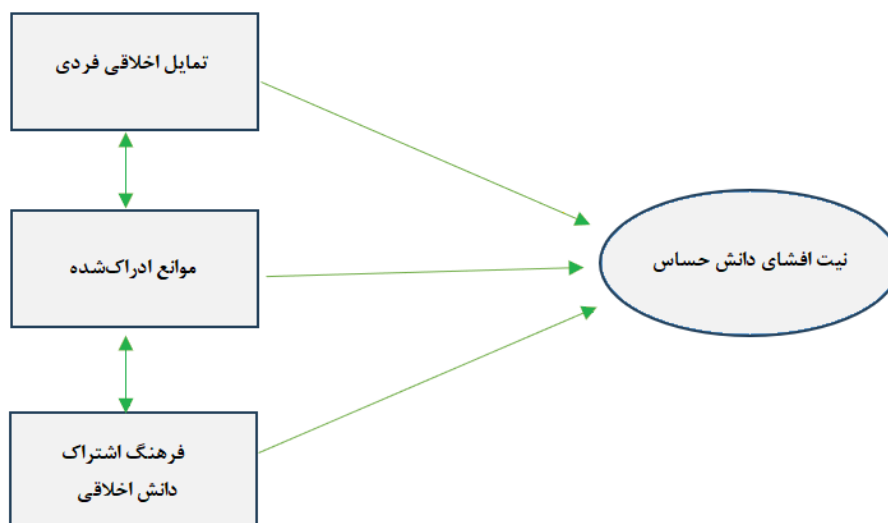
1. Revealed Knowledge
2. Knowledge Sharing
3. Security Considerations
4. Cognitive & Perceptual Abilities

پژوهش حاضر با برخورداری از چندین جنبه نوآورانه، تلاش کرده است تحلیلی متفاوت و کاربردی از نیت افشای دانش حساس در بستر سازمان‌های دولتی ارائه دهد. نوآوری‌های پژوهش به شرح زیر می‌باشند.

- ترکیب هم‌زمان متغیرهای فرهنگی، فردی و ساختاری با رویکردی تلفیقی در تحلیل نیت افشای دانش حساس: پژوهش حاضر برای نخستین بار به صورت هم‌زمان سه مؤلفه کلیدی شامل عوامل فرهنگی (فرهنگ اشتراک دانش اخلاقی)، فردی (تمایل درونی^۱) و ساختاری (موانع اشتراک دانش) را به عنوان متغیرهای مؤثر بر نیت افشای دانش حساس در سازمان‌های دولتی بررسی می‌کند، در حالی که اغلب مطالعات پیشین صرفاً به یکی از این ابعاد پرداخته و تحلیل تلفیقی آن‌ها مغفول مانده است.
- به کارگیری روش سطح پاسخ و طراحی مرکب مرکزی در کنار تحلیل زبانی هوش مصنوعی: در این پژوهش برخلاف رویه غالب پژوهش‌های منابع انسانی که به روش‌های آماری کلاسیک یا تحلیل‌های کیفی محدود می‌شوند، از تلفیق روش سطح پاسخ و طرح مرکب مرکزی^۲ در کنار مقایسه تحلیلی نتایج با خروجی مدل زبانی گروک-۳ استفاده شده است. این ترکیب نوآورانه، امکان شناسایی نقطه بهینه ترکیب عوامل انسانی - دانشی و نیز درک تطبیقی از نحوه تحلیل شرایط افشاگری توسط انسان و هوش مصنوعی را فراهم کرده است.
- تمرکز بر بستر واقعی ادارات خدماتی دولتی در استان لرستان با هدف بومی‌سازی تحلیل رفتار افشاگرانه: مطالعه حاضر با تمرکز بر داده‌های واقعی جمع‌آوری شده از کارکنان ادارات خدماتی دولتی در استان لرستان، ضمن توجه به ویژگی‌های فرهنگی و سازمانی محلی، زمینه ارائه یافته‌هایی بومی‌سازی شده در خصوص رفتار افشاگرانه را فراهم می‌سازد. این تمرکز جغرافیایی و سازمانی، پژوهش را از حیث مدل عمومی و توسعه راهکارهای مدیریتی در سطح منطقه‌ای و ملی متمایز می‌سازد.

مدل مفهومی

بر اساس مرور نظام‌مند ادبیات و پیشینه پژوهش، مدل مفهومی پژوهش شکل ۱ طراحی شده است. در این مدل، سه متغیر مستقل شامل «تمایل اخلاقی فردی»، «فرهنگ اشتراک دانش اخلاقی» و «موانع ادراک شده» در نظر گرفته شده‌اند که تأثیر آن‌ها بر متغیر وابسته یعنی «نیت افشای دانش حساس» مورد بررسی قرار می‌گیرد.



شکل ۱. مدل مفهومی ارتباط متغیرهای پژوهش

با توجه به مدل مفهومی و مرور ادبیات، سه فرضیه زیر برای آزمون در این مطالعه تدوین شدند:

- فرضیه اول: تمایل اخلاقی فردی کارکنان رابطه خطی و مثبت با میزان نیت افشای دانش حساس دارد و قوی‌ترین متغیر است.
- فرضیه دوم: فرهنگ اشتراک دانش اخلاقی در سازمان با نیت افشای دانش حساس کارکنان رابطه خطی و قوی دارد.

- فرضیه سوم: موانع ساختاری اشتراک دانش رابطه خطی منفی با نیت افشای دانش حساس کارکنان دارد.

روش‌شناسی پژوهش

مطابق با لایه نخست «پیاژ پژوهش»^۱ ساندرز و همکاران (Saunders, Lewis, & Thornhill, 2019)، این پژوهش بر مبنای فلسفه اثبات‌گرایی انجام شده است که بر وجود واقعیتی عینی و قابل سنجش درباره تأثیر متغیرهای فردی، فرهنگی و ساختاری بر نیت افشای دانش حساس تأکید دارد و با بهره‌گیری از روش‌های کمی و استدلال آماری، در پی کشف این واقعیت است. رویکرد تحقیق، ترکیبی از قیاس نظری^۲ (بر پایه مبانی علمی حوزه‌های مرتبط) و استقرا تجربی^۳ (با گردآوری و تحلیل داده‌های میدانی) است که به‌صورت مکمل به آزمون نظریه‌ها و شناسایی الگوهای واقعی می‌پردازد. راهبرد پژوهش نیز از نوع شبه‌آزمایشی و پیمایشی است که در آن از سناریوهای فرضی (وینیت‌ها) و طراحی مرکب مرکزی استفاده شده و داده‌ها با روش مدل‌سازی سطح پاسخ تحلیل شده‌اند تا اثرات خطی، غیرخطی و تعاملی بین متغیرها مشخص شود. همچنین افق زمانی پژوهش مقطعی است و داده‌ها در یک بازه زمانی مشخص از جامعه آماری جمع‌آوری شده‌اند تا روابط علی بین متغیرها بدون در نظر گرفتن تغییرات زمانی بررسی شوند.

در این پژوهش، داده‌ها به‌صورت میدانی و از طریق پرسش‌نامه سناریومحور^۴ (وینیت) گردآوری شد. ابزار شامل ۱۵ سناریوی ساختگی اما واقع‌گرایانه بود که با ترکیب سطوح پایین، متوسط و بالا از سه متغیر اصلی (تمایل فردی به افشاگری، فرهنگ اشتراک دانش اخلاقی و موانع ساختاری اشتراک دانش) طراحی شد و از پاسخ‌دهندگان خواسته شد میزان نیت خود برای افشای دانش حساس در هر وضعیت را بر اساس طیف لیکرت ۷ درجه‌ای مشخص کنند. در این پژوهش دانش حساس، اطلاعات مربوط به تخلفات مالی سازمانی مانند پنهان‌کاری در مناقصات^۵، دریافت رشوه^۶، جعل اسناد مالی^۷، یا تخصیص منابع برخلاف آیین‌نامه‌ها بوده و در پرسش‌نامه این موضوع ذکر شده است. پرسش‌نامه بر اساس ادبیات پژوهش و مشورت با متخصصان طراحی و پس از اجرای آزمایشی، اصلاح شد. نسخه نهایی به‌صورت حضوری و آنلاین، میان کارکنان ادارات خدماتی استان لرستان توزیع و پس از پالایش داده‌ها تحلیل شد.

جامعه آماری این پژوهش شامل ۷۰۰ نفر از کارکنان ادارات خدماتی دولتی استان لرستان (شهرداری، تأمین اجتماعی، آموزش و پرورش، مراکز بهداشتی، ادارات برق و آبفا) بود. با استفاده از فرمول کوکران (Cochran, 1977) و سطح اطمینان ۹۵٪، با توجه به محدود بودن جامعه آماری، حداقل ۲۴۸ نفر برآورد شد. با این حال، برای افزایش دقت مدل غیرخطی سطح پاسخ و با در نظر گرفتن ریزش، ۴۰۰ پرسش‌نامه توزیع و ۳۸۵ مورد معتبر جمع‌آوری و تحلیل شد. نمونه‌گیری به روش تصادفی طبقه‌ای انجام شد تا نمایندگی متوازن ادارات حفظ شود. در اجرای سناریوها، از طراحی آماری «بلوک ناقص متوازن» (Hinkelmann & Kempthorne, 1994)، استفاده شد؛ به هر پاسخ‌دهنده تنها ۵ سناریوی تصادفی از مجموع ۱۵ سناریو اختصاص یافت تا از خستگی ذهنی جلوگیری شود. هر سناریو به‌طور متوسط توسط ۱۲۸ نفر ارزیابی شد که برای مدل‌سازی پاسخ، تحلیل اثرات و تعامل‌ها کافی بود (Atzmüller & Steiner, 2010). توزیع پرسش‌نامه‌ها با تضمین محرمانگی انجام شد تا از سوگیری پاسخ‌ها جلوگیری شود (Dillman, Smyth, & Christian, 2014). این ساختار نمونه‌گیری، شرایط لازم برای تحلیل دقیق در چارچوب طرح مرکب مرکزی را فراهم کرد.

در این پژوهش، یک متغیر وابسته و سه متغیر مستقل تعریف شده‌اند و سناریوها با ترکیب سطوح متغیرهای مستقل طراحی شدند. پاسخ‌دهندگان میزان نیت افشای دانش حساس (متغیر وابسته) را برای هر سناریو ارزیابی کردند. این ساختار امکان تحلیل اثرات اصلی و تعاملی متغیرها را در قالب طرح مرکب مرکزی و روش سطح پاسخ فراهم کرد. متغیر اول، تمایل اخلاقی فردی، نشان‌دهنده انگیزه درونی کارکنان برای اقدام اخلاقی، حتی در شرایط پرریسک است و شامل ابعادی مانند مسئولیت‌پذیری و قضاوت اخلاقی است (Potipiroon & Wongpreedee, 2020). متغیر دوم، فرهنگ اشتراک دانش اخلاقی، به هنجارهای سازمانی حامی افشای مسئولانه اطلاعات حساس اشاره دارد و شامل سیاست‌های حمایتی، شفافیت و حمایت از افشاگران است (Fischer & Döring, 2022). متغیر سوم، موانع ساختاری اشتراک دانش، به محدودیت‌هایی چون سلسله‌مراتب بسته، نبود چارچوب‌های حمایتی و ترس از تلافی می‌پردازد (Wen et al., 2021). متغیر وابسته یعنی نیت افشای دانش حساس، آمادگی ذهنی کارکنان برای گزارش اطلاعات مهم یا تخلفات را بیان می‌کند و به عواملی مانند ارزیابی خطرات، پیامدهای اجتماعی و

1. Research Onion
2. Theoretical Analogy
3. Empirical Induction
4. Scenario-Based
5. Secrecy In Tenders
6. Receiving A Bribe
7. Forgery Of Financial Documents

اعتبار فرد گزارشگر وابسته است (Cheliatsidou et al., 2023). استفاده از سناریوهای طراحی شده چارچوبی دقیق برای تحلیل روابط میان این متغیرها در سازمان‌های خدماتی دولتی فراهم ساخته است.

برای سنجش روایی محتوایی، نسخه اولیه پرسش‌نامه شامل ۱۵ سناریو مرتبط با متغیرهای پژوهش در اختیار پنج کارشناس حوزه مدیریت دانش و رفتار سازمانی قرار گرفت. آن‌ها هر گویه را از نظر وضوح، انطباق مفهومی و ارتباط با نیت افشای دانش حساس ارزیابی کردند. اصلاحات پیشنهادی آنان در نسخه نهایی اعمال شد. مطابق با توصیه (Rossi & Nock, 1982)، این روش، معیار معتبری برای ارزیابی روایی محتوا در مطالعات سناریومحور محسوب می‌شود. شاخص روایی محتوا برای تمامی گویه‌ها برابر با ۰/۹۹ بود. برای بررسی پایایی، پرسش‌نامه روی نمونه‌ای آزمایشی از (۵۰) نفر اجرا شد. آلفای کرونباخ برای نیت افشای دانش حساس برابر با ۰/۸۸ به دست آمد. همچنین، همبستگی بین آیتم‌های هر متغیر مستقل بالای ۰/۴۰ بود که نشان‌دهنده انسجام درونی مناسب و ثبات پاسخ‌ها در مواجهه با سناریوهای مختلف است.

برای تحلیل داده‌های حاصل از (۱۵) سناریو، از روش سطح پاسخ با طرح مرکب مرکزی استفاده شد. (۳۸۵) کارمند ادارات دولتی لرستان پرسش‌نامه‌های وینیتهمحور را تکمیل کرده و نیت افشای دانش حساس را در مقیاس لیکرت ۷ درجه‌ای سنجیدند. داده‌ها در نرم‌افزار دیزاین اکسپرت^۱ نسخه (۱۳) و براساس مدل درجه دوم شامل اثرات خطی، مربعی و تعاملی تحلیل شد. شایستگی مدل با ضریب تبیین، آنالیز واریانس و تحلیل باقی‌مانده‌ها تأیید گردید. سپس برای تفسیر تعامل‌ها، نمودارهای سه‌بعدی ترسیم شد. در گام نهایی، از تابع مطلوبیت برای تعیین ترکیب بهینه متغیرها استفاده شد. نمرات نیت افشا به بازه [۰/۱] نرمال‌سازی و میانگین هندسی مطلوبیت‌ها محاسبه شد. این رویکرد، روابط تعاملی پیچیده را آشکار کرده و راهکارهای سیاستی مؤثر برای ارتقای افشای دانش حساس در سازمان‌های دولتی ارائه می‌دهد.

این پژوهش طی یک فرایند مرحله‌ای و منسجم انجام شد. ابتدا با مرور ادبیات، سه متغیر مستقل (تمایل اخلاقی فردی، فرهنگ اشتراک دانش اخلاقی و موانع ساختاری) و متغیر وابسته (نیت افشای دانش حساس) استخراج شدند. سپس با تدوین چارچوب نظری و فرضیه‌ها، (۱۵) سناریوی وینیتهمحور بر اساس طرح مرکب مرکزی طراحی و با نظر خبرگان اصلاح شد. نمونه‌گیری از (۳۸۵) کارمند ادارات دولتی لرستان انجام و هر پاسخ‌دهنده به پنج سناریوی تصادفی پاسخ داد. داده‌ها در نرم‌افزار دیزاین اکسپرت وارد و پس از آماده‌سازی، مدل سطح پاسخ درجه دوم برازش شد. تحلیل آماری شامل آنالیز واریانس، بررسی شاخص‌های برازش و تحلیل باقی‌مانده‌ها بود. تعامل متغیرها از طریق نمودارهای سه‌بعدی بررسی و در نهایت، با استفاده از تابع مطلوبیت^۲، ترکیب بهینه متغیرها برای بیشینه‌سازی نیت افشای دانش حساس شناسایی و مبنای ارائه پیشنهادها سیاستی قرار گرفت.

افزون بر تحلیل داده‌های انسانی، در این پژوهش برای نخستین‌بار از مدل زبانی پیشرفته «گروک-۳» نیز استفاده شد. پاسخ‌های این مدل به سناریوهای طراحی شده، به‌طور مستقل استخراج و سپس با داده‌های انسانی از نظر جهت و شدت تأثیر متغیرها مقایسه گردید. این مقایسه به پژوهشگران امکان داد تا تفاوت‌ها و شباهت‌های درک انسانی و ماشینی از تصمیم‌گیری اخلاقی را بررسی کنند.

یافته‌های پژوهش

برای مدل‌سازی روابط غیرخطی میان سه متغیر مستقل شامل «تمایل اخلاقی فردی (A)، موانع ساختاری اشتراک دانش (B) و فرهنگ اشتراک دانش اخلاقی (C) با متغیر وابسته «نیت افشای دانش حساس»، از روش سطح پاسخ با بهره‌گیری از نرم‌افزار دیزاین اکسپرت نسخه ۱۳ استفاده شد. این روش یکی از تکنیک‌های پیشرفته طراحی آزمایش است که امکان تحلیل هم‌زمان اثرات اصلی و تعاملی متغیرهای مستقل را در قالب مدل‌های درجه دوم فراهم می‌سازد. طراحی آزمایش‌ها بر اساس طرح مرکب مرکزی انجام گرفت و برای هر یک از متغیرهای مستقل، سه سطح «کم»، «متوسط» و «زیاد» تعریف شد. در مجموع، (۱۵) سناریوی منحصر به فرد از ترکیب این سطوح ایجاد و در قالب پرسش‌نامه‌های وینیتی طراحی گردید. برای هر سناریو، از پاسخ‌دهندگان خواسته شد میزان «نیت افشای دانش حساس» را بر اساس محتوای سناریو در مقیاس لیکرت ۷ درجه‌ای (از ۱ = بسیار پایین تا ۷ = بسیار بالا) ارزیابی کنند. میانگین نمرات هر سناریو، به‌عنوان مقدار مشاهده‌شده متغیر وابسته، در مدل وارد شد.

در این پژوهش، علاوه بر جمع‌آوری داده‌های انسانی، از پاسخ‌های مدل هوش مصنوعی گروک-۳ نیز برای همان (۱۵) سناریو بهره گرفته شد. بدین ترتیب، داده‌ها در دو سطح تحلیل شدند: (۱) داده‌های واقعی کارمندان ادارات خدماتی دولتی استان لرستان و (۲) داده‌های تولیدشده توسط مدل هوش مصنوعی. هدف از این رویکرد، بررسی تطابق یا تفاوت در درک و ارزیابی نیت افشاگرانه در مواجهه با شرایط مختلف اخلاقی و ساختاری بیت انسان و هوش مصنوعی است.

تحلیل مدل‌ها برای هر دو منبع داده به صورت جداگانه انجام گرفت و نتایج حاصل از مدل سطح پاسخ در دو حالت مقایسه شد. این مقایسه، علاوه بر ارزیابی اعتبار و انسجام ادراک انسانی در مواجهه با سناریوهای طراحی شده، قابلیت تحلیلی مدل هوش مصنوعی را در درک و بازنمایی ساختارهای رفتاری و اخلاقی مرتبط با افشای دانش حساس را به آزمون می‌گذارد.

تحلیل داده های انسانی

برای انتخاب مناسب‌ترین مدل جهت تحلیل روابط میان متغیرهای مستقل و متغیر وابسته (نیت افشای دانش حساس)، چهار مدل رگرسیونی مختلف شامل مدل خطی ساده، مدل اثرات متقابل، مدل درجه دوم و مدل درجه سوم طراحی و با یکدیگر مقایسه شدند. مقایسه مدل‌ها در جدول (۲) انجام شد. از میان این گزینه‌ها، مدل درجه دوم به عنوان مناسب‌ترین مدل انتخاب گردید؛ زیرا این مدل بالاترین میزان ضریب تبیین و ضریب تبیین تعدیل شده را در بین سایر مدل‌ها نشان داد، در حالی که آزمون عدم برازش مدل در سطح خطای ۱ درصد معنادار نبود، و آزمون معناداری کلی مدل در همان سطح معنادار تشخیص داده شد. لازم به ذکر است فاصله کم ضریب تاثیر تبیین تعدیل شده (۰/۹۰۶۴) و ضریب تبیین (۰/۹۵۰۷) نشان دهنده اعتبار بالای مدل هست اما به دلیل احتمال اندک وجود بیش برازش بهتر است یافته‌ها با احتیاط تفسیر شوند.

جدول ۲. آنالیز تعیین بهترین مدل آماری

منبع	سطح معنی داری	انحراف استاندارد	ضریب تبیین	ضریب تبیین تعدیل یافته	عدم برازش ^۱	مجموع مربعات خطاهای پیش‌بینی شده
مدل خطی	۰/۰۰۰۱	۰/۵۴۰۳	۰/۸۲۳۱	۰/۷۸۹۹	۰/۰۲۶۲	۸/۵۷
مدل اثر متقابل	۰/۸۹۰۹	۰/۵۸۵۶	۰/۸۳۱۱	۰/۷۵۳۱	۰/۰۱۵۶	۳۲/۹۳
مدل درجه ۲	۰/۰۰۵	۰/۳۶۰۶	۰/۹۵۰۷	۰/۹۰۶۴	۰/۱۲۷	۱۸/۳۶
مدل درجه ۳	۰/۰۵۰۷	۰/۲۳۲۸	۰/۹۸۷۷	۰/۹۶۱۰	۰/۶۲۴۷	۲۰/۹۶

برای بررسی معنی داری مدل رگرسیون انتخاب شده (مدل درجه دوم)، از تحلیل واریانس استفاده شد. جدول (۳) نتایج این تحلیل را ارائه می‌دهد.

جدول ۳. تحلیل واریانس مدل درجه دوم

منبع تغییرات	درجه آزادی	مجموع مربعات	میانگین مربعات	آماره F ^۲	سطح معنی داری
مدل	۴	۲۴/۶۱	۶/۱۵	۵۱/۷۹	< ۰/۰۰۰۱
تمایل، A	۱	۹/۲۲	۹/۲۲	۷۷/۵۸	< ۰/۰۰۰۱
موانع، B	۱	۳/۴۸	۳/۴۸	۲۹/۳۰	< ۰/۰۰۰۱
فرهنگ، C	۱	۹/۰۲	۹/۰۲	۷۵/۹۷	< ۰/۰۰۰۱
A ^۲	۱	۲/۸۹	۲/۸۹	۲۴/۳۱	۰/۰۰۰۲
باقی مانده	۱۵	۱/۷۸	۰/۱۱۸۸	-	-
عدم برازش	۱۰	۱/۴۷	۰/۱۴۷۴	۰/۶۳۰۲	۰/۱۷۴۲
خطای خالص	۵	۰/۳۰۸۳	۰/۰۶۱۷	-	-

براساس جدول (۳) که نتایج تحلیل واریانس مدل درجه دوم را نشان می‌دهد، اثرهای خطی هر سه متغیر مستقل شامل تمایل اخلاقی فردی (A)، موانع ساختاری اشتراک دانش (B) و فرهنگ اشتراک دانش اخلاقی (C) بر متغیر وابسته «نیت افشای دانش حساس» در سطح ۱ درصد معنی دار بوده‌اند. در مقابل، تعامل‌های دوجانبه بین این متغیرها معنادار نبوده‌اند، که نشان می‌دهد اثرات ترکیبی آن‌ها بر متغیر وابسته تأثیر افزایشی یا تقویتی معناداری نداشته است. از میان مؤلفه‌های درجه دوم، تنها اثر درجه دوم متغیر تمایل اخلاقی فردی (A^۲) معنادار بود، در حالی که اثرات درجه دوم موانع ساختاری (B^۲) و فرهنگ اشتراک دانش اخلاقی (C^۲) معنی دار تشخیص داده نشدند. همچنین، آزمون عدم برازش مدل در سطح ۱

1. Lack Of Fit
2. F-value

درصد بی معنی بود که نشان می دهد مدل رگرسیون درجه دوم دارای برازش مناسب با داده های گردآوری شده از کارکنان ادارات خدماتی استان لرستان بوده و می توان به دقت ضرایب آن اطمینان داشت.

مدل رگرسیون درجه دوم برای متغیر وابسته «نیت افشای دانش حساس» به صورت کدگذاری شده، به شکل رابطه (۱) نمایش داده می شود.

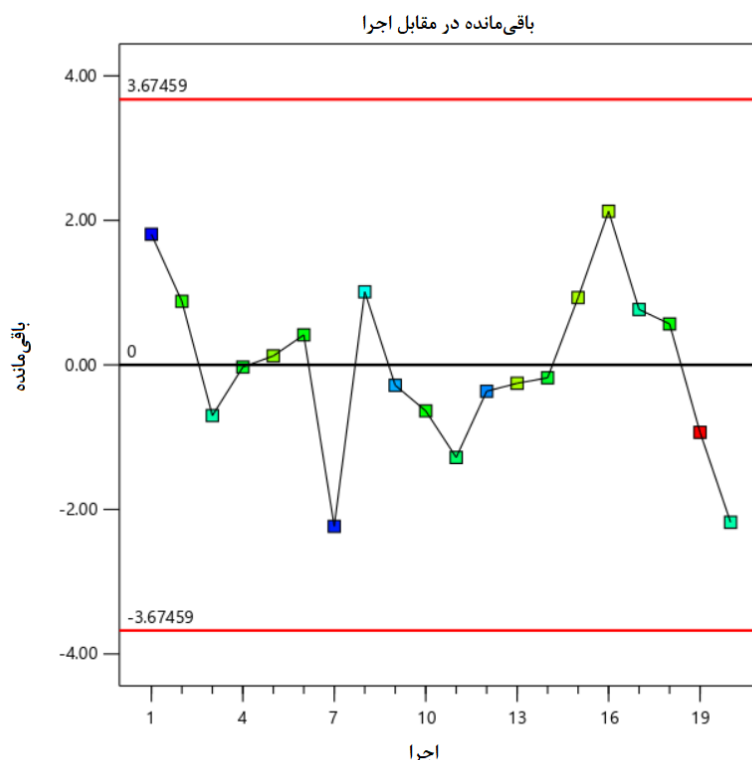
$$Y = 3.81 + 0.96 \times A - 0.59 \times B + 0.95 \times C - 0.76 \times A^2 \quad \text{رابطه ۱.}$$

در این رابطه، Y نیت افشای دانش حساس، A تمایل اخلاقی فردی، B موانع ساختاری اشتراک دانش و C فرهنگ اشتراک دانش اخلاقی است.

تحلیل رگرسیون درجه دوم حاکی از آن است که ضرایب خطی A و C مثبت هستند، و تأثیر A اندکی برتر از C برآورد شده است. در مقابل، ضریب خطی B منفی است که نقش بازدارندگی موانع ساختاری را در کاهش نیت افشاگری برجسته می کند. در میان مؤلفه های درجه دوم، تنها A^2 معنی دار و جهت منفی است؛ این موضوع نشان می دهد که پس از عبور «تمایل اخلاقی فردی» از نقطه ای بهینه، هرگونه افزایش بیشتر در A رشد اثربخشی آن کاهش می دهد و نشانه ای از اشباع را به تصویر می کشد. نتیجه آزمون «عدم برازش» مدل تأیید می کند که ساختار رگرسیونی انتخاب شده به خوبی با داده های کارمندان ادارات خدماتی لرستان تطابق دارد و می توان به اعتبار ضرایب برآوردی اعتماد کرد.

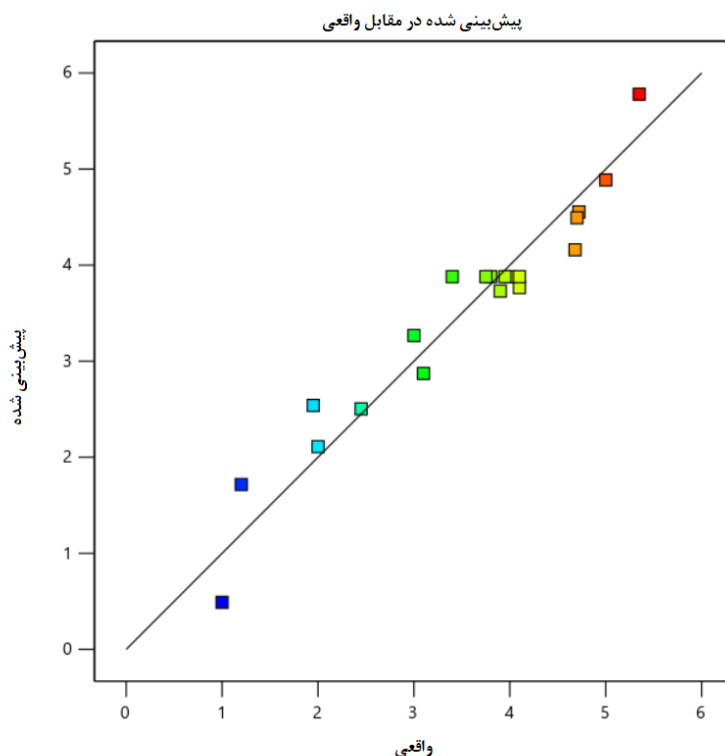
در مجموع، این یافته ها حاکی از آن است که برای تقویت پایدار «نیت افشای دانش حساس»، باید به صورت هم زمان «تمایل اخلاقی فردی» را افزایش، «موانع ساختاری اشتراک دانش» را کاهش و «فرهنگ اشتراک دانش اخلاقی» را تقویت کرد. چنین رویکرد تلفیقی می تواند راهنمای تصمیم سازی مبتنی بر شواهد در مدیریت دانش سازمان های دولتی باشد.

برای بررسی اعتبار آماری ساختار مدل رگرسیونی و تحلیل نحوه پراکندگی خطاها، نمودار باقی مانده ها در برابر مقادیر پیش بینی شده در شکل (۲) ارائه شده است. این نمودار نشان می دهد که باقی مانده ها به طور پراکنده و بدون الگوی منظم در اطراف محور افقی قرار گرفته اند. این پراکندگی تصادفی از عدم وجود روند خاص یا الگوی سیستماتیک در خطاهاست و مؤید آن است که مفروضاتی نظیر نرمال بودن و استقلال خطاها برقرار هستند. علاوه بر این، نبود الگوی مشخص در نمودار بیانگر آن است که مدل در تمامی سطوح متغیرهای مستقل، رفتاری پایدار و یکنواخت دارد و نشانی از وجود واریانس ناهمسان یا نقص ساختاری دیده نمی شود. در نتیجه، این نمودار گواهی بر اعتبار و اتکاپذیری مدل رگرسیونی از منظر مبانی آماری ارائه می دهد.



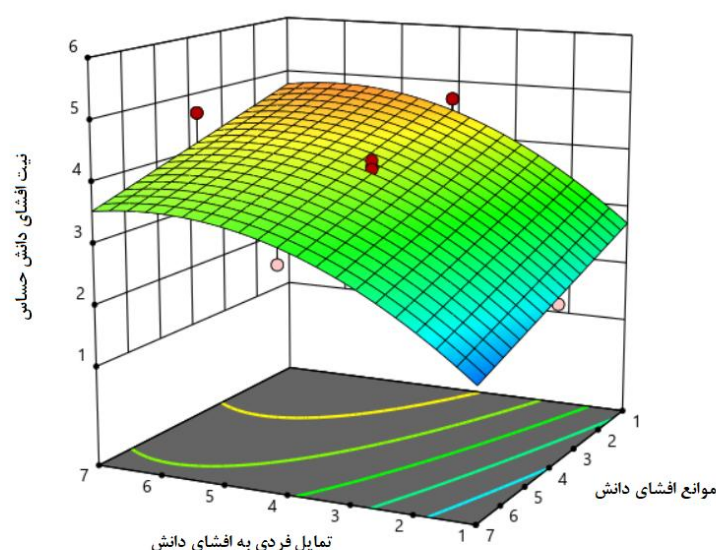
شکل ۲. مقادیر باقی مانده در مقابل اجرا

در شکل (۳) نمودار پراکندگی مقادیر واقعی «نیت افشای دانش حساس» در برابر مقادیر پیش‌بینی شده توسط مدل رگرسیون درجه دوم ترسیم شده است. خط نیم‌ساز، نشانگر ایده‌آل مطابقت کامل پیش‌بینی‌ها با مقادیر مشاهده شده است. هر چه نقاط داده (مقادیر واقعی) به این خط نزدیک‌تر باشند، نشان از دقت بالاتر مدل در تخمین نیت افشاگری دارد. همان‌طور که مشاهده می‌شود، نقاط در حاشیه‌ی اندکی از این خط مستقر هستند، که بیانگر برازش مطلوب مدل، توانایی بالا در پیش‌بینی و حداقل انحراف پیش‌بینی‌ها از مقادیر واقعی می‌باشد. این نتیجه تأییدی بر قابلیت اطمینان مدل رگرسیون درجه دوم در پیش‌بینی نیت افشای دانش حساس بر اساس ترکیب‌های متنوع سطوح «تمایل اخلاقی فردی»، «موانع ساختاری اشتراک دانش» و «فرهنگ اشتراک دانش اخلاقی» اشتراک دانش می‌باشد.



شکل ۳. مقدار متغیر نیت افشای دانش حساس پیش‌بینی شده در مقابل واقعی

در شکل ۴، نمودار سه‌بعدی تأثیر موانع ساختاری اشتراک دانش (B) و تمایل اخلاقی فردی (A) بر نیت افشای دانش حساس (Y) ارائه شده است. از آنجا که فرهنگ اشتراک دانش اخلاقی (C) در مدل به صورت خطی عمل می‌کند و تعامل غیرخطی با دو متغیر دیگر ندارد، در این شکل مقدار آن در سطح میانی مقیاس لیکرت (عدد ۴) ثابت در نظر گرفته شده است. بر مبنای این نمودار، با افزایش مداوم موانع ساختاری، میزان نیت افشاگری به صورت خطی کاهش می‌یابد و انحنای نمودار در جهت محور تمایل فردی نمایانگر حضور اثر درجه دوم تمایل اخلاقی فردی است؛ به گونه‌ای که رشد تأثیر تمایل اخلاقی فردی تا یک نقطه بهینه ادامه یافته و پس از آن در سطوح بالای تمایل سرعت افزایش اثر آن کند شده و وارد ناحیه اشباع می‌شود. در شرایطی که تمایل اخلاقی فردی در سطوح پایین و موانع ساختاری در سطوح بالا قرار داشته باشد، کمترین مقدار نیت افشای دانش حساس مشاهده می‌شود. برعکس، هنگامی که موانع ساختاری حداقلی و تمایل اخلاقی فردی حداکثری باشد، بیشینه نیت افشای دانش حساس به ثبت می‌رسد. این الگو بر اهمیت هم‌افزایی کاهش موانع و تقویت انگیزه اخلاقی برای ارتقای مؤثر نیت افشای دانش حساس تأکید می‌کند.



شکل ۴. اثر تغییرات تمایل اخلاقی فردی و موانع ساختاری اشتراک دانش بر سطح نیت افشای دانش حساس سازمانی

تحلیل داده های هوش مصنوعی

برای انتخاب بهترین مدل در تحلیل داده های پیش بینی شده توسط هوش مصنوعی، چهار ساختار رگرسیونی شامل مدل خطی، مدل اثرات متقابل، مدل درجه دوم و مدل درجه سوم در جدول (۴)، با یکدیگر مقایسه شدند. مدل کامل درجه دوم، به دلیل کسب بیشترین ضریب تبیین و ضریب تبیین تعدیل شده، عدم برازش غیر معنی دار در سطح ۵ درصد، پایین ترین مجموع مربعات خطاهای پیش بینی و معناداری کلی مدل در سطح ۱ درصد، به عنوان بهترین گزینه انتخاب گردید. ضریب تبیین تعدیل یافته (۰/۹۷۵۴) فاصله کمی با ضریب تبیین (۰/۹۸۷۰) دارد که نشان دهنده اعتبار بالای مدل است اما به دلیل احتمال وجود بیش برازش بهتر است نتایج با احتیاط تفسیر شوند.

جدول ۴. آنالیز تعیین بهترین مدل آماری

منبع	سطح معنی داری	انحراف استاندارد	ضریب تبیین	ضریب تبیین تعدیل یافته	عدم برازش	مجموع مربعات خطاهای پیش بینی شده
مدل خطی	۰/۰۰۰۱	۰/۴۵۲۵	۰/۹۱۸۴	۰/۹۰۳۱	۰/۰۸۵۵	۶/۴۱
مدل اثر متقابل	۰/۰۳۱۸	۰/۳۶۱۸	۰/۹۳۸۱	۰/۹۳۸۱	۰/۱۹۳۶	۸/۸۱
مدل درجه ۲	۰/۰۰۶۳	۰/۲۲۸۲	۰/۹۸۷۰	۰/۹۷۵۴	۰/۸۲۶۴	۲/۳۲
مدل درجه ۳	۰/۶۹۲۹	۰/۲۵۰۶	۰/۹۷۰۳	۰/۹۷۰۳	۰/۷۷۳۷	۸/۸۵

برای بررسی معنی داری مدل رگرسیون درجه دوم مدل، از تحلیل واریانس استفاده شد. جدول ۵ نتایج این تحلیل را ارائه می دهد.

جدول ۵. تحلیل واریانس مدل درجه دوم

منبع تغییرات	درجه آزادی	مجموع مربعات	میانگین مربعات	آماره F	سطح معنی داری
مدل	۷۶	۳۹/۳۵	۶/۵۶	۱۰۳/۷۲	< ۰/۰۰۰۱

منبع تغییرات	درجه آزادی	مجموع مربعات	میانگین مربعات	آماره F	سطح معنی داری
تمایل، A	۱	۲۶/۹	۲۶/۹	۴۲۵/۳۶	< ۰/۰۰۰۱
موانع، B	۱	۳/۶	۳/۶	۵۶/۹۳	< ۰/۰۰۰۱
فرهنگ، C	۱	۶/۴	۶/۴	۱۰۱/۲۲	< ۰/۰۰۰۱
AB	۱	۱/۱۲	۱/۱۲	۱۷/۷۹	۰/۰۰۱۰
AC	۱	۰/۴۰۵	۰/۴۰۵	۶/۴۱	۰/۰۲۵۱
A ²	۱	۰/۹۲۴۵	۰/۹۲۴۵	۱۴/۶۲	۰/۰۰۲۱
باقی مانده	۱۳	۰/۸۲۲	۰/۰۶۳۲	-	-
عدم برازش	۸	۰/۴۵۲	۰/۰۵۶۵	۰/۷۶۳۵	۰/۶۵۰۹
خطای خالص	۵	۰/۳۷	۰/۰۷۴	-	-

بر اساس جدول شماره (۵)، می‌توان دریافت که ضریب‌های خطی هر سه متغیر مستقل شامل تمایل اخلاقی فردی (A)، موانع ساختاری اشتراک دانش (B) و «فرهنگ اشتراک دانش اخلاقی» (C) همگی تأثیر معناداری بر نیت افشای دانش حساس داشته‌اند. همچنین، دو اثر متقابل $A \times B$ (تعامل تمایل و موانع) و $A \times C$ (تعامل تمایل و فرهنگ) و نیز اثر درجه دوم A^2 (تمایل اخلاقی فردی) نیز در سطح (۱) درصد معنادار ارزیابی شده‌اند، که نشان‌دهنده وجود روابط پیچیده‌تر و غیرخطی میان متغیرهاست. در مقابل، آزمون عدم برازش مدل در سطح (۵) درصد معنادار نبوده است، که این موضوع تأییدکننده برازش مطلوب مدل درجه دوم، دقت ضرایب و اعتبار آماری ساختار مدل پیشنهادی در داده‌های حاصل از هوش مصنوعی است. بر این اساس، رابطه رگرسیونی درجه دوم مدل به‌صورت کدگذاری‌شده در قالب رابطه (۲) ارائه شده است.

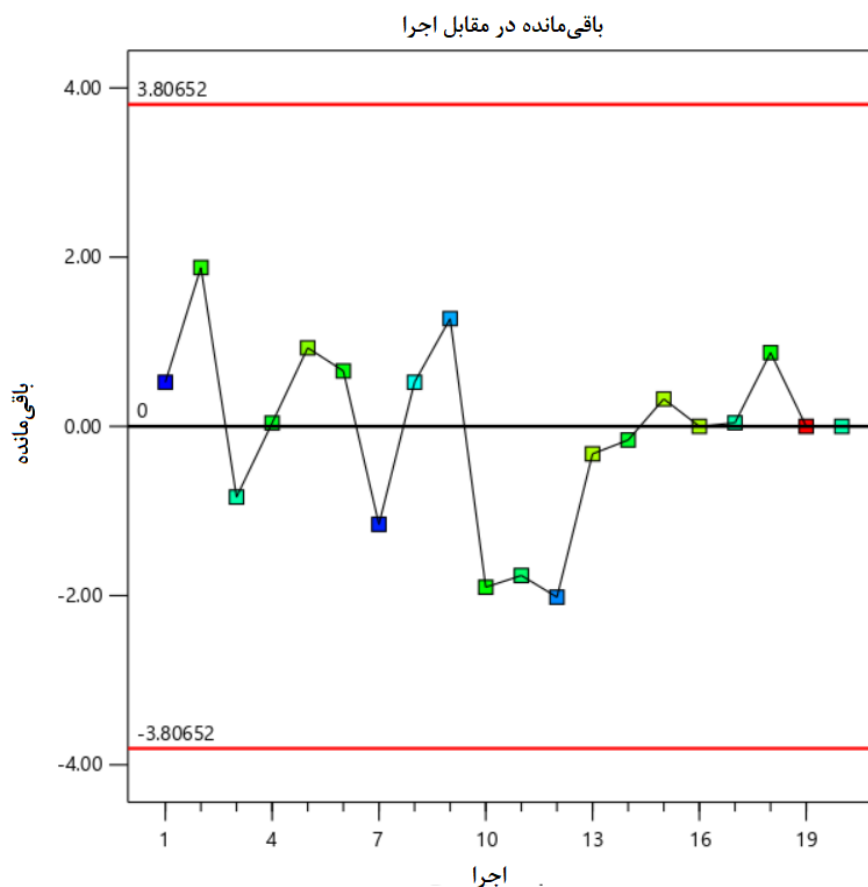
$$Y = 3.79 + 1.64 \times A - 0.6 \times B + 0.8 \times C - 0.375 \times AB + 0.225 \times AC - 0.43 \times A^2 \quad \text{رابطه ۲.}$$

در این رابطه، Y نیت افشای دانش حساس، A تمایل اخلاقی فردی، B موانع ساختاری اشتراک دانش و C فرهنگ اشتراک دانش اخلاقی است.

بر اساس رابطه ۲، متغیر «تمایل اخلاقی فردی» (A) و «فرهنگ اشتراک دانش اخلاقی» (C) هر دو دارای ضریب خطی مثبت هستند و ضریب خطی تمایل اخلاقی فردی بیش از دو برابر فرهنگ اشتراک دانش است، که این موضوع نشان می‌دهد اثرگذاری تمایل فردی بر نیت افشای دانش حساس بسیار پررنگ‌تر از فرهنگ سازمانی اخلاق‌مدار است. در مقابل، «موانع ساختاری اشتراک دانش» (B) دارای ضریب خطی منفی است که بیانگر نقش بازدارنده آن در تقویت نیت افشاگرانه است. در میان اثرات متقابل، تعامل میان تمایل اخلاقی فردی و موانع ساختاری (AB) ضریبی منفی دارد، به این معنا که ترکیب هم‌زمان تمایل بالا و وجود موانع ساختاری، به جای خنثی‌سازی یکدیگر، به‌صورت هم‌افزا شدت بازدارندگی را افزایش می‌دهد. این در حالی است که تعامل بین تمایل اخلاقی فردی و فرهنگ اشتراک دانش اخلاقی (AC) دارای ضریب مثبت است، به‌گونه‌ای که هم‌افزایی بین این دو عامل می‌تواند به تقویت قابل توجهی در نیت افشای دانش حساس منجر شود. تنها مؤلفه غیرخطی حاضر در این مدل، اثر درجه دوم تمایل اخلاقی فردی (A^2) است که ضریب آن منفی است. این موضوع نشان می‌دهد که با وجود اثر مثبت اولیه تمایل اخلاقی فردی، در سطوح بالاتر، شیب رشد اثر آن کاهش یافته و به ناحیه اشباع می‌رسد. به عبارت دیگر، اگرچه افزایش تمایل فردی در ابتدا نیت افشاگرانه را تقویت می‌کند، اما این اثر در سطوح بالا کند شده و به مرور به سقف تأثیر خود نزدیک می‌شود.

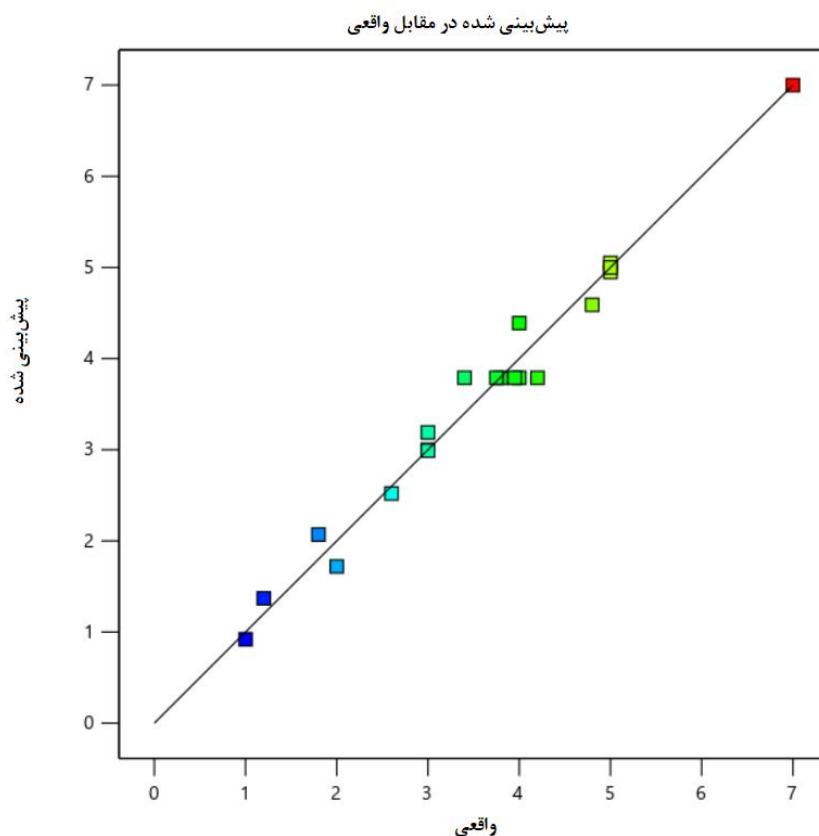
به منظور ارزیابی آماری مدل رگرسیونی و بررسی نحوه توزیع خطاها، نمودار باقی‌مانده‌ها در مقابل مقادیر پیش‌بینی‌شده در شکل ۵ رسم گردید. تحلیل این نمودار نشان می‌دهد که باقی‌مانده‌ها به شکلی پراکنده و بدون جهت‌گیری خاص در اطراف محور افقی قرار گرفته‌اند. این نوع پراکندگی نشان می‌دهد که خطاها ساختار منظم یا الگوی مشخصی ندارند و می‌توان فرض استقلال و نرمال بودن آن‌ها را قابل قبول دانست. همچنین، فقدان

الگو در نمودار به این معناست که مدل در سطوح مختلف متغیرهای مستقل رفتاری متوازن دارد و نشانه‌ای از واریانس ناهمسان یا ضعف ساختاری در مدل دیده نمی‌شود. بنابراین، این نمودار نشان‌دهنده آن است که مدل از نظر مفروضات آماری پایه، از اعتبار مناسبی برخوردار است.



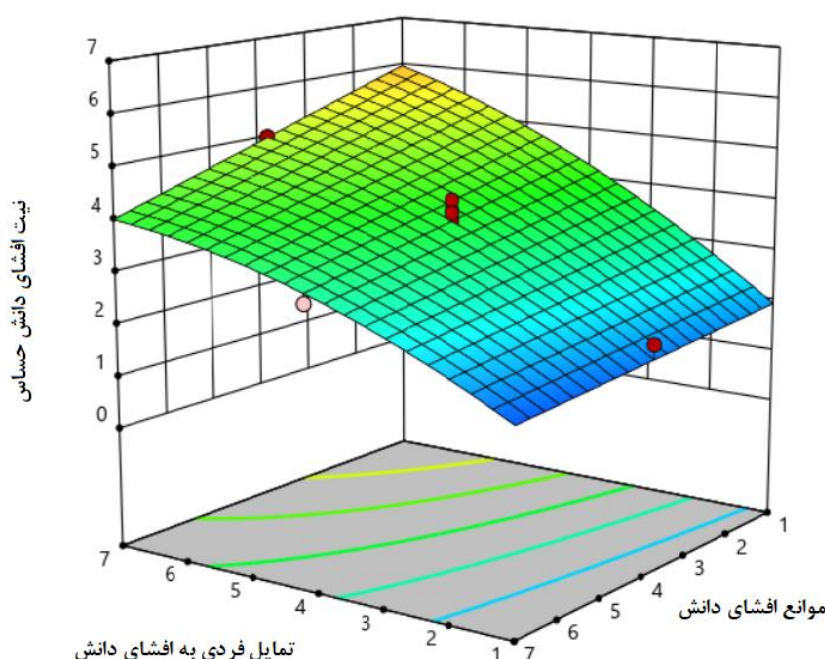
شکل ۵. تغییرات مقدار باقی مانده در مقابل اجرا

در شکل (۶)، نمودار مقایسه‌ای بین مقادیر واقعی نیت افشای دانش حساس و مقادیر پیش‌بینی‌شده توسط مدل رگرسیون درجه دوم نمایش داده شده است. در این نمودار، خط (۴۵) درجه نمایانگر حالتی است که در آن مقدار پیش‌بینی‌شده دقیقاً با مقدار واقعی برابر باشد. نزدیکی نقاط رنگی به این خط، نشان‌دهنده دقت بالای مدل در پیش‌بینی متغیر وابسته است. با توجه به توزیع نقاط در اطراف این خط، می‌توان دریافت که مدل مورد استفاده توانسته است مقادیر نیت افشاگرانه را با خطای اندک و دقت مطلوبی بر اساس سطوح مختلف تمایل اخلاقی فردی، موانع ساختاری اشتراک دانش و فرهنگ اشتراک دانش اخلاقی پیش‌بینی کند. این مسئله گواهی بر برازش مناسب مدل و اعتبار آن در تحلیل روابط بین متغیرهای مورد مطالعه است.



شکل ۶. مقدار متغیر نیت افشای دانش حساس پیش‌بینی شده در مقابل واقعی

شکل (۷) نمودار سه‌بعدی تأثیر تمایل اخلاقی فردی و موانع ساختاری اشتراک دانش بر نیت افشای دانش حساس را نمایش می‌دهد، می‌توان روند تغییرات این متغیرها را به‌وضوح مشاهده کرد. در این نمودار، متغیر فرهنگ اشتراک دانش اخلاقی به دلیل اثر خطی، در مقدار میانه (عدد ۴ در مقیاس لیکرت) ثابت در نظر گرفته شده است و نقش آن در تعیین سطح نمودار منعکس می‌شود؛ به‌گونه‌ای که افزایش یا کاهش این متغیر صرفاً باعث جابه‌جایی عمودی سطح نیت افشاگری در نمودار می‌شود. بر اساس این شکل مشاهده می‌شود که کاهش موانع ساختاری به‌طور پیوسته با افزایش نیت افشای دانش حساس همراه است. از سوی دیگر، انحنای سطح در امتداد محور تمایل اخلاقی فردی نشان‌دهنده وجود اثر درجه دوم این متغیر است؛ بدین معنا که در سطوح بالاتر تمایل، شدت اثرگذاری آن به‌تدریج کاهش یافته و وارد ناحیه اشباع می‌شود. در مجموع، پایین‌ترین مقدار نیت افشاگری در شرایطی رخ می‌دهد که تمایل فردی اندک و موانع ساختاری در بالاترین سطح خود قرار دارند؛ در مقابل، بیشترین سطح نیت افشای دانش حساس در شرایط تمایل حداکثری و موانع حداقلی مشاهده می‌شود.



شکل ۷. اثر تغییرات تمایل اخلاقی فردی و موانع ساختاری اشتراک دانش بر سطح نیت افشای دانش حساس سازمانی

در جدول (۶)، مقادیر تابع مطلوبیت برای تعیین سطوح بهینه متغیرها با هدف بیشینه‌سازی نیت افشای دانش حساس، بر اساس داده‌های انسانی و هوش مصنوعی ارائه شده است.

جدول ۶. بهینه‌ترین شرایط مقدار نیت افشای دانش حساس

مطلوبیت	فرهنگ اشتراک دانش اخلاقی	موانع ساختاری اشتراک دانش	تمایل اخلاقی فردی	نیت افشای دانش حساس (هوش مصنوعی)	نیت افشای دانش حساس (انسانی)
۰/۸۷۱	۷	۱	۷	۷	۵/۵۵

در این پژوهش با بهره‌گیری از روش سطح پاسخ و طراحی مرکب مرکزی در دو بستر داده‌های انسانی و مدل هوش مصنوعی گروک-۳، نیت افشای دانش حساس کارکنان ادارات خدماتی استان لرستان مورد تحلیل قرار گرفت. تابع مطلوبیت برای ۱۵ سناریوی ترکیبی سطوح «تمایل اخلاقی فردی»، «فرهنگ اشتراک دانش اخلاقی» و «موانع ساختاری اشتراک دانش» محاسبه گردید.

بر اساس جدول ۶، در شرایط بهینه، یعنی «فرهنگ اشتراک دانش اخلاقی» در بالاترین سطح (۷)، «تمایل اخلاقی فردی» در بالاترین سطح (۷) و «موانع ساختاری اشتراک دانش» در کمترین سطح (۱)، مقدار تابع مطلوبیت پیش‌بینی شده برای نیت افشای دانش حساس توسط هوش مصنوعی برابر ۷ و مقدار متناظر انسانی برابر ۵.۵۵ ثبت شد. این میزان تفاوت تقریباً معادل ۲۶ درصد است و نشان می‌دهد که در شرایط ایده‌آل، هوش مصنوعی جسارت بیشتری به افشای دانش حساس از خود بروز می‌دهد، در حالی که انسان‌ها با احتیاط و ملایمت بیشتری اقدام به افشاگری می‌کنند.

نتیجه‌گیری

این پژوهش با هدف تحلیل و مدل‌سازی نیت افشای دانش حساس کارکنان ادارات خدماتی استان لرستان، با بهره‌گیری از روش سطح پاسخ و طرح مرکب مرکزی انجام شد. در این مسیر، از دو منبع تحلیلی - داده‌های انسانی مبتنی بر سناریوهای واقعی و خروجی‌های مدل هوش مصنوعی گروک-

(۳) استفاده گردید تا تصویری دقیق، ترکیبی و مقایسه‌پذیر از رفتار افشاگرانه در مواجهه با سه متغیر «تمایل اخلاقی فردی»، «فرهنگ اشتراک دانش اخلاقی» و «موانع ساختاری اشتراک دانش» ترسیم شود.

نتایج نشان داد که هر سه متغیر دارای اثر معنادار بر نیت افشاگری هستند، اما شدت و شکل تأثیرگذاری آن‌ها متفاوت است. به‌طور مشخص، تمایل اخلاقی فردی قوی‌ترین عامل پیش‌بینی‌کننده نیت افشاگری بود، اما رابطه آن با نیت افشای دانش حساس غیرخطی و اشباع‌شونده است. این بدان معناست که پس از رسیدن به سطحی معین از انگیزش اخلاقی، احتمالاً افزایش بیشتر آن دیگر به رشد معنادار در نیت افشاگرانه منجر نمی‌شود. چنین الگویی، که برخلاف رویکردهای خطی کلاسیک است، چشم‌اندازی تازه برای تحلیل پیچیدگی‌های اخلاقی در رفتارهای سازمانی فراهم می‌سازد. یافته ما با نتایج پژوهش‌هایی چون (Potipiroon و Wongpreedee 2020) در تأیید نقش انگیزش اخلاقی همسو است، اما در عین حال با افزودن بُعد غیرخطی، دروازه‌ای نو به سوی نظریه‌پردازی در افشای دانش حساس می‌گشاید.

از سوی دیگر، فرهنگ اشتراک دانش اخلاقی رابطه‌ای خطی، مثبت و پایدار با نیت افشاگری دارد. این متغیر در سطوح بالای افشاگری، احتمالاً نقشی کلیدی ایفا می‌کند و می‌تواند حتی در نبود تمایل کامل فردی، محرک عمل افشاگرانه باشد. چنین نقشی تأکیدی است بر اهمیت ساختارهای فرهنگی سازمان در تقویت رفتارهای مسئولانه و افشاگرانه. یافته‌هایی که با دیدگاه‌های (Cegarra-Navarro et al 2021) در اهمیت بستر فرهنگی هم‌راستا هستند.

سومین متغیر، یعنی موانع ساختاری اشتراک دانش، تأثیری منفی و پایدار بر نیت افشاگری داشت. در هر دو منبع داده انسانی و هوش مصنوعی ضریب این متغیر مشابه برآورد شد، که نشان می‌دهد احتمالاً ساختارهای سازمانی ناکارآمد یکی از بازدارنده‌های جدی در مسیر افشاگری محسوب می‌شوند که با مطالعات (Chua et al 2021) همسو است. این نتایج بر ضرورت بازنگری در فرآیندها، قوانین و خطوط گزارش‌دهی سازمانی برای رفع موانع ادراک‌شده تأکید دارند.

در شرایط بهینه، که در آن «فرهنگ اشتراک دانش اخلاقی» و «تمایل اخلاقی فردی» در بالاترین سطح (۷) و «موانع ساختاری» در پایین‌ترین سطح (۱) قرار دارند، تابع مطلوبیت پیش‌بینی‌شده نیت افشای دانش حساس در مدل هوش مصنوعی برابر (۷) و در مدل انسانی برابر ۵/۵۵ برآورد شد. این اختلاف (۲۶) درصدی میان دو مدل، بیانگر نوعی واگرایی در نحوه تحلیل و پیش‌بینی رفتار انسانی توسط انسان و ماشین است؛ واگرایی‌ای که از تفاوت بنیادین در منطق تصمیم‌گیری آن‌ها نشأت می‌گیرد. در حالی که مدل هوش مصنوعی عمدتاً بر روابط آماری و الگوهای داده‌ای تکیه دارد و جسورانه‌ترین پیش‌بینی ممکن را ارائه می‌دهد، داده‌های انسانی تحت تأثیر ملاحظات فرهنگی، روان‌شناختی و اجتماعی قرار دارند؛ عواملی نظیر ترس از پیامدهای اجتماعی، فشارهای جمعی، هنجارهای سازمانی و ویژگی‌های شخصیتی، که به‌طور مستقیم در تصمیم‌گیری انسان نقش دارند ولی در منطق فعلی مدل‌های زبانی بازنمایی کامل نیافته‌اند. از سوی دیگر، نمی‌توان نقش سوگیری‌های پاسخ‌دهی کارکنان، تمایل به ارائه پاسخ‌های مقبول اجتماعی یا ترس از افشای نیت واقعی را در تضعیف دقت داده‌های انسانی نادیده گرفت. تحلیل‌های سطح پاسخ نشان داد که هر دو منبع داده (انسان و هوش مصنوعی) اثر درجه دوم «تمایل اخلاقی فردی» را به‌درستی شناسایی کرده‌اند. ضریب مربعی این متغیر (A^2) منفی برآورد شد که نشان‌دهنده وجود یک الگوی اشباع است؛ به‌گونه‌ای که با افزایش تمایل اخلاقی تا سطوح بالا، شدت تأثیر آن کاهش می‌یابد. با این حال، مدل هوش مصنوعی ضریب خطی «تمایل اخلاقی فردی» را تقریباً دو برابر «فرهنگ اشتراک دانش اخلاقی» برآورد کرده، در حالی که در مدل انسانی، این دو ضریب تقریباً برابر گزارش شده‌اند. این تفاوت، بخشی دیگر از همان واگرایی را توضیح می‌دهد: ماشین، بیش از انسان به اخلاق فردی وزن می‌دهد، در حالی که انسان‌ها فرهنگ سازمانی را هم‌ارز آن تلقی می‌کنند.

علاوه‌براین، ضریب منفی «موانع ساختاری اشتراک دانش» در هر دو مدل تقریباً یکسان بود؛ موضوعی که نشان می‌دهد مدل هوش مصنوعی در شناسایی نقش بازدارنده ساختارهای رسمی و غیررسمی عملکردی نسبتاً دقیق دارد. از منظر انسانی، الگوی تأثیرگذاری متغیرها به‌گونه‌ای است که در سطوح پایین نیت افشای دانش حساس، «تمایل اخلاقی فردی» عامل اصلی محرک رفتار است، اما با افزایش سطح تمایل، به دلیل اثر اشباع، نقش آن کاهش می‌یابد و در سطوح بالا، این «فرهنگ اشتراک دانش اخلاقی» است که به‌عنوان عامل مسلط وارد عمل می‌شود. به‌بیان دیگر، حتی اگر تمایل فردی کاهش یابد، در صورت وجود حمایت فرهنگی کافی در سازمان، افراد همچنان به افشای دانش حساس تمایل پیدا می‌کنند. در مجموع، این نتایج تأکیدی است بر این که مدل‌های هوش مصنوعی علی‌رغم توانایی بالا در تحلیل روندهای کلی، در مواجهه با پیچیدگی‌های رفتاری و بین‌فردی انسانی نیازمند تکمیل با تحلیل انسانی‌اند. به‌کارگیری ترکیبی منابع انسانی و ماشینی، می‌تواند تصویر دقیق‌تری از رفتار افشاگرانه در محیط‌های سازمانی ارائه دهد.

از منظر روش پژوهش، بهره‌گیری از رویکرد سطح پاسخ و طرح مرکب مرکزی، امکان تحلیل غیرخطی و تعاملی میان متغیرها را فراهم ساخت، امکانی که در اغلب پژوهش‌های پیشین با اتکای صرف بر روش‌های توصیفی یا رگرسیونی، مورد غفلت قرار گرفته بود. این مطالعه یکی از نخستین

نمونه‌های به‌کارگیری روش ترکیبی داده انسانی و مدل زبانی در حوزه تحلیل رفتار افشاگرانه در ایران است و می‌تواند الگویی قابل استفاده برای خط‌مشی‌گذاران حوزه مدیریت دانش، منابع سازمانی و رفتار سازمانی فراهم سازد.

یافته‌های این پژوهش بیانگر آن است که تصمیم به افشای دانش حساس در سازمان‌های عمومی، نتیجه برهم‌کنش چندلایه میان انگیزش‌های اخلاقی فردی، موانع ساختاری و سازوکارهای فرهنگی حاکم بر اکوسیستم دانشی سازمان است. از منظر خط‌مشی‌گذاری دانشی، این یافته‌ها نشان می‌دهد که خط‌مشی‌گذاران نمی‌توانند صرفاً با اتکا به تقویت هنجارهای اخلاقی فردی، رفتارهای دانشی مسئولانه را نهادینه کنند، بلکه لازم است به‌صورت هم‌زمان، زیرساخت‌های تسهیل‌گر افشای ایمن دانش نظیر سامانه‌های گزارش‌دهی محرمانه، چارچوب‌های حمایتی قانونی و کاهش پیچیدگی‌های بوروکراتیک را توسعه دهند. همچنین، ترویج فرهنگ اشتراک دانش اخلاقی دانشی از طریق شفاف‌سازی مأموریت‌ها، ارتقای سرمایه اجتماعی درون‌سازمانی، و تقویت نظام‌های پاداش‌دهی به کنش‌های دانشی مسئولانه، نقشی بنیادین در شکل‌گیری اعتماد و امنیت روانی در میان کارکنان ایفا می‌کند. اگرچه ابزارهای هوش مصنوعی مانند گروک-۳ می‌توانند در شبیه‌سازی الگوهای تصمیم‌گیری و تحلیل رفتاری کارکنان مفید واقع شوند، اما یافته‌ها نشان می‌دهد این مدل‌ها فاقد ظرفیت درک بافت‌های ظریف روان‌شناختی و اجتماعی هستند که تصمیم‌گیری‌های دانشی را در محیط واقعی هدایت می‌کنند. بنابراین، هوش مصنوعی باید نه به‌عنوان جایگزین، بلکه به‌مثابه مکملی برای تحلیل‌های انسانی در فرآیند خط‌مشی‌گذاری دانشی به‌کار گرفته شود. در مجموع، نتایج این مطالعه می‌تواند به طراحی و پیاده‌سازی سیاست‌هایی منجر شود که زیرساخت‌های امن و فرهنگ‌های حامی را در اکوسیستم دانشی سازمان‌های عمومی تقویت کرده و مسیر را برای ارتقای شفافیت، عدالت دانشی، یادگیری مستمر و خلق ارزش‌های پایدار سازمانی هموار سازد.

پیشنهادهای پژوهشی و کاربردی

با توجه به محدودیت‌های پژوهش حاضر و خلأهای نظری شناسایی‌شده در بررسی هم‌زمان عوامل دانشی، اخلاقی و ساختاری، انجام پژوهش‌های آینده در مسیرهای زیر می‌تواند به توسعه دانش در این حوزه کمک کند.

- کاوش کیفی مؤلفه‌های روان‌شناختی و فرهنگی در رفتار افشاگری: یافته‌ها نشان دادند که علیرغم سطح بالای تمایل اخلاقی، ملاحظات روان‌شناختی و هنجارهای گروهی مانع از افشای دانش حساس می‌شوند. پیشنهاد می‌شود پژوهش‌های کیفی، مانند مصاحبه‌های نیمه‌ساختاریافته یا گروه‌های کانونی، با کارکنان ادارات خدماتی انجام گیرد تا ابعاد عمیق‌تر متغیرهایی چون «ترس از تلافی»، «واکنش جامعه حرفه‌ای» و «فشارهای میان‌فردی» روشن شود. چارچوب نظری «عدالت اطلاع‌رسانی» (Xu, Huang, & Huang, 2023) می‌تواند مبنای تحلیل این ادراکات باشد.
- مدل‌سازی نقش تعدیل‌گرهای فردی و شغلی: با توجه به تنوع ویژگی‌های جمعیتی و حرفه‌ای کارکنان، پیشنهاد می‌شود متغیرهایی مانند «سابقه خدمت»، «سطح تحصیلات» و «پست سازمانی» به‌عنوان تعدیل‌گر یا میانجی وارد مدل شوند. به‌ویژه، بررسی این‌که آیا تجربه شغلی بالا موجب کاهش اثر موانع ساختاری یا تشدید اشباع اخلاقی می‌شود، می‌تواند به سیاست‌گذاران در طراحی برنامه‌های هدفمند کمک کند (Potipiroon & Wongpreedee, 2020).
- بررسی نقش «شدت اخلاقی» در افشاگری دانشی: پیشنهاد می‌شود متغیر «شدت اخلاقی» (Valentine & Godkin, 2019)، به‌عنوان عامل میانجی یا کنترل در مدل لحاظ شده و تأثیر آن بر رابطه میان «تمایل اخلاقی فردی» و «نیت افشاگری» تحلیل گردد. استفاده از طرح‌های آزمایشی غیرخطی مانند طراحی مرکب مرکزی می‌تواند به شناسایی الگوی دقیق اثر این متغیر کمک کند. پژوهش آینده می‌تواند این متغیر را به‌عنوان میانجی یا کنترل وارد مدل نموده و اثر تغییرات آن را بر رابطه «تمایل اخلاقی فردی» و «نیت افشاگری» بررسی کند. همچنین کاربست طراحی آزمایش مبتنی بر طرح مرکب مرکزی در این چارچوب می‌تواند الگوی غیرخطی تأثیر شدت اخلاقی را آشکار سازد.
- بومی‌سازی و اعتبارسنجی مدل هوش مصنوعی: با توجه به محدودیت گروک-۳ در بازنمایی دقیق شرایط فرهنگی-روانی کارکنان، پیشنهاد می‌شود مدل هوش مصنوعی با استفاده از داده‌های بومی (شامل گزارش‌ها و روایت‌های واقعی از افشاگران در ادارات ایرانی) بازآموزی شده و در قالب ریزتنظیم بهینه‌سازی گردد. این اقدام می‌تواند دقت پیش‌بینی مدل در محیط‌های بومی را ارتقا داده و زمینه توسعه سامانه‌های پشتیبان تصمیم‌گیری فراهم آورد.
- ارزیابی تجربی اثر مداخلات ساختاری-فرهنگی: یافته‌های پژوهش نشان‌دهنده تأثیر قوی موانع ساختاری و حمایت فرهنگی بر نیت افشاگری هستند. در نتیجه، انجام پژوهش‌های مداخله‌ای، مانند اصلاح رویه‌های گزارش‌دهی یا برگزاری کارگاه‌های ارتقاء شفافیت، و سنجش اثر آن‌ها بر متغیر وابسته در قالب مطالعات شبه‌آزمایشی پیشنهاد می‌شود.

بر اساس یافته‌های پژوهش، به‌ویژه تحلیل ترکیبی متغیرها و تفاوت در رفتار انسانی و هوش مصنوعی، چند راهکار مشخص و عملی برای بهبود سازوکار افشای دانش حساس در سازمان‌های خدماتی پیشنهاد می‌شود.

- کاهش ساختاری موانع اشتراک دانش: با توجه به تأثیر منفی و معنادار موانع ساختاری بر نیت افشاگری دانشی ($B = -0.59$) در رابطه ۱، پیشنهاد می‌شود سازمان‌های دولتی، به‌ویژه ادارات خدماتی، با همکاری واحدهای فناوری اطلاعات و حراست، یک سامانه گزارش‌دهی الکترونیکی یکپارچه و رمزگذاری شده طراحی کنند که امکان ارسال ناشناس، پیگیری وضعیت گزارش، و بارگذاری مستندات را فراهم آورد. این سامانه باید به‌صورت مستقل از ساختار سلسله‌مراتبی ادارات عمل کرده و مستقیماً به یک نهاد نظارتی فرابخشی (مانند دیوان عدالت یا وزارت کشور) متصل باشد. همچنین، ابلاغ رسمی و آموزش کارکنان درباره مصونیت قانونی، پشتیبانی قضایی، و کانال‌های امن افشاگری ضروری است تا هزینه روانی-اداری اقدام افشاگرانه به حداقل برسد.
- تقویت انگیزش اخلاقی فردی: نتایج رابطه ۱ نشان دادند که «تأمیل اخلاقی فردی» قوی‌ترین پیش‌بینی‌کننده نیت افشاگری است ($A^2 = 0.96$; -0.76). با این حال، به دلیل اثر اشباع در سطوح بالا، مداخلات سازمانی باید هدفمند و متعادل باشند. پیشنهاد می‌شود سازمان‌ها با طراحی کارگاه‌های تعاملی اخلاق حرفه‌ای مبتنی بر سناریوهای واقعی، ظرفیت تحلیل اخلاقی کارکنان را تقویت کنند. علاوه بر این، ایجاد یک نظام انگیزشی غیرمالی مانند اعطای «نشان افتخار صداقت سازمانی» به کارکنانی که گزارش‌های مؤثر ثبت کرده‌اند، می‌تواند انگیزه درونی را زنده نگه دارد. استفاده از خدمات مشاوره اخلاقی و روان‌شناختی محرمانه برای کارمندان درگیر تعارضات اخلاقی نیز به‌عنوان مکمل ضروری است.
- گسترش فرهنگ اخلاق‌مدار اشتراک دانش: بر اساس رابطه ۱ ضریب مثبت و معنادار فرهنگ اشتراک دانش اخلاقی ($C = 0.95$) نشان می‌دهد که تقویت این فرهنگ نقش کلیدی در تسهیل رفتار افشاگرانه دارد. پیشنهاد می‌شود ارزش‌هایی نظیر شفافیت، صداقت، و مسئولیت‌پذیری در قالب منشور اخلاق سازمانی تدوین و به‌صورت رسمی ابلاغ شود. در سطح عملیاتی، مدیران میانی باید با رفتار الگویی و حمایت علنی از گزارش‌دهندگان، فضای روانی امنی فراهم کنند. همچنین، برگزاری منظم جلسات بازخورد، انتشار خبرنامه‌های داخلی حاوی مثال‌های موفق افشاگری، و تقدیر عمومی از رفتارهای شجاعانه، ابزارهای مؤثری برای نهادینه‌سازی این فرهنگ به‌شمار می‌آیند.
- به‌کارگیری محتاطانه هوش مصنوعی: با وجود قابلیت مدل زبانی گروک-۳ در تحلیل سناریوهای اخلاقی، اختلاف حدود ۲۶ درصد میان برآوردهای آن و انسان نشان داد که هوش مصنوعی هنوز در درک مؤلفه‌های روان‌شناختی و فرهنگی انسانی محدودیت دارد. بنابراین، توصیه می‌شود نتایج تحلیل‌های هوش مصنوعی صرفاً به‌عنوان ابزار مکمل برای سناریونویسی و تحلیل سیاست استفاده شوند، نه جایگزین تحلیل انسانی. برای ارتقاء دقت، لازم است مدل‌های زبانی با استفاده از متون بومی، گزارش‌های واقعی افشاگرانه، و داده‌های رفتاری داخلی ریزتنظیم شوند. ایجاد یک نهاد پژوهشی میان‌رشته‌ای برای اعتبارسنجی اخلاقی مدل‌های زبانی در حوزه مدیریت دولتی می‌تواند گام مؤثری در این مسیر باشد.

محدودیت‌ها

- مقطعی بودن طراحی مطالعه: تمامی داده‌ها در یک دوره زمانی مشخص گردآوری شدند و تغییرات یا پویایی‌های درازمدت «نیت افشای دانش حساس» از طریق طراحی پی‌درپی یا طولی مورد بررسی قرار نگرفت. بنابراین امکان ردیابی تحول نگرش‌ها و تأثیرات بلندمدت مداخلات وجود ندارد.
- محدودیت جامعه و قلمرو جغرافیایی: جامعه آماری پژوهش به کارکنان ادارات خدماتی استان لرستان محدود بود. این موضوع می‌تواند عاملیت‌پذیری نتایج را به سایر سازمان‌های دولتی یا بخش‌های غیردولتی در استان‌های دیگر کشور کاهش دهد.
- استفاده از سناریوهای وینیت‌محور: هرچند روش طرح مرکب مرکزی با وینیت‌ها امکان تحلیل دقیق اثرات متغیرها را فراهم ساخت، اما ماهیت فرضی و ساخته‌شده سناریوها ممکن است رفتار واقعی و عینی کارکنان را کاملاً بازتاب ندهد. آزمون میدانی و مشاهده مستقیم تصمیم‌گیری‌های واقعی می‌تواند بینش متفاوتی ارائه کند.
- سوگیری احتمالی پاسخ‌دهی: جمع‌آوری داده‌ها از طریق پرسش‌نامه حتی با تضمین ناشناس ماندن در معرض سوگیری مطلوبیت اجتماعی و خستگی پاسخ‌دهی بود. این امر به ویژه در تخصیص چندمتغیره‌ی سناریوها و پاسخ به ۵ وینیت می‌تواند دقت ادراک واقعی کارکنان را تحت تأثیر قرار دهد.
- عدم شمول برخی متغیرهای میانجی و تعدیل‌گر: این مطالعه تنها سه متغیر اصلی فرهنگی، فردی و ساختاری را بررسی کرد و از متغیرهای روان‌شناختی میانجی‌گر (مانند شدت اخلاقی) یا تعدیل‌گرهای جمعیتی-حرفه‌ای (مانند سابقه خدمت یا موقعیت سازمانی) صرف‌نظر نمود.

- محدودیت‌های فناوری هوش مصنوعی: پیش‌بینی‌های گروک-۳ مبتنی بر معماری بسته است و محققان امکان دسترسی یا اصلاح لایه‌های میانی و وزن‌دهی‌های داخلی آن را نداشتند. علاوه بر این، مدل بر پایه داده‌های عمومی آموزش یافته و سازگاری دقیق آن با فرهنگ و زبان فارسی و بافت سازمانی ایرانی انجام نشده است؛ موضوعی که ممکن است دقت درک ادراکات پیچیده انسانی را کاهش دهد.

مشارکت‌های نویسندگان

علی اصغر عبدشاهی: مفهوم‌سازی، مدیریت اجرایی پروژه، روش‌شناسی، تحلیل رسمی، مدیریت داده‌ها، محاسبات آماری، تهیه پیش‌نویس اولیه مقاله.

حجت رضایی‌ارجمند: مرور پیشینه نظری، طراحی ابزار گردآوری داده، گردآوری داده‌ها، بازبینی و ویرایش مقاله.

شروین وحیدیان قزوینی: انتخاب، انجام آزمایش‌ها و تفسیر نتایج حاصل از هوش مصنوعی.

تضاد منافع

نویسندگان تضاد منافع ندارند.

References

- Abdullah Sani, N., Sallehuddin, A. S., Salim, A. S. A., Devi, N., Mohanadas, & Yaakup, K., & Azman, H. (2020). The influence of individual factors on whistleblowing intention: The perspective of future internal auditors. *International Journal of Innovation, Creativity and Change*, 11(12), 198–217.
- Arabshehi, M., Kabiri, A., & Behboudi, O. (2022). The effect of top managers' knowledge value on knowledge sharing practices, open innovation, and organizational performance. *Strategic Management of Organizational Knowledge*, 5(1), 165–191. <https://doi.org/10.47176/smok.2022.1413> (In Persian)
- Atzmüller, C., & Steiner, P. M. (2010). Experimental vignette studies in survey research. *Methodology: European Journal of Research Methods for the Behavioral and Social Sciences*, 6(3), 128–138.
- Azimi, H., Morshedi, A., & Abbasi, D. (2020). The role of organizational learning and business processes on knowledge-based performance of organizations (Case study: Zanzan Tax Office and Bojnourd Industry and Mine Organization). *Strategic Management of Organizational Knowledge*, 3(2), 167–209. <https://doi.org/10.47176/smok.2020.1117> (In Persian)
- Cegarra-Navarro, J. G., Jiménez-Jiménez, D., & García-Pérez, A. (2021). An integrative view of knowledge processes and a learning culture for ambidexterity: Toward improved organizational performance in the banking sector. *IEEE Transactions on Engineering Management*, 68(2), 408–417. <https://doi.org/10.1109/TEM.2019.2917430>
- Cheliatsidou, A., Sariannidis, N., Garefalakis, A., Passas, I., & Spinthiropoulos, K. (2023). Exploring attitudes towards whistleblowing in relation to sustainable municipalities. *Administrative Sciences*, 13(9), 199. <https://doi.org/10.3390/admsci13090199>
- Chua, K., Thinakaran, R., & Vasudevan, A. (2023). Knowledge sharing barriers in organizations – A review. *TEM Journal*, 12(1), 184–191. <https://doi.org/10.18421/TEM121-24>
- Cochran, W. G. (1977). *Sampling Techniques* (3rd ed.). John Wiley & Sons.
- Dalkir, K. (2011). *Knowledge management in theory and practice* (2nd ed.). MIT Press.
- Dillman, D. A., Smyth, J. D., & Christian, L. M. (2014). *Internet, phone, mail, and mixed-mode surveys: The tailored design method* (4th ed.). John Wiley & Sons.
- Feng, N., Huang, H., Tian, X., & Wang, Y. (2025). Configurational effects of organizational context factors on knowledge sharing in CoPS development. *Kybernetes*. <https://doi.org/10.1108/K-09-2024-2412>
- Financial Times. (2025, February 18). Elon Musk's risqué Grok AI chatbot offers challenge to risk-averse rivals. *Financial Times*.

- Fischer, C., & Döring, M. (2022). Thank you for sharing! How knowledge sharing and information availability affect public employees' job satisfaction. *International Journal of Public Sector Management*, 35(1), 76–93. <https://doi.org/10.1108/IJPSM-10-2020-029>
- Hakimi, I. (2019). Investigating the effect of IT support for knowledge management on business performance: The mediating role of dynamic capabilities. *Strategic Management of Organizational Knowledge*, 2(2), 147–172. <https://doi.org/10.47176/smok.2019.1015> (In Persian)
- Hassani Ahangar, M. R., & Maqdisi, A. (2013). A study and introduction of artificial intelligence algorithms for solving optimization problems: Applications, advantages, and disadvantages. *Passive Defence*, 4(3), 13–24. <https://doi.org/10.47176/pnd.2013.03> (In Persian)
- Hassani Ahangar, M. R., Tavallaei, R., & Shadmanfar, M. H. (2023). Presenting the model of a wise core - a capable network for the role-playing of universities in the knowledge management of modern Islamic civilization issues based on imam khamenei's thought. *Strategic Management of Organizational Knowledge*, 6(4), 21–48. <https://doi.org/10.47176/smok.2023.1659>
- He, H., & Wei, L. (2022). Preventing knowledge hiding behaviors through workplace friendship and altruistic leadership: The mediating role of positive emotions. *Frontiers in Psychology*, 13, Article 905890. <https://doi.org/10.3389/fpsyg.2022.905890>
- Hinkelmann, K., & Kempthorne, O. (1994). *Design and Analysis of Experiments, Volume I: Introduction to Experimental Design*. John Wiley & Sons.
- Kang, M. M. (2023). Whistleblowing in the public sector: A systematic literature review. *Review of Public Personnel Administration*, 43(2), 381–406. <https://doi.org/10.1177/0734371X221078784>
- Karimi, A., Ghaderi, S., & Moradinejad, E. (2020). A sociological study of the perception of administrative corruption and its social factors in particularistic cultures (Case study: Khorramabad city). *Social Problems of Iran*, 10(1), 241–262. <https://doi.org/10.22059/ijsp.2020.76175> (In Persian)
- Laihonen, H., Kork, A. A., & Sinervo, L. M. (2023). Advancing public sector knowledge management: Towards an understanding of knowledge formation in public administration. *Knowledge Management Research & Practice*, 22(3), 223–233. <https://doi.org/10.1080/14778238.2023.2187719>
- Lu, Y., & Sinnott, R. O. (2020). Security and privacy solutions for smart healthcare systems. In *Innovation in Health Informatics: A smart healthcare primer* (pp. 189–216). Elsevier. <https://doi.org/10.1016/B978-0-12-819043-2.00008-3>
- Moulton, E. E. (2020). *The disclosure of sensitive information* (Doctoral dissertation). Columbia University.
- Naeiji, M. J., Ghorchi Beigi, E., Ghorbani Farab, M. T., & Seyed Alijaanlou, M. H. (2022). Examining the moderating role of Strategic Management of Organizational Knowledge in the relationship between intellectual capital dimensions and innovative performance (Case study: Dairy companies). *Strategic Management of Organizational Knowledge*, 5(4), 145–171. <https://doi.org/10.47176/smok.2022.1460> (In Persian)
- Nazari Zadeh, F., Farsijani, E., & Khazaiel, H. (2023). The future of artificial intelligence in organizational strategic management. *Vital Sciences and Emerging Fields* 1(1), 119–156. (In Persian)
- Nicholls, A. R., et al. (2021). Snitches get stitches and end up in ditches: A systematic review of the factors associated with whistleblowing intentions. *Frontiers in Psychology*. <https://doi.org/10.3389/fpsyg.2021.631538>
- Potipiroon, W. (2024). Reward expectancy and external whistleblowing: Testing the moderating roles of public service motivation, seriousness of wrongdoing, and whistleblower protection. *Public Personnel Management*, 53(2), 309–345. <https://doi.org/10.1177/00910260231222814>
- Potipiroon, W., & Wongpreedee, A. (2020). Ethical climate and whistleblowing intentions: Testing the mediating roles of public service motivation and psychological safety among local government employees. *Public Personnel Management*, 50(3), 327–355. <https://doi.org/10.1177/0091026020944547>
- Rashid Alipour, Z., Ansari, M., & Seyed Javadin, S. R. (2019). Investigating the impact of knowledge management implementation on organizational performance (Case study: Pegah Dairy Company). *Strategic Management of Organizational Knowledge*, 2(4), 113–151. <https://doi.org/10.47176/smok.2019.1105> (In Persian)

- Razmerita, L., Kirchner, K., & Nielsen, P. (2016). What factors influence knowledge sharing in organizations? A social dilemma perspective of social media communication. *Journal of Knowledge Management*, 20(6), 1225–1246. <https://doi.org/10.1108/JKM-03-2016-0112>
- Reuters. (2025, February 18). Musk's xAI unveils Grok-3 AI chatbot to rival ChatGPT, China's DeepSeek. Reuters.
- Saeed, I., Khan, J., Zada, M., Zada, S., Vega-Muñoz, A., & Contreras-Barraza, N. (2022). Linking ethical leadership to followers' knowledge sharing: Mediating role of psychological ownership and moderating role of professional commitment. *Frontiers in Psychology*, 13, Article 841590. <https://doi.org/10.3389/fpsyg.2022.841590>
- Saunders, M., Lewis, P., & Thornhill, A. (2019). *Research Methods for Business Students* (8th ed.).
- Torabi, M. A., & Eghbal, H. (2024). Proposing an AI governance model for government governance in the Islamic Republic of Iran. *State Studies of Contemporary Iran*, 10(4), 11–42. <https://doi.org/10.47176/cigs.2024.10411> (In Persian)
- Valentine, S., & Godkin, L. (2019). Moral intensity, ethical decision making, and whistleblowing intention. *Journal of Business Research*, 98, 277–288. <https://doi.org/10.1016/j.jbusres.2019.01.009>
- Valivand Zamani, H., & Mortezaazadeh, A. (2024). Investigating the impact of artificial intelligence usage on managers' decision making processes in organizational crisis management. *Social Crisis Management*, 16(2), 11–26. <https://doi.org/10.47176/cms.2024.16211> (In Persian)
- Wang, N., Chen, B., Wang, L., Ma, Z., & Pan, S. (2024). Big data analytics capability and social innovation: The mediating role of knowledge exploration and exploitation. *Humanities and Social Sciences Communications*, 11, Article 288. <https://doi.org/10.1057/s41599-024-03288-8>
- Wen, J., Zheng, J., & Ma, R. (2021). Antecedents of knowledge hiding and their impact on organizational performance. *Frontiers in Psychology*, 12, Article 796976. <https://doi.org/10.3389/fpsyg.2021.796976>
- Zia, B., Rezvani, M., & Ahmadian, Y. (2021). Investigating the factors affecting knowledge transfer and its impact on the innovation performance of international strategic alliances in small and medium-sized pharmaceutical businesses. *Entrepreneurship Development*, 13(4), 601–619. <https://doi.org/10.22059/jed.2021.307469.653430> (In Persian)

