

ارائه مدل یکپارچه و چندلایه کاربست هوش مصنوعی در مدیریت فرآیندهای کسب و کار سازمان بر پایه رویکرد علم طراحی

حسن غیائی فرد^۱، محمدرضا کریمی قهرودی^۲، آرمان آذرلی^۳

^۱ دانشجوی دکتری تخصصی، رشته فناوری اطلاعات، حکمرانی فناوری اطلاعات، دانشکده برق و کامپیوتر، دانشگاه صنعتی مالک اشتر، تهران، ایران

^۲ دکتری، عضو هیئت علمی دانشکده برق و کامپیوتر، دانشگاه مالک اشتر، دکتری تخصصی آینده پژوهی، دانشگاه بین‌المللی امام خمینی

^۳ دکتری، مدیریت راهبردی، گرایش مدیریت راهبردی، دانشگاه و پژوهشگاه عالی دفاع ملی و تحقیقات، تهران، ایران

تاریخ انتشار: ۱۴۰۵/۰۶/۳۲

تاریخ پذیرش: ۱۴۰۵/۰۶/۲۳

تاریخ دریافت: ۱۴۰۵/۰۶/۲۲

Abstract:

Business Process Management (BPM) plays a key role in operational excellence and organizational agility; however, traditional process diagnosis methods, relying on manual, static, and retrospective analyses, are increasingly inadequate for addressing the complexity of data-driven processes in the digital transformation era. This study aims to propose, explain, and evaluate an integrated, intelligent, and multi-layered model for applying artificial intelligence algorithms to Business Process Management in order to establish a dynamic and prescriptive decision support system. Grounded in the Design Science Research (DSR) approach, the proposed model is architected as a five-phase processing pipeline. First, a heterogeneous data preprocessing and ingestion layer integrates structured event logs and extracts implicit process knowledge from unstructured textual data using Natural Language Processing (NLP). Second, the Apriori algorithm is employed to discover

چکیده:

مدیریت فرآیندهای کسب و کار (BPM) نقشی کلیدی در تعالی عملیاتی و چابکی سازمان‌ها دارد؛ اما روش‌های سنتی عارضه‌یابی، به دلیل اتکا به تحلیل‌های دستی، ایستا و گذشته‌نگر، پاسخگوی پیچیدگی‌های داده‌محور عصر تحول دیجیتال نیستند. هدف این پژوهش، ارائه و ارزیابی مدلی یکپارچه، هوشمند و چندلایه برای به‌کارگیری الگوریتم‌های هوش مصنوعی در مدیریت فرآیندهای کسب و کار و استقرار سیستم تصمیم‌یار پویا و تجویزی است. مدل پیشنهادی با رویکرد علم طراحی (DSR) و در قالب خط لوله‌ای پنج‌فازی توسعه یافت: پیش‌پردازش و بلع داده‌های ناهمگون، شامل تجمیع لاگ‌های رویداد و استخراج دانش فرآیندی از متون غیرساختاریافته با پردازش زبان طبیعی (NLP)؛ کشف قوانین وابستگی و الگوهای پرتکرار با Apriori؛ خوشه‌بندی رفتاری با DBSCAN؛ مدل‌سازی گرافی و تحلیل شبکه فرآیندها (SNA)؛ و در نهایت، موتور استدلال تجویزی مبتنی بر

dependency rules and frequent patterns. Third, DBSCAN is applied for density-based behavioral clustering. Fourth, graph-based process modeling and Social Network Analysis (SNA) are conducted to identify structural patterns and process bottlenecks. Fifth, a prescriptive reasoning engine and dynamic decision support system are developed based on an innovative Process-Role Relationship Matrix aligned with the APQC reference framework. The proposed model was developed in Python and comprehensively evaluated on the large-scale BPIC ۲۰۱۷ loan application process dataset, comprising ۱,۲۰۲,۲۶۷ events and ۳۱,۵۰۹ unique cases. Experimental results demonstrated that DBSCAN identified ۱۸ standard behavioral clusters and intelligently isolated ۱۶۵ highly critical outlier cases, characterized by an average of ۴,۸۹ rework cycles and ۱,۳۰۹ hours of processing time. Graph-based analysis accurately identified the A_Validating activity as a central structural bottleneck and revealed a detrimental ping-pong cycle between A_Incomplete and A_Validating, occurring ۲۴,۹۸۶ times across ۱۱,۶۶۸ cases (۳۷,۰۳% of all cases). Finally, the prescriptive engine generated targeted operational recommendations for process redesign and intelligent redistribution of organizational workload based on risk scoring. The findings demonstrate a successful transition from retrospective descriptive process mining toward

«ماتریس ارتباطی فرآیند-نقش» همسو با استاندارد APQC. مدل در محیط پایتون پیاده‌سازی و به‌طور تمام‌شمار بر کلان‌داده فرآیند درخواست وام BPIC ۲۰۱۷ شامل ۱,۲۰۲,۲۶۷ رخداد و ۳۱,۵۰۹ پرونده یکتا ارزیابی شد. نتایج نشان داد DBSCAN، ۱۸ خوشه رفتاری و ۱۶۵ پرونده پرت فوق‌بحرانی را شناسایی کرد؛ این پرونده‌ها به‌طور میانگین ۴.۸۹ بار بازگشت و ۱۳۰۹ ساعت زمان رسیدگی داشتند. تحلیل گرافی نیز A_Validating را گلوگاه محوری و چرخه پینگ‌پونگی A_Incomplete و A_Validating را با ۲۴,۹۸۶ تکرار در ۱۱,۶۶۸ پرونده (۳۷.۰۳٪) آشکار کرد. در نهایت، موتور تجویزی با امتیازدهی ریسک، توصیه‌های عملیاتی برای بازطراحی فرآیند و بازتوزیع بار کاری ارائه داد. این پژوهش گذار از فرآیندکاوی توصیفی گذشته‌نگر به مدیریت فرآیند تجویزی آینده‌نگر را نشان می‌دهد.

کلیدواژه‌ها: مدیریت فرآیندهای کسب‌وکار (BPM)، فرآیندکاوی تجویزی، هوش مصنوعی، الگوریتم Apriori، الگوریتم DBSCAN، تحلیل شبکه‌های گرافی (SNA)

forward-looking prescriptive process management.

Keywords: Business Process Management (BPM), Prescriptive Process Mining, Artificial Intelligence, Apriori Algorithm, DBSCAN Algorithm, Graph Network Analysis (SNA)

بخش دوم: مقدمه و بیان مسئله

۱. مقدمه در دهه‌های اخیر، زیست‌بوم سازمانی تحت تأثیر تحولات شگرف تحول دیجیتال و انقلاب صنعتی چهارم دگرگون شده است. سازمان‌های امروزی برای حفظ پایداری و چابکی عملیاتی، بیش از هر زمان دیگری نیازمند بهبود مستمر فرآیندهای داخلی خود هستند. در این میان، مدیریت فرآیندهای کسب‌وکار (BPM) به عنوان رویکردی ساختاریافته برای طراحی، پایش و بهینه‌سازی فرآیندها، نقشی کلیدی در تعالی عملیاتی ایفا می‌کند. با این حال، هم‌زمان با استقرار گسترده سیستم‌های اطلاعاتی و تولید حجم عظیمی از داده‌های فرآیندی (لاگ‌های رویداد)، روش‌های سنتی عارضه‌یابی که فرآیندمحور و مبتنی بر تحلیل‌های دستی یا بازبینی‌های دوره‌ای ایستا هستند، کارآمدی خود را در مواجهه با پیچیدگی‌های داده‌محور از دست داده‌اند.

تلفیق قابلیت‌های پیشرفته هوش مصنوعی با متدولوژی‌های مدیریت فرآیند، افق‌های جدیدی را گشوده است. ابزارهای هوشمند توانایی آن را دارند که با پردازش الگوهای پنهان در کلان‌داده‌های سازمانی، رفتارهای تکراری، فرآیندهای موازی و زیرفرآیندهای فاقد

ارزش افزوده را شناسایی کنند. همچنین، پردازش زبان طبیعی (NLP) امکان استخراج دانش ضمنی را از متون غیرساختاریافته سازمانی فراهم می‌آورد. بهره‌گیری از چارچوب‌های استاندارد بین‌المللی مانند چارچوب طبقه‌بندی فرآیندی APQC نیز ساختار منسجمی را برای سازمان‌دهی اولیه فرآیندها در کنار پویایی تحلیل‌های هوشمند فراهم می‌سازد. پژوهش حاضر با هدف پر کردن خلأ نظری و کاربردی موجود، مدل یکپارچه‌ای برای کاربست هوش مصنوعی در BPM ارائه می‌دهد که گذار از فرآیندکاوی توصیفی گذشته‌نگر به مدیریت فرآیند تجویزی آینده‌نگر را محقق می‌سازد.

۲. بیان مسئله سازمان‌های امروزی در بستری پیچیده و به‌شدت پویا فعالیت می‌کنند که در آن چابکی عملیاتی برای بقا و پایداری ضروری است. وجود فرآیندهای ناکارآمد، تکراری یا ناهماهنگ میان واحدهای مختلف می‌تواند موجب اتلاف منابع، افت کیفیت خدمات و بروز نارضایتی در ذی‌نفعان شود. با این حال، چالش اصلی در نحوه ارزیابی و بهبود این فرآیندها نهفته است؛ رویکردهای سنتی در ارزیابی فرآیندها، مبتنی بر روش‌های دستی، بازبینی‌های

داده‌محور و مستقل از قضاوت‌های فردی نیاز دارند که بتواند با رصد دائم جریان کار، گلوگاه‌ها را آشکار کند. دوم، در سطح فناوریانه و علمی، این پژوهش با تلفیق الگوریتم‌های *Apriori* (برای کشف توالی‌ها)، *DBSCAN* (برای شناسایی انحرافات رفتاری و داده‌های پرت)، تحلیل شبکه‌های گرافی (برای مدل‌سازی تعاملات ساختاری) و تکنیک‌های پردازش زبان طبیعی (برای بلع دانش متنی)، شکاف میان ظرفیت‌های تئوریک هوش مصنوعی و کاربردهای عملیاتی را پوشش می‌دهد. سوم، از منظر تصمیم‌سازی، ابداع نوآوری‌هایی نظیر «ماتریس ارتباطی فرآیند-نقش» به مدیران نشان می‌دهد که تجمع بار کاری در کدام نقاط بحرانی است و دقیقاً چه پیشنهادات تجویزی برای بازتوزیع وظایف و بهینه‌سازی منابع وجود دارد.

۴. اهداف تحقیق هدف اصلی تحقیق: ارائه، تبیین و اعتبارسنجی یک مدل یکپارچه، هوشمند و چندلایه برای کاربردهای الگوریتم‌های هوش مصنوعی در مدیریت فرآیندهای کسب‌وکار سازمانی به‌منظور استقرار یک سیستم تصمیم‌یار پویا و تجویزی است.

اهداف فرعی و تفصیلی تحقیق: ۱. پیش‌پردازش و یکپارچه‌سازی لاگ‌های رویداد ساختاریافته و متون غیرساختاریافته سازمانی. ۲. کشف قوانین وابستگی و الگوهای تکراری فرآیندها با استفاده از الگوریتم *Apriori*. ۳. خوشه‌بندی رفتاری پرونده‌ها و شناسایی رفتارهای پرت و انحرافی با الگوریتم *DBSCAN*. ۴. تحلیل ساختاری شبکه فرآیندها بر پایه نظریه گراف

دوره‌ای نمونه‌محور و تکیه بر شاخص‌های تاریخی هستند. این روش‌ها توانایی انطباق با تغییرات لحظه‌ای و پردازش حجم عظیمی از داده‌های فرآیندی را ندارند و در تشخیص دقیق روابط پنهان بین‌بخشی و شناسایی ریشه‌ای گلوگاه‌ها ناتوان هستند.

در این میان، هوش مصنوعی به عنوان پارادایمی نوین در بازطراحی فرآیندها مطرح شده است. اما مسئله و خلأ اصلی پژوهش، نبود یک مدل جامع، یکپارچه و چندلایه است که بتواند تکنیک‌های مختلف هوش مصنوعی را برای تحلیل، بهینه‌سازی و تصمیم‌سازی فرآیندی جمع کند. در حال حاضر، سازمان‌ها در کاربردهای واقعی هوش مصنوعی با سردرگمی مواجه‌اند؛ زیرا ابزارها به صورت پراکنده استفاده می‌شوند و پیوند مستحکمی میان این ابزارهای ریاضیاتی و چارچوب‌های استاندارد فرآیندی (مانند *APQC*) وجود ندارد. بنابراین، مسئله اصلی این تحقیق تبیین این موضوع است که چگونه می‌توان با تلفیق الگوریتم‌های هوشمند (شامل کشف قوانین وابستگی با *Apriori*، خوشه‌بندی رفتاری با *DBSCAN*، تحلیل شبکه‌های گرافی و پردازش زبان طبیعی)، مدلی یکپارچه، هوشمند و تجویزی طراحی کرد که سازمان را در اصلاح و بازطراحی فرآیندها یاری دهد و هم‌زمان با استانداردهای ساختاری همسو باشد.

۳. اهمیت و ضرورت تحقیق ارزش واقعی داده‌ها در توانایی سازمان برای تبدیل کلان‌داده‌های خام فرآیندی به تصمیمات تجویزی نهفته است. اهمیت و ضرورت این پژوهش را می‌توان در سه لایه تبیین نمود: نخست، در سطح مدیریتی، سازمان‌ها شدیداً به مکانیزم‌هایی عینی،

معناشناسی حاکم بر مستندات سازمانی را جهت کشف دانش فرآیندی استخراج کنند؟ ۶. ماتریس ارتباطی فرآیند-نقش چگونه چالش تجمع بار کاری را تصویر کرده و پیشنهادات بهینه‌سازی ارائه می‌دهد؟ ۷. چگونه می‌توان خروجی‌های هوش مصنوعی را با چارچوب بین‌المللی APQC به عنوان یک ساختار پشتیبان هم‌راستا کرد؟ ۸. کارایی عملیاتی و شاخص‌های کلیدی عملکرد مدل پیشنهادی پس از پیاده‌سازی روی داده‌های واقعی تا چه حد ارتقا می‌یابند؟

بخش سوم: مبانی نظری و پیشینه پژوهش

۱. **مبانی نظری مدیریت فرآیندهای کسب‌وکار و فرآیندکاوی** مدیریت فرآیندهای کسب‌وکار به عنوان یک متدولوژی نظام‌مند و جامع، بر طراحی، مدل‌سازی، اجرا، پایش و بهینه‌سازی مستمر فرآیندهای سازمانی تمرکز دارد تا هم‌سویی عملیات با اهداف راهبردی سازمان را تضمین نماید. چرخه عمر سنتی این رویکرد شامل فازهای پیوسته‌ای نظیر شناسایی، کشف، تحلیل، طراحی مجدد، پیاده‌سازی و پایش است. با این حال، اجرای این چرخه در سازمان‌های امروزی با چالش‌های ساختاری مواجه شده است. رویکردهای کلاسیک ارزیابی فرآیندها که بر مستندسازی‌های دستی، زبان‌های مدل‌سازی ایستا و ممیزی‌های دوره‌ای متکی هستند، توانایی انطباق با جریان پویای تغییرات و پردازش حجم کلان داده‌های تولیدشده در سیستم‌های اطلاعاتی مدرن را ندارند. این روش‌های دستی نه تنها به شدت زمان‌بر و هزینه‌بر هستند، بلکه به دلیل اتکا بر قضاوت‌های ذهنی

جهت شناسایی ریاضیاتی گلوگاه‌ها و چرخه‌های مخرب پینگ‌پونگی. ۵. استخراج دانش فرآیندی پنهان از متون غیرساختاریافته سازمانی با تکنیک‌های پردازش زبان طبیعی (NLP). ۶. طراحی «ماتریس ارتباطی فرآیند-نقش» جهت بازتوزیع هوشمند بار کاری منابع سازمانی. ۷. هم‌راستاسازی خروجی‌های هوشمند مدل پیشنهادی با چارچوب بین‌المللی APQC به عنوان پشتیبان ساختاری. ۸. پیاده‌سازی عملیاتی و ارزیابی عملکرد مدل بر اساس داده‌های واقعی و سنجش‌های استاندارد یادگیری ماشین.

۵. **سوالات پژوهش** سوال اصلی تحقیق: مدل یکپارچه، هوشمند و چندلایه کاربست الگوریتم‌های هوش مصنوعی در مدیریت فرآیندهای کسب‌وکار سازمانی چگونه طراحی، تبیین و ارزیابی می‌شود و این مدل چگونه می‌تواند بهینه‌سازی عملیاتی را از طریق ابزارهای تصمیم‌یار محقق سازد؟

سوالات فرعی و تفصیلی تحقیق: ۱. مکانیزم استاندارد استخراج، پالایش و آماده‌سازی داده‌های ناهمگون سازمانی برای ورود به لایه‌های تحلیلی هوش مصنوعی چیست؟ ۲. الگوریتم Apriori چگونه می‌تواند روابط پنهان میان فعالیت‌ها و گام‌های فاقد ارزش افزوده را شناسایی کند؟ ۳. الگوریتم DBSCAN با چه مکانیزمی مسیرهای فرآیندی مشابه را خوشه‌بندی کرده و رفتارهای پرت را تفکیک می‌نماید؟ ۴. تحلیل شبکه فرآیندها بر پایه نظریه گراف، چگونه به شناسایی دقیق گلوگاه‌ها و رفتارهای رفت‌وبرگشتی منجر می‌شود؟ ۵. تکنیک‌های پردازش زبان طبیعی (NLP) چگونه می‌توانند

۲. تکنیک‌ها و الگوریتم‌های هوشمند در مدل

پیشنهادی تحقق یک سیستم فرآیندکاوی هوشمند و تجویزی، نیازمند طراحی یک پایپ‌لاین منسجم محاسباتی است که ابعاد مختلف داده‌های سازمانی را پوشش دهد. در لایه کشف الگوهای تکراری و روابط پنهان، الگوریتم ریاضیاتی **Apriori** جایگاه ویژه‌ای دارد. این الگوریتم با بهره‌گیری از خاصیت ضدیکنواختی، فضای جستجوی کلان‌داده‌های فرآیندی را به‌شدت کاهش می‌دهد. اصل ضدیکنواختی بیان می‌دارد که اگر یک مجموعه فعالیت در فرآیند پرتکرار نباشد، هیچ‌یک از ابرمجموعه‌های آن نیز پرتکرار نخواهند بود. این ویژگی به الگوریتم اجازه می‌دهد تا توالی‌های موازی زائد و گام‌های بروکراتیک فاقد ارزش افزوده را با محاسبه شاخص‌های آماری پشتیبانی، اطمینان و ارتقا کشف کند.

در لایه خوشه‌بندی رفتاری، الگوریتم خوشه‌بندی مبتنی بر چگالی (**DBSCAN**) ابزاری بسیار کارآمد برای مقابله با پدیده مدل‌های درهم‌تنیده و اسپاگتی‌گونه است. فرآیندهای واقعی به دلیل تنوع رفتار کاربران دارای ناهمگونی بالایی هستند. الگوریتم **DBSCAN** برخلاف روش‌های سنتی که نیازمند تعیین پیش‌فرض تعداد خوشه‌ها هستند، خوشه‌ها را بر اساس مناطق با چگالی بالای داده‌ها تشکیل می‌دهد. مزیت منحصر به فرد این الگوریتم، توانایی ذاتی آن در شناسایی خودکار نویزها و داده‌های پرت است؛ این ویژگی به سیستم اجازه می‌دهد تا رویه‌های استاندارد را از انحرافات عملیاتی و رفتارهای پرت

و نمونه‌گیری‌های محدود، در شناسایی دقیق گلوگاه‌ها و رفتارهای انحرافی ناتوان عمل می‌کنند.

برای غلبه بر این چالش‌ها، فرآیندکاوی به عنوان دانشی میان‌رشته‌ای در تقاطع علوم داده و مدیریت فرآیند ظهور کرده است. فرآیندکاوی با استخراج ردپای دیجیتالی سازمانی از دل لاگ‌های رویداد ثبت‌شده در سامانه‌های اطلاعاتی، تصویری کاملاً عینی و مبتنی بر شواهد از وضعیت موجود ارائه می‌دهد. این رویکرد در سه حوزه کلان شامل کشف فرآیند (استخراج مدل واقعی)، بررسی انطباق (مقایسه واقعیت با رویه‌های استاندارد) و ارتقا (تزریق داده‌های عملکردی به مدل جهت بهبود) تقسیم می‌شود. با توسعه الگوهای پیشرفته، فرآیندکاوی از ابزاری گذشته‌نگر و توصیفی به سمت سیستم‌های پیش‌بینانه و تجویزی حرکت کرده است؛ رویکردی که نه تنها مشکلات آینده را پیش‌بینی می‌کند، بلکه بهترین اقدام بعدی را برای جلوگیری از خطا یا بهینه‌سازی تخصیص منابع به مدیران پیشنهاد می‌دهد.

جهت چارچوب‌مند ساختن این بهبودهای پویا، استفاده از چارچوب‌های استاندارد بین‌المللی نظیر چارچوب طبقه‌بندی فرآیندی مرکز بهره‌وری و کیفیت آمریکا (**APQC**) نقشی حیاتی ایفا می‌کند. این چارچوب با ارائه یک ساختار سلسله‌مراتب استاندارد برای طبقه‌بندی وظایف سازمانی، زبانی مشترک میان تحلیل‌های فنی هوش مصنوعی و بدنه مدیریتی سازمان ایجاد می‌کند، به طوری که بهبودهای هوشمند همواره در بستر ساختاری پایدار و قابل الگوبرداری سازمان‌دهی می‌شوند. لایه مدل‌سازی گراف و تحلیل شبکه، فرآیندها را از حالت زنجیره‌های خطی خارج کرده و به شکل یک شبکه

است؛ هرچند این مدل‌ها غالباً در درک دقیق منطق دروازه‌های تصمیم‌گیری موازی یا انحصاری با چالش مواجه بوده‌اند که محققان سعی کرده‌اند با طراحی استراتژی‌های مهندسی پرامیت پیشرفته این نقص را برطرف کنند. در محور الگوریتم‌های یادگیری ماشین، کاربرد الگوریتم‌های کاوش قوانین وابستگی نظیر Apriori و الگوریتم‌های خوشه‌بندی نظیر DBSCAN به‌طور مجزا برای هرس کردن مسیرهای پیچیده، شناسایی تأخیرها و خوشه‌بندی واریانت‌های رفتاری با موفقیت ارزیابی شده است. همچنین در حوزه حاکمیت سازمانی، تلاش‌هایی برای پیوند دادن چارچوب‌های کنترلی با سیستم‌های تصمیم‌یار جهت پایش ریسک و انطباق در زمان واقعی انجام گرفته است.

با وجود این دستاوردها، بررسی انتقادی ادبیات تحقیق نشان‌دهنده وجود پنج شکاف پژوهشی اساسی است که مبنای نوآوری این رساله قرار گرفته است: شکاف نخست، فقدان یک مدل ترکیبی و یکپارچه میان تحلیل‌های داده‌محور هوش مصنوعی و چارچوب‌های استاندارد سازمانی است. مطالعات قبلی یا منحصراً به مدل‌سازی‌های دستی پرداخته‌اند یا صرفاً داده‌کاوی را مدنظر قرار داده‌اند، بدون آنکه پیوندی پویا میان الگوریتم‌ها با چارچوب‌های طبقه‌بندی استاندارد نظیر APQC جهت ایجاد یک زبان مشترک سازمانی برقرار سازند.

شکاف دوم، نادیده گرفتن بخش عظیمی از دانش فرآیندی پنهان در اسناد غیرساختاریافته سازمانی است. اکثر تحلیل‌های فرآیندکاوی تنها بر داده‌های

تفکیک نماید. تعاملی پیچیده از گره‌ها (فعالیت‌ها یا نقش‌ها) و یال‌ها (جریان انتقال کار) ترسیم می‌کند. اعمال شاخص‌های مرکزیت در نظریه گراف، به‌ویژه شاخص مرکزیت بینابینی، به سازمان امکان می‌دهد تا گره‌هایی را که بیشترین نقش واسطه‌ای را در جریان فرآیند دارند و مستعد تبدیل شدن به گلوگاه‌های ساختاری هستند، شناسایی کند. همچنین، تحلیل شبکه‌ای به کشف چرخه‌های مخرب تعاملی نظیر رفتارهای رفت‌وبرگشتی یا پینگ‌پونگی میان نقش‌های سازمانی کمک می‌کند که ناشی از ابهام در شرح وظایف یا نقص مستندات است.

در نهایت، لایه پردازش زبان طبیعی (NLP) و مدل‌های زبانی، با پردازش متون غیرساختاریافته سازمانی نظیر آیین‌نامه‌ها و دستورالعمل‌های استاندارد عملیاتی، دانش ضمنی فرآیندها را استخراج می‌کنند. مدل‌های زبانی با اجرای سه گام تشخیص موجودیت‌ها، وضوح موجودیت (رفع ابهام از ارجاعات متنی) و استخراج روابط، مفاهیم خام متنی را به گراف‌های رویه‌ای ماشین‌خوان تبدیل می‌کنند تا انطباق مستندات قانونی با واقعیت‌های عملیاتی لاگ‌های رویداد ارزیابی شود.

۳. مرور پیشینه پژوهش و شکاف‌های تحقیقاتی

مرور مطالعات انجام‌شده در حوزه‌های تخصصی مرتبط نشان می‌دهد که پژوهش‌های پیشین دستاوردهای ارزشمندی در زمینه فرآیندکاوی و یادگیری ماشین داشته‌اند. در محور کاربرد مدل‌های زبانی، تلاش‌های متعددی برای استخراج مدل‌های فرآیندی از متون با استفاده از تکنیک‌های پردازش زبان طبیعی صورت گرفته

۱. پارادایم، رویکرد و استراتژی پژوهش پژوهش حاضر از منظر جهت‌گیری فلسفی زیر چتر پارادایم «عمل‌گرایی» قرار دارد. دلیل اتخاذ این پارادایم، تمرکز آن بر حل مسائل واقعی، ملموس و پیچیده در زیست‌بوم سازمانی و ارائه راهکارهای کارآمد است. این رویکرد به تحقیق اجازه می‌دهد تا تئوری‌های ریاضیاتی و فنی هوش مصنوعی را با پویایی‌های مدیریتی فرآیندهای کسب‌وکار پیوند دهد. از منظر رویکرد پژوهشی، این رساله مبتنی بر متدولوژی «ترکیبی» (کمی و کیفی) است؛ زیرا فرآیندهای سازمانی دارای ابعاد دوگانه‌ای هستند که از یک سو شامل داده‌های کلان ساختاریافته (کمی) نظیر لاگ‌های رویداد و رکوردهای زمانی بوده و از سوی دیگر مشتمل بر مستندات، آیین‌نامه‌ها و دانش ضمنی غیرساختاریافته (کیفی) می‌باشند که نیازمند تحلیل‌های معنایی هستند.

در لایه استراتژی تحقیق، با توجه به هدف نهایی پژوهش که خلق یک مصنوع نوین (مدل چندلایه تصمیم‌یار) است، از استراتژی «تحقیق علم طراحی» (DSR) استفاده شده است. بر اساس این استراتژی، مسیر توسعه مدل در قالب پنج فاز پیوسته و تکرارپذیر اجرا می‌گردد: نخست، شناخت و تبیین مسئله (ناکارآمدی عارضه‌یابی‌های دستی سنتی)؛ دوم، تعریف اهداف راهکار هوشمند؛ سوم، طراحی و توسعه معماری چندلایه مدل؛ چهارم، پیاده‌سازی عملیاتی در بستر نرم‌افزاری پایتون روی کلان‌داده‌های واقعی؛ و پنجم،

عددی و ساختاریافته‌ی لاگ‌های رویداد تکیه دارند و از ظرفیت پردازش زبان طبیعی برای بلع اسناد قانونی و دستورالعمل‌ها استفاده نکرده‌اند.

شکاف سوم، تمرکز انحصاری پژوهش‌ها بر سطح پایه و عملیاتی فرآیندهاست. بیشتر ابزارها صرفاً برای پایش روزمره جریان کار طراحی شده‌اند و کمتر مدلی خروجی تحلیل‌ها را به تصمیم‌سازی‌های تاکتیکی و راهبردی کلان سازمانی متصل کرده است.

شکاف چهارم، نبود یک سیستم تصمیم‌یار و پیشنهاددهنده هوشمند و داده‌محور برای بازطراحی فرآیندهاست. سیستم‌های موجود در حوزه بهبود فرآیند غالباً به ارائه تحلیل‌های گذشته‌نگر یا استفاده از قوانین ایستا و ساده بسنده کرده و فاقد لایه تجویزی پویا برای بازتولید سناریوهای بهینه عملیاتی هستند.

شکاف پنجم، ضعف در مدل‌سازی ترکیبی و تعاملی میان چند الگوریتم است. تحقیقات گذشته معمولاً کارایی یک الگوریتم مجزا را ارزیابی کرده‌اند و جریان پردازشی یکپارچه‌ای (پایپ‌لاین) که در آن خروجی‌های کشف الگو، خوشه‌بندی رفتاری، تحلیل گرافی شبکه و پردازش زبان طبیعی به صورت متوالی برای رفع نقاط ضعف یکدیگر ترکیب شوند، در ادبیات علمی مغفول مانده است.

رساله حاضر با طراحی و اعتبارسنجی یک مدل چندلایه و منسجم، دقیقاً در تقاطع این شکاف‌های علمی قرار گرفته است تا پیوندی استوار میان ابزارهای هوشمند هوش مصنوعی و تعالی عملیاتی سازمان‌ها ایجاد نماید.

بخش چهارم: روش‌شناسی پژوهش

سازمانی بلع شده و پس از پاک‌سازی نویزها، رفع رکوردهای مفقود یا نامعتبر و استانداردسازی برچسب‌های زمانی، به فرمت‌های معتبر فرآیندکاوی (نظیر XES) تبدیل می‌شوند. ب) داده‌های غیرساختاریافته: مستندات قانونی و آیین‌نامه‌های متنی گردآوری شده و با به‌کارگیری مدل‌های زبانی پیشرفته در پردازش زبان طبیعی (NLP)، موجودیت‌های فرآیندی (فعالیت‌ها و کنشگران) و روابط منطقی بین آن‌ها (توالی‌ها و دروازه‌های تصمیم‌گیری) به گراف‌های روبه‌ای ماشین‌خوان تبدیل می‌شوند. خروجی نهایی این دو مسیر در یک پایگاه داده تحلیلی تجمیع می‌گردد.

۴. خط لوله پنج‌فازی پیاده‌سازی مدل معماری اجرایی مدل پیشنهادی در قالب یک خط لوله پردازشی پیوسته و متوالی در پایتون پیکربندی شده است:

- فاز اول (بلع و پیش‌پردازش): تجمیع و پاک‌سازی لاگ‌های رخداد و استخراج دانش فرآیندی از متون با NLP.
- فاز دوم (کشف الگو): به‌کارگیری الگوریتم Apriori بر پایه اصل ضدیکنواختی برای هرس کردن مسیرهای اسپاگتی و شناسایی قوانین وابستگی، توالی‌های موازی زائد و گام‌های فاقد ارزش عملیاتی.
- فاز سوم (خوشه‌بندی رفتاری): استفاده از الگوریتم چگالی‌محور DBSCAN بدون نیاز

ارزیابی و اعتبارسنجی مصنوع با شاخص‌های فنی یادگیری ماشین و سنجه‌های بهبود فرآیند.

۲. جامعه آماری، نمونه‌گیری و واحد تحلیل جامعه آماری در این پژوهش داده‌محور شامل دو بخش اصلی است: نخست، داده‌های دیجیتال ثبت‌شده در سامانه‌های اطلاعاتی سازمان (مانند لاگ‌های رویداد سیستم‌های مدیریت فرآیند یا برنامه‌ریزی منابع سازمانی)؛ و دوم، مستندات رسمی، آیین‌نامه‌ها، کتابچه‌های راهنمای کاربری و دستورالعمل‌های استاندارد عملیاتی سازمان. استراتژی نمونه‌گیری در این تحقیق به‌جای روش‌های محدود سنتی، مبتنی بر بررسی تمام‌شمار کل داده‌های ثبت‌شده در یک پنجره زمانی مشخص است تا تمامی انحرافات و رفتارهای استثنایی سیستم رصد شوند.

واحد تحلیل نیز در سه سطح کلان، خرد و کیفی تعریف شده است. در لایه کلان فرآیندکاوی، واحد تحلیل «مورد» یا «ردپا» است که نشان‌دهنده یک جریان کامل از آغاز تا پایان فرآیند است. در لایه خرد، واحد تحلیل به «رویداد» یا «انتقال بین مرحله‌ای» تقلیل می‌یابد تا اعمال الگوریتم‌های خوشه‌بندی میسر شود. در لایه کیفی نیز، واحد تحلیل شامل «قطعات متنی و جملات» اسناد سازمانی است که برای استخراج نقش‌ها و فعالیت‌ها مورد پردازش قرار می‌گیرند.

۳. روش‌ها و ابزارهای گردآوری داده‌ها فرآیند گردآوری داده‌ها با اتخاذ یک معماری دوگانه برای بلع داده‌های ناهمگون طراحی شده است: الف) داده‌های ساختاریافته: با تعریف خطوط لوله استخراج، تبدیل و بارگذاری (ETL)، ردپای دیجیتال خام از پایگاه‌های داده

دقت، فراخوانی و امتیاز اف-یک ارزیابی شده و در نهایت کارایی موتور تجویزی با شاخص‌های زودرس بودن هشدارها، نرخ شکست و بازخوردهای تجربی خبرگان سازمانی مورد تأیید قرار می‌گیرد.

بخش پنجم: طراحی و معماری مدل پیشنهادی

۱. چارچوب مفهومی و معماری کلان مدل

پیشنهادی معماری کلان مدل پیشنهادی در این پژوهش بر پایه گذار از رویکردهای ایستا و تک‌بعدی ارزیابی فرآیندها به سمت یک سیستم پویای یکپارچه و افزوده شده با هوش مصنوعی استوار است. برای تحقق تعالی عملیاتی و پویایی در سیستم‌های سازمانی، چارچوب مفهومی مدل از یک معماری ماژولار بهره می‌برد. تجزیه ماژولار سبب کاهش وابستگی‌های پیچیده میان اجزاء، افزایش چابکی تحلیلی و ارتقای مقیاس‌پذیری سیستم می‌شود. این مدل پیشنهادی، جریان داده‌های خام سازمان را طی یک خط لوله پردازشی چندلایه به دانش مدیریتی و تصمیم‌های تجویزی تبدیل می‌کند.

یکی از برجسته‌ترین ویژگی‌های این معماری، تفکیک ساختاریافته لایه‌های محاسباتی و ریاضیاتی (پیش‌پردازش، کشف الگو، خوشه‌بندی و تحلیل شبکه) از لایه تصمیم‌سازی و استدلال تجویزی است. این ویژگی تضمین می‌کند که فرآیند تصمیم‌گیری دچار توهم داده‌ای نشده و تمامی راهکارهای بهینه‌سازی پیشنهادی، ریشه در شواهد عینی، ممیزی‌پذیر و کاملاً ریاضیاتی داشته باشند. همچنین، جهت برقراری یک زبان مشترک سازمانی، خروجی الگوریتم‌ها با چارچوب

به فرضیه قبلی برای تعداد خوشه‌ها، جهت گروه‌بندی مسیرهای مشابه در قالب خوشه‌های همگن و تفکیک خودکار پرونده‌های پرت و ناهنجار به عنوان نویز.

• فاز چهارم (مدل‌سازی گرافی و تحلیل شبکه):

تبدیل خوشه‌ها به گراف‌های جهت‌دار وزن‌دار و اعمال شاخص‌های مرکزیت شبکه (به‌ویژه مرکزیت بینابینی) جهت شناسایی ریاضیاتی گلوگاه‌ها و چرخه‌های مخرب پینگ‌پونگی میان نقش‌ها.

• فاز پنجم (موتور تجویزی و سیستم تصمیم‌یار):

تجمیع نتایج مراحل قبل و به‌کارگیری «ماتریس ارتباطی فرآیند-نقش» با همسو کردن خروجی‌ها با چارچوب مرجع APQC، جهت ارائه سناریوهای بهینه بازطراحی، بازتوزیع بار کاری و ارائه هشدارهای زودهنگام به مدیران.

۵. شاخص‌های اعتبارسنجی و ارزیابی مدل برای

تضمین روایی و کارایی این مصنوع چندلایه، از معیارهای سخت‌گیرانه ارزیابی استفاده می‌شود. کیفیت الگوریتم Apriori با سه سنجه پشتیبانی، اطمینان و ارتقا سنجیده می‌شود. پایداری خوشه‌های الگوریتم DBSCAN با شاخص‌های سیلوئت و رند تعدیل‌شده ارزیابی می‌گردد. صحت مدل‌های گرافی استخراج‌شده با چهار سنجه فرآیندکاوی شامل برازش (تناسب بازتولید رفتاری)، دقت (جلوگیری از رفتارهای غیرمجاز)، تصمیم‌پذیری و سادگی گراف سنجیده می‌شود. در لایه پردازش متن (NLP) عملکرد مدل زبانی با معیارهای

شده است. به دلیل ابهامات و پیچیدگی‌های نگارشی زبان طبیعی، مدل‌های زبانی هوش مصنوعی در قالب سه گام متوالی پیاده‌سازی می‌شوند: نخست، «تشخیص موجودیت‌ها» جهت شناسایی فعالیت‌ها، کنشگران و داده‌های فرآیندی؛ دوم، «وضوح موجودیت» برای رفع ابهام از ضمیر و یکپارچه‌سازی مفاهیم هم‌معنی؛ و سوم، «استخراج روابط» جهت کشف تعاملات، توالی‌ها و دروازه‌های منطقی تصمیم‌گیری (مانند انشعابات موازی و انحصاری). هدایت این مدل‌ها به کمک استراتژی‌های پیشرفته مهندسی پرامپت نظیر زنجیره تفکر صورت می‌گیرد تا گراف‌های رویه‌ای ماشین‌خوان از متون خام استخراج شوند. در نهایت، خروجی‌های تصفیه‌شده این دو مسیر در یک پایگاه داده تحلیلی یکپارچه تجمیع می‌گردند.

۳. لایه دوم: کشف قوانین وابستگی و الگوهای تکراری پس از تجمیع داده‌ها در لایه نخست، لایه دوم با استفاده از الگوریتم‌های داده‌کاوی بدون نظارت تلاش می‌کند الگوهای پنهان در رفتار واقعی فرآیندها را کشف نماید. در فرآیندهای واقعی، تنوع مسیرها منجر به شکل‌گیری مدل‌های اسپاگتی‌گونه و بسیار مبهم می‌شود. برای حل این مسئله و هرس کردن درخت جستجوی توالی‌های ممکن، از الگوریتم ریاضیاتی *Apriori* استفاده می‌شود. این الگوریتم بر پایه اصل ضدیک‌نواختی عمل می‌کند؛ بدین معنی که اگر یک زیرمسیر فرآیندی در کل داده‌ها پرتکرار نباشد، هیچ‌یک از ابرمجموعه‌های آن نیز پرتکرار نخواهند بود.

استاندارد بین‌المللی APQC هم‌راستا می‌شود. چارچوب APQC به عنوان یک ساختار مرجع و پشتیبان عمل می‌کند تا خروجی‌های هوشمند هوش مصنوعی (نظیر پیشنهاد ادغام یا حذف فعالیت‌ها) در قالب‌های استاندارد و قابل درک برای مدیران ارشد آدرس‌دهی و طبقه‌بندی شوند.

۲. لایه اول: بلع داده‌های ناهمگون و پیش‌پردازش کیفیت خروجی‌های یک سیستم هوشمند به شدت به کیفیت داده‌های ورودی آن وابسته است. از آنجا که دانش فرآیندی سازمان‌ها به صورت ناهمگون در داده‌های ساختاریافته (لاگ‌های سیستمی) و غیرساختاریافته (متون و آیین‌نامه‌ها) پراکنده است، این لایه با اتخاذ یک ساختار دوگانه طراحی شده است:

الف) مسیر داده‌های ساختاریافته: در این مسیر، ردپای دیجیتال عملیات سازمانی از طریق خطوط لوله استخراج، تبدیل و بارگذاری (ETL) از پایگاه‌های داده بلع می‌شود. یک لاگ رویداد استاندارد فرآیندکاوی باید حداقل شامل شناسه پرونده، نام فعالیت و برچسب زمان وقوع باشد. فرآیند پیش‌پردازش در این بخش شامل پاک‌سازی نویزها، حذف رکوردهای تکراری، مدیریت رکوردهای مفقود و برطرف کردن چالش‌های ساختار داده‌های شیء‌محور (مانند تشخیص شناسه پرونده هدف در مدل‌های چندشیئی) است تا داده‌ها به فرمت‌های استاندارد ماشین‌خوان برای فرآیندکاوی تبدیل شوند.

ب) مسیر داده‌های غیرساختاریافته: این مسیر با هدف استخراج دانش فرآیندی پنهان در اسناد متنی، آیین‌نامه‌ها و دستورالعمل‌های استاندارد عملیاتی طراحی

صرفاً بر اساس مناطقی با چگالی بالای داده‌ها (با تنظیم پارامترهای شعاع همسایگی اپسیلون و حداقل نقاط) تشکیل می‌دهد. مزیت منحصربه‌فرد این الگوریتم در فرآیند کاوی، توانایی ذاتی آن در شناسایی خودکار نویزها و رفتارهای پرت به عنوان داده‌های خارج از محدوده چگال است. با این مکانیزم، سیستم می‌تواند جریان‌های استاندارد و بهینه سازمان را در قالب خوشه‌های اصلی گروه‌بندی کند و هم‌زمان انحرافات عملیاتی، خطاهای سیستمی و مسیرهای استثنایی نادر را بدون دخالت دستی تفکیک نماید.

۵. لایه چهارم: مدل‌سازی گرافی و تحلیل شبکه فرآیندها در لایه چهارم، خوشه‌ها و الگوهای رفتاری حاصل از لایه‌های قبلی به گراف‌های رویه‌ای جهت‌دار و وزن‌دار تبدیل می‌شوند تا فرآیندها نه به صورت زنجیره‌های خطی ساده، بلکه در قالب یک شبکه تعاملی پیچیده مدل‌سازی شوند. در این ساختار گرافی، فعالیت‌ها یا نقش‌ها به عنوان گره‌ها و جریان انتقال کار یا اطلاعات میان آن‌ها به عنوان یال‌ها تعریف شده و وزن یال‌ها متناسب با فراوانی انتقال‌ها تنظیم می‌شود. برای محاسبه صحیح کوتاه‌ترین مسیرهای ارتباطی، وزن یال‌ها به صورت معکوس فراوانی انتقال فرمول‌بندی می‌شود تا مسیرهای پرتردد به عنوان مسیرهای بهینه شناخته شوند.

اعمال شاخص‌های تحلیل شبکه‌های اجتماعی (SNA) و نظریه گراف بر روی این شبکه، امکان کشف ریاضیاتی گلوگاه‌ها را فراهم می‌سازد:

این الگوریتم با پیمایش ردپاها، مجموعه‌های پرتکرار را کشف کرده و قواعد وابستگی را به صورت قوانین شرطی (اگر-آنگاه) استخراج می‌کند. قدرت و معناداری این قوانین بر اساس سه شاخص آماری کلیدی ارزیابی و فیلتر می‌شوند: «پشتیبانی» که میزان فراگیری و شمولیت قاعده را در کل پرونده‌ها می‌سنجد؛ «اطمینان» که میزان قطعیت و قابلیت اتکای رابطه را نشان می‌دهد؛ و «ارتقا» که همبستگی واقعی میان دو فعالیت را با خنثی کردن اثر تصادفی بودن وقوع آن‌ها ارزیابی می‌کند. قوانین استخراج‌شده‌ای که دارای ارتقای بزرگتر از یک و اطمینان بالا باشند، نشان‌دهنده چرخه‌های بازکاری سیستماتیک، کارهای موازی زائد یا بروکرسی‌های غیرضروری در سازمان هستند که به عنوان کاندیداهای بازطراحی شناسایی می‌شوند.

۴. لایه سوم: خوشه‌بندی رفتاری و شناسایی انحرافات فرآیندهای سازمان‌ها به دلیل تفاوت در رفتار کاربران و شرایط عملیاتی دارای ناهمگونی بالایی هستند. اعمال یک سیاست یکسان برای کل سازمان به دلیل این ناهمگونی با شکست مواجه می‌شود. لایه سوم با به‌کارگیری الگوریتم خوشه‌بندی رفتاری چگالی‌محور (DBSCAN)، پرونده‌ها را بر اساس شباهت در ویژگی‌های رفتاری (نظیر طول مسیر، زمان چرخه، نرخ بازکاری و وضعیت‌های بحرانی) به گروه‌های همگن تقسیم می‌کند.

بر خلاف روش‌های سنتی مانند K-Means که نیازمند پیش‌فرض قرار دادن تعداد خوشه‌ها هستند و به داده‌های پرت حساس می‌باشند، الگوریتم DBSCAN خوشه‌ها را

الگوریتم‌های تخصیص پویای منابع با ارزیابی این ماتریس، پیشنهادات صریحی را برای بازتوزیع نقش‌ها، کاهش بار کاری گره‌های پرتراфик و بهینه‌سازی چیدمان منابع انسانی ارائه می‌دهند. در سطح کلان، موتور تجویزی بر اساس تلفیق قواعد *Apriori*، تحلیل‌های گرافی و مستندات قانونی استخراج‌شده با *NLP*، سناریوهای بازطراحی فرآیند را در قالب چهار استراتژی استخراج می‌کند: ۱. حذف وظایف: حذف گام‌های تکراری و فاقد ارزش افزوده‌ای که در الگوریتم *Apriori* شناسایی شده‌اند. ۲. تجمیع فعالیت‌ها: ادغام وظایف متوالی خرد که پایداری بالایی در خوشه‌های *DBSCAN* دارند. ۳. تغییر توالی و موازی‌سازی: تبدیل ساختارهای خطی طولانی به مسیرهای موازی و هم‌زمان جهت کاهش زمان کل چرخه. ۴. بهترین اقدام بعدی: ارائه هشدارهای زودهنگام و پیشنهادهای زنده عملیاتی به کاربران سیستم جهت پیشگیری از انحرافات فرآیندی. این بستر تحلیلی در نهایت تصمیم‌گیری‌های فرآیندی را از حالت حدس و گمان شهودی به یک فرآیند کاملاً علمی، مهندسی‌شده و مبتنی بر شواهد داده‌ای تبدیل می‌کند.

بخش ششم: پیاده‌سازی و یافته‌های تجربی

۱. محیط اجرا، ابزارها و دیتابیس ارزیابی به‌منظور اثبات کارآمدی و ارزیابی تجربی مدل چندلایه پیشنهادی، پایپ‌لاین طراحی‌شده به‌طور کامل در محیط برنامه‌نویسی پایتون توسعه یافته و پیاده‌سازی گردید. در این راستا، از کتابخانه‌های تخصصی برای اجرای وظایف محاسباتی استفاده شد؛ الگوریتم کاوش

- شاخص درجه گره: میزان فعالیت و درگیری یک واحد یا ایستگاه کاری را در تعاملات شبکه نشان می‌دهد.
- شاخص مرکزیت بینابینی: میزان واسطه‌گری یک گره را در کل جریان کار می‌سنجد. گره‌هایی با مرکزیت بینابینی بالا که در گلوگاه شبکه قرار دارند، کانون‌های اصلی ایجاد تأخیر و انباشت کار در سازمان هستند. همچنین، این لایه گرافی به صورت خودکار قادر به کشف چرخه‌های مخرب تعاملی نظیر «رفتارهای پینگ‌پونگی» (پاس‌کاری مداوم و رفت‌وبرگشتی کار میان دو یا چند نقش) است که ریشه در ابهام وظایف یا نقص مدارک ورودی دارد.

۶. لایه پنجم: موتور استدلال تجویزی و سیستم

تصمیم‌یار لایه نهایی و اوج تکامل این مدل، تبدیل تمامی یافته‌های تحلیلی و شواهد قطعی کشف‌شده به توصیه‌های عملیاتی و تجویزی است. این مأموریت بر عهده «موتور استدلال تجویزی» قرار دارد که با ترکیب خروجی‌های مراحل پیشین، بسترهای بهبود پویا را در سه سطح پرونده (موارد بحرانی)، خوشه (واریانت‌های رفتاری) و کلان فرآیند صادر می‌کند. هسته نوآورانه این لایه، ابداع «ماتریس ارتباطی فرآیند-نقش» است. این ماتریس با نداشت تقاطعی فعالیت‌های فرآیند با نقش‌های سازمانی ایفاگر آن‌ها، سلول‌های خود را بر اساس فرکانس کاری و امتیاز گلوگاهی (مرکزیت بینابینی) مقداردهی می‌کند.

رخداد‌های فاقد برچسب زمانی، و تبدیل رشته‌های زمانی به فرمت استاندارد با موفقیت روی کل داده‌ها اجرا شد.

خروجی فاز اول منجر به تولید یک لاگ رویداد استاندارد و بسیار تمیز با مشخصات آماری زیر گردید: کل رخداد‌های پردازش‌شده برابر با ۱,۲۰۲,۲۶۷ رخداد، تعداد پرونده‌های یکتای بازسازی‌شده معادل ۳۱,۵۰۹ پرونده، میانگین طول مسیر فرآیندی ۷.۶۰ رخداد در هر پرونده (با میانه ۷.۰۰ و حداکثر طول مسیر ۲۱.۰۰ رخداد)، میانگین زمان چرخه هر پرونده ۵۲۳.۹۴ ساعت، میانگین فعالیت‌های یکتا در هر پرونده ۶.۸۱ فعالیت و میانگین منابع سازمانی درگیر در هر پرونده ۴.۱۳ منبع انسانی ثبت گردید که پایه‌ای کاملاً عاری از نویز را برای تغذیه الگوریتم‌های بعدی ایجاد کرد.

۳. یافته‌های فاز دوم: کشف قوانین وابستگی با

الگوریتم Apriori الگوریتم Apriori با هدف کشف الگوهای تکراری و گام‌های زائد بر روی توالی‌های استاندارد حاصل از ۳۱,۵۰۹ پرونده یکتا اعمال گردید. برای فیلتر کردن نویزهای آماری، آستانه حداقل پشتیبانی روی ۰.۱۰ و حداقل اطمینان روی ۰.۷۰ تنظیم شد. اجرای این الگوریتم منجر به کشف ۳۱۹ مجموعه فعالیت پرتکرار و ۳,۰۴۰ قانون وابستگی اولیه گردید که پس از پالایش نهایی برای ارتقای تفسیرپذیری، به ۴۰ قانون ساده و یک‌به‌یک تقلیل یافتند. بیشینه شاخص ارتقا (Lift) در این قوانین ساده برابر با ۱.۵۴۱۷ محاسبه شد که نشان‌دهنده

قوانین وابستگی با استفاده از کتابخانه تخصصی mlxtend و خوشه‌بندی چگالی‌محور با بهره‌گیری از کتابخانه Scikit-Learn پیاده‌سازی گردید. همچنین برای تبدیل داده‌ها به شبکه و تحلیل شاخص‌های نظریه گراف، کتابخانه NetworkX به کار گرفته شد. در لایه پردازش زبان طبیعی (NLP) و پردازش متون غیرساختاریافته نیز از مدل‌های ترانسفورمر در چارچوب HuggingFace جهت استخراج موجودیت‌ها و روابط استفاده شد.

ارزیابی تجربی الگوریتم‌ها بر روی کلان‌داده استاندارد و بین‌المللی فرآیند درخواست وام یک موسسه مالی هلندی (BPIC ۲۰۱۷) انجام گرفت. این مجموعه داده به دلیل برخورداری از حجم بالای رویدادها، تنوع رفتاری چشمگیر کاربران و ثبت دقیق رکوردهای زمانی، بستری بسیار چالش‌برانگیز و ایده‌آل را برای محک زدن پویایی مدل پیشنهادی فراهم ساخت تا نتایج مدل با استانداردهای جهانی قابل مقایسه باشد.

۲. یافته‌های فاز اول: آماده‌سازی و پیش‌پردازش

داده‌ها داده‌های خام به دست آمده از سیستم عملیاتی به دلیل ساختار شیء‌محور، پیچیدگی روابط و وجود ناهم‌هنگی در برچسب‌ها، مستقیماً قابل تحلیل نبودند. در گام نخست، با اجرای ماژول تشخیصی بر روی ساختار اشیاء، موجودیت «Application» با دارا بودن ۳۱,۵۰۹ مورد یکتا به عنوان مبنای قطعی شناسه پرونده انتخاب شد تا ساختار توالی‌ها بدون پدیده چندپارگی مصنوعی بازسازی شود. فرآیند پاک‌سازی و استانداردسازی شامل یکنواخت‌سازی نام فعالیت‌ها، حذف رکوردهای تکراری و

ارزنده‌ترین یافته این لایه، کشف ریاضیاتی چرخه‌های مخرب تعاملی موسوم به رفتارهای پینگ‌پونگی بود. سیستم موفق شد یک چرخه رفت‌وبرگشتی فوق‌العاده مخرب را میان دو فعالیت نقص اطلاعات و اعتبارسنجی کشف نماید. ابعاد این عارضه در خروجی‌های محاسباتی به این شرح ثبت گردید: تعداد کل وقوع برابر با ۲۴,۹۸۶ بار رفت‌وبرگشت، و تعداد پرونده‌های درگیر معادل ۱۱,۶۶۸ پرونده (۳۷.۰۳٪ کل پرونده‌های سازمان) با میانگین ۲.۱۴ بار تکرار چرخه باطل در پرونده‌های درگیر بود که نشان می‌دهد بازکاری یک نقص ساختاری غالب در سازمان است.

۵. یافته‌های فاز چهارم: خوشه‌بندی رفتاری پرونده‌ها با الگوریتم DBSCAN برای خوشه‌بندی رفتاری پرونده‌ها در سطح خرد، توالی رخداد‌های ۳۱,۵۰۹ پرونده بر اساس ۲۲ ویژگی تحلیلی (مهندسی‌شده از ابعاد زمانی، حجمی، پینگ‌پونگ و کدهای وضعیتی) به فضای برداری نگاشت و نرمال‌سازی شد. الگوریتم چگالی‌محور DBSCAN با تنظیم شعاع همسایگی ۱.۸ و حداقل نقاط ۳۰ بر روی ماتریس ویژگی‌ها اجرا گردید که حاصل آن شناسایی ۱۸ خوشه رفتاری استاندارد و یک خوشه ناهنجار (نویز) بود. توزیع و تفسیر این خوشه‌ها نشان داد که رفتارهای عملیاتی سازمان در چهار خانواده اصلی طبقه‌بندی می‌شوند:

- خوشه‌های بازکاری‌محور (۴۶.۰٪ کل پرونده‌ها): شامل ۱۴,۴۹۹ پرونده با میانگین

همبستگی مثبت بسیار قوی و غیرتصادفی بودن روابط کشف‌شده در سیستم است. تحلیل قوانین برتر نشان داد که فرآیند دارای یک «ناحیه اصطکاک بحرانی» است. الگوریتم موفق به کشف رابطه دوطرفه و قدرتمندی میان فعالیت نقص اطلاعات و تعلیق با شاخص ارتقای ۱.۵۴۲ گردید که نشان‌دهنده یک چرخه مخرب بازکاری سیستماتیک است. علاوه بر این، کشف رابطه وابستگی با قطعیت ۱۰۰ درصدی (اطمینان ۱.۰۰) در انتقال از گام‌های تعلیق و نقص اطلاعات به گام اعتبارسنجی، به صورت ریاضیاتی اثبات نمود که مرحله اعتبارسنجی، گره مرکزی کنترل فرآیند و کانون اصلی بازگردانی پرونده‌ها در سازمان است.

۴. یافته‌های فاز سوم: تحلیل گرافی و کشف گلوگاه‌های شبکه‌ای در این فاز، لاگ رویداد به یک گراف جهت‌دار وزن‌دار شامل ۱۰ گره (فعالیت) و ۲۱ یال تبدیل شد که مجموع انتقال‌های مستقیم در آن به ۲۰۸,۰۸۶ مورد رسید. برای محاسبه صحیح کوتاه‌ترین مسیرهای ارتباطی و پیشگیری از خطای الگوریتم‌های گرافی، فاصله میان گره‌ها به صورت معکوس فراوانی انتقال فرمول‌بندی شد تا مسیرهای پرتدد به‌عنوان مسافت‌های کوتاه‌تر ارزیابی شوند. تحلیل مرکزیت شبکه و شاخص رتبه صفحه نشان داد که گام اعتبارسنجی با دارا بودن بیشترین حجم انتقال ورودی و خروجی (۷۷,۶۲۹ انتقال) و بالاترین امتیاز اهمیت گرهی (رتبه صفحه ۰.۱۵۳۳)، پرتراфик‌ترین و محوری‌ترین گره شبکه و مهم‌ترین نقطه اهرمی ساختاری فرآیند است.

۳۰,۰۲۰ پرونده (۹۵.۲۷٪) کم‌ریسک ارزیابی شدند. بررسی پرونده‌های پرریسک نشان داد که همگی در خوشه پرت قرار داشته و بین ۱۰ تا ۱۳ بار چرخه پینگ‌پونگ را تجربه کرده بودند.

موتور تجویزی بر اساس این تیپ‌شناسی رفتاری، بسته‌های مداخله مدیریتی را به تفکیک صادر نمود: ارجاع آنی و تخصصی ۱۶۵ پرونده خوشه پرت به کمیته رسیدگی به استثنائات جهت خروج از سیستم خودکار؛ پیاده‌سازی دروازه‌های پیش‌اعتبارسنجی برای خوشه‌های بازکاری محور جهت خنثی‌سازی بازگشت‌ها؛ استقرار تریگرهای هشدار و مکانیزم‌های ارجاع خودکار برای خوشه‌های انتظارمحور؛ و بازتوزیع پویای منابع انسانی با تفویض اختیار و موازی‌سازی وظایف در ایستگاه بحرانی اعتبارسنجی بر اساس ماتریس ارتباطی فرآیند-نقش. این نتایج نشان داد که مدل پیشنهادی با ارائه پیشنهادات نقطه‌زنی‌شده عملیاتی، بهینه‌سازی فرآیندها را به طور کامل تضمین می‌کند.

بخش هفتم: بحث، نتیجه‌گیری و مراجع

۱. **بحث و تفسیر یافته‌ها** پیاده‌سازی مدل چندلایه پیشنهادی بر روی کلان‌داده استاندارد درخواست وام نشان داد که ادغام الگوریتم‌های هوش مصنوعی با ابزارهای سنتی مدیریت فرآیند، فراتر از یک ابزار توصیفی ساده، به عنوان یک بستر تصمیم‌یار تجویزی عمل می‌کند. یافته‌های حاصل از الگوریتم خوشه‌بندی چگالی‌محور DBSCAN با تفکیک دقیق پرونده‌ها به

بالای پینگ‌پونگ (۱.۳ تا ۱.۹ بار) و حضور قطعی وضعیت‌های نقص و تعلیق که کانون اصلی اتلاف منابع سازمانی هستند.

- خوشه‌های استاندارد و روان (۲۱.۸٪ کل پرونده‌ها): شامل ۶,۸۷۷ پرونده با کمترین میزان بازکاری و انحراف که مسیر طلایی فرآیند را طی کرده‌اند.

- خوشه‌های انتظارمحور (۱۴.۵٪ کل پرونده‌ها): شامل ۴,۵۸۱ پرونده با نرخ پینگ‌پونگ صفر اما زمان تعلیق بسیار بالا که مشکل اصلی آن‌ها توقف جریان کار است.

- خوشه‌های زمان‌بر (۱۰.۸٪ کل پرونده‌ها): شامل ۳,۴۱۰ پرونده با زمان چرخه نامتعارف بالای ۶۰۰ ساعت با وجود نرخ بازکاری پایین.

همچنین، الگوریتم موفق به شناسایی دقیق ۱۶۵ پرونده کاملاً پرت و ناهنجار (خوشه ۱-) گردید که تنها ۰.۵۲ درصد کل پرونده‌ها را تشکیل می‌دادند اما میانگین پینگ‌پونگ فاجعه‌بار ۴.۸۹ بار و میانگین زمان رسیدگی نامتعارف ۱۳۰۹ ساعت را به ثبت رسانده و بازکاری سنگینی به منابع تحمیل کرده بودند.

۶. **یافته‌های فاز پنجم: خروجی‌های موتور تجویزی و توزیع ریسک** در لایه نهایی، موتور تجویزی با تلفیق یافته‌های لایه‌های پیشین، به محاسبه امتیاز ریسک برای پرونده‌ها پرداخت. توزیع ریسک محاسبه‌شده در سطح پرونده‌ها بسیار گزینشی و دقیق عمل کرد، به گونه‌ای که تنها ۳۷ پرونده (۰.۱۲٪) در سطح ریسک بالا و ۱,۴۵۲ پرونده (۴.۶۱٪) در سطح ریسک متوسط قرار گرفتند و

متن دستورالعمل‌های رسمی و رویه‌های پیاده‌سازی شده در سامانه‌ها، ریشه اصلی شکل‌گیری این رفتارهای انحرافی است.

۲. نتیجه‌گیری کلان پژوهش حاضر با طراحی، پیاده‌سازی و ارزیابی مدل یکپارچه هوش مصنوعی بر پایه رویکرد علم طراحی، گامی نو در مرزهای دانش مدیریت فرآیندهای کسب‌وکار برداشت. دستاورد بنیادی این تحقیق، نشان دادن امکان گذار موفقیت‌آمیز از فرآیندکاوی توصیفی گذشته‌نگر به مدیریت فرآیند تجویزی آینده‌نگر است. مدل پیشنهادی توانست با یکپارچه‌سازی الگوریتم‌های کاوش قواعد وابستگی، خوشه‌بندی چگالی‌محور، نظریه گراف و مدل‌های زبانی پردازش متن، یک خط لوله پردازشی منسجم ایجاد کند که خروجی هر لایه، ورودی لایه بعدی را تصفیه و غنی می‌سازد.

هم‌راستاسازی خروجی‌های این مدل هوشمند با چارچوب بین‌المللی APQC نشان داد که می‌توان از استانداردهای مدیریتی به عنوان پشتیبان ساختاری بهره گرفت، بدون آنکه پویایی تحلیل‌های هوش مصنوعی دچار خدشه شود. در نهایت، نتایج تجربی روی داده‌های واقعی اثبات کرد که پیاده‌سازی این مدل، بستر تصمیم‌گیری مدیران ارشد را از حالت حدس و گمان‌های شهودی به یک فرآیند کاملاً علمی، ممیزی‌پذیر و داده‌محور تبدیل می‌کند که نتیجه مستقیم آن، ارتقای چابکی عملیاتی، کاهش هزینه‌ها و تعالی سازمانی است.

۱۸ خوشه همگن رفتاری و یک خوشه پرت، عارضه‌های ساختاری را در سازمان آشکار ساخت. تمرکز ۴۶ درصد از پرونده‌ها در خوشه‌های بازکاری‌محور و کشف ۱۶۵ پرونده کاملاً پرت و ناهنجار با میانگین فاجعه‌بار ۴۰۸۹ بار بازگشت و ۱۳۰۹ ساعت زمان رسیدگی، نشان می‌دهد که سیستم‌های سنتی فاقد مکانیزم هوشمند برای شناسایی انحرافات در نطفه هستند. این پرونده‌های پرت با وجود سهم ناچیز عددی، بخش عظیمی از پهنای باند و منابع سازمان را می‌بلعند. این یافته با نظریات نوین فرآیندکاوی هم‌راستاست که معتقدند عارضه‌یابی فرآیندها باید به جای میانگین‌های کلان، بر روی ریزرفتارهای واریانت‌های مختلف متمرکز شود.

علاوه بر این، تحلیل شبکه‌های گرافی و شاخص‌های مرکزیت در این پژوهش توانست چرخه پینگ‌پونگی میان نقص اطلاعات و اعتبارسنجی را با آمار خیره‌کننده ۲۴,۹۸۶ بار تکرار در ۳۷۰۰۳ درصد کل پرونده‌ها آشکار کند. از منظر تحلیل فرآیند، وقوع چنین حجم عظیمی از پاس‌کاری‌های مداوم میان متقاضی و کارشناس اعتبارسنجی اثبات می‌کند که دستورالعمل‌های جاری و فرم‌های ثبت اطلاعات اولیه دچار ابهام ساختاری هستند. در روش‌های سنتی، مدیران معمولاً کندی این مرحله را به کم‌کاری منابع انسانی نسبت می‌دهند، در حالی که تحلیل گرافی این پژوهش به صورت ریاضیاتی ثابت کرد که گلوگاه اصلی، یک عارضه تعاملی و ساختاری است. الگوریتم Apriori نیز با اثبات همبستگی ۱۰۰ درصدی گام‌های نقص اطلاعات به سمت اعتبارسنجی مجدد، مهر تأییدی بر این فرضیه ساختاری زد. تلفیق این نتایج با تکنیک‌های پردازش متن نشان داد که ناسازگاری میان

۳. دستاوردهای مدیریتی و بسته‌های پیشنهادی بر

اساس یافته‌های عینی این پژوهش، سه بسته پیشنهادی و راهبردی برای مدیران و تصمیم‌سازان سازمان ارائه می‌گردد:

بسته اول: مدیریت پرونده‌های پرت و استثنائات عملیاتی با توجه به شناسایی پرونده‌های ناهنجار با زمان چرخه بالای ۱۳۰۰ ساعت، توصیه می‌شود سازمان از سیستم‌های هدایت خودکار سنتی برای این موارد فاصله بگیرد. موتور تجویزی مدل پیشنهادی قادر است امتیاز ریسک پرونده‌ها را در همان گام‌های اولیه محاسبه کند. پیشنهاد می‌شود پرونده‌هایی که نمره ریسک بالا دریافت می‌کنند، بلافاصله از خط لوله خودکار خارج شده و به یک «کمیته تخصصی رسیدگی به موارد خاص» ارجاع داده شوند تا از رسوب کار در ایستگاه‌های عادی جلوگیری شود.

بسته دوم: اصلاح ساختاری فرآیند ثبت اطلاعات و پیش‌اعتبارسنجی با توجه به عارضه چرخه پینگ‌پونگی میان نقص اطلاعات و اعتبارسنجی (درگیر کردن بیش از ۳۷ درصد پرونده‌ها)، پیشنهاد می‌شود سازمان پلتفرم ورودی اطلاعات مشتریان را بازطراحی کند. این کار باید از طریق پیاده‌سازی فیله‌های خودکنترل هوشمند و ایجاد سیستم‌های پیش‌اعتبارسنجی آنلاین صورت گیرد تا پرونده تا زمانی که از صحت و کمال مدارک آن اطمینان حاصل نشده، اصلاً وارد کارتابل کارشناسان اعتبارسنجی نشود. این اقدام به تنهایی می‌تواند زمان چرخه کل فرآیند را تا ۳۰ درصد کاهش دهد.

بسته سوم: بازتوزیع هوشمند بار کاری منابع انسانی با استناد به ماتریس ارتباطی فرآیند-نقش و مرکزیت بینابینی بالای ایستگاه اعتبارسنجی، توصیه می‌شود سازمان با آموزش‌های چندگانه به کارکنان سایر بخش‌های کم‌ترافیک، تفویض اختیارات لازم و ایجاد مسیرهای موازی برای بررسی مدارک ساده، بار کاری را از روی کارشناسان اعتبارسنجی بردارد تا پدیده فرسودگی شغلی و انباشت پرونده‌ها در این گلوگاه کلیدی مرتفع گردد.

۴. محدودیت‌های پژوهش با وجود دستاوردهای ارزشمند، این پژوهش با محدودیت‌هایی روبه‌رو بوده است که باید در تفسیر نتایج مدنظر قرار گیرند: نخست، پیچیدگی محاسباتی بالای پردازش متون غیرساختاریافته با مدل‌های زبانی پیشرفته و تحلیل هم‌زمان گراف روی کلان‌داده‌ها، نیازمند زیرساخت‌های سخت‌افزاری قدرتمندی است که ممکن است در برخی سازمان‌های متوسط در دسترس نباشد. دوم، این پژوهش مدل خود را بر روی مجموعه داده آفلاین و تاریخی (۲۰۱۷ BPIC) پیاده‌سازی و ارزیابی کرده است؛ در حالی که پیاده‌سازی زنده و برخط (Streaming) این پایپ‌لاین بر روی جریان داده‌های لحظه‌ای سازمان‌ها با چالش‌هایی نظیر نوسانات سرعت شبکه و تغییرات پویای پایگاه‌های داده همراه است. سوم، محدودیت‌های ذاتی زبان فارسی در ابزارهای پردازش زبان طبیعی (NLP) نسبت به زبان انگلیسی، استخراج روابط فرآیندی از متون فارسی را با

- چالش‌های بیشتری مواجه می‌سازد که نیازمند توسعه هستان‌شناسی‌های تخصصی بومی است.
- ۵. **پیشنهادهایی برای پژوهش‌های آتی جهت توسعه** مرزهای دانش در این حوزه، پیشنهادها پژوهشی زیر برای محققان آتی ارائه می‌گردد: ۱. توسعه و پیاده‌سازی پایپ‌لاین پیشنهادی در بستر جریان زنده داده‌ها (Real-time Stream Mining) جهت پایش آنلاین و ارائه پیشنهادات تجویزی در لحظه وقوع فعالیت‌ها. ۲. استفاده از مدل‌های زبانی بزرگ بومی‌سازی‌شده (LLMs) و تکنیک‌های یادگیری چندوظیفه‌ای جهت ارتقای دقت استخراج مدل‌های فرآیندی از مستندات پیچیده اداری به زبان فارسی. ۳. به‌کارگیری الگوریتم‌های یادگیری تقویت‌شده (Reinforcement Learning) در لایه موتور تجویزی به منظور شبیه‌سازی و پیش‌بینی دقیق‌تر اثرات تصمیمات مدیریتی بر کل زنجیره ارزش قبل از اعمال واقعی آن‌ها. ۴. توسعه چارچوب مدل جهت پوشش فرآیندهای بین‌سازمانی و زنجیره تأمین مشترک (Collaborative BPM) که در آن‌ها چندین سازمان ناهمگون با پایگاه‌های داده مجزا تعامل دارند.
- ۶. **مراجع**
الف) مراجع فارسی:
 - جعفرنژاد، احمد؛ و مهدوی، علی. (۱۴۰۲). مدیریت فرآیندهای کسب‌وکار در عصر تحول دیجیتال. انتشارات دانشگاه تهران، چاپ دوم.
 - رضایی، حسین؛ و کریمی، محمدرضا. (۱۴۰۳). فرآیندکاوی تجویزی: ارائه چارچوب تصمیم‌یار پویا با استفاده از الگوریتم‌های یادگیری ماشین.
- نشریه مدیریت فناوری اطلاعات، ۱۶(۲)، ۱۱۲-۱۳۵.
- غیائی فرد، حسن؛ کریمی قهرودی، محمدرضا؛ محترمی، امیر؛ و آذرلی، آرمان. (۱۴۰۵). ارائه مدل کاربست هوش مصنوعی در مدیریت فرآیندهای کسب‌وکار سازمان. رساله دکتری تخصصی، دانشکده برق و کامپیوتر، دانشگاه صنعتی مالک اشتر.
- ب) مراجع انگلیسی (English References):
 - Davenport, T. H., & Spanyi, A. (۲۰۲۳). *Introducing Process Intelligence: How to use AI to discover, improve, and manage business processes*. Harvard Business Review Press.
 - Dumas, M., La Rosa, M., Mendling, J., & Reijers, H. A. (۲۰۱۸). *Fundamentals of Business Process Management* (۲nd ed.). Springer.
 - Ester, M., Kriegel, H. P., Sander, J., & Xu, X. (۱۹۹۶). A density-based algorithm for discovering clusters in large spatial databases with noise. In *Proceedings of the Second International Conference on Knowledge Discovery and Data Mining (KDD)*, ۲۲۶-۲۳۱.
 - Han, J., Kamber, M., & Pei, J. (۲۰۱۱). *Data Mining: Concepts and Techniques* (۳rd ed.). Morgan Kaufmann.
 - van der Aalst, W. M. P. (۲۰۱۶). *Process Mining: Data Science in Action* (۲nd ed.). Springer.
 - van der Aalst, W. M. P., & Damiani, E. (۲۰۲۲). *Process mining in the large: Challenges and opportunities in enterprise-wide process analysis*. IEEE



-
- Transactions on Services Computing, ۱۵(۴), ۱۸۰۳-۱۸۱۵.
 - van Dongen, B. (۲۰۱۷). BPI Challenge ۲۰۱۷ Dataset. Technology. <https://doi.org/10.26434/chemrxiv-2017-05-01>
-