



AI-enabled fraud detection in health insurance: a multi-module framework based on anomaly analysis

Mojtaba Farrokh^{1*}, Syrus Sharifi², Nasrin Hozarmoghadam³

¹ Assistant Professor, Operations and Information Technology Management Department, Faculty of Management, Kharazmi University, Tehran, Iran.

² Assistant Professor, Banking, Insurance and Customs Department, Faculty of Management, Kharazmi University, Tehran, Iran.

³ Assistant Professor, Department, Faculty, Insurance Research center, Tehran, Iran.

Abstract

BACKGROUND AND OBJECTIVES: The dramatic increase in health insurance losses in recent years and the growth of various fraudulent behaviors by policyholders, physicians, and other healthcare providers have doubled the necessity of utilizing intelligent methods for controlling and detecting fraud. The limitations of traditional methods, the absence of labeled data, and the complexity of fraud patterns represent the primary challenges faced by insurance companies, particularly in Iran. Healthcare fraud manifests in multifaceted ways, requiring sophisticated detection mechanisms. The literature identifies practices such as incorrect coding, upcoding, submitting duplicate bills, and charging for unnecessary services. Fraud is often a collaborative effort or anomaly involving patients, doctors, and institutions. Because historical data lacks explicit fraud labels, supervised learning models are often impractical for deployments. Consequently, this study leverages unsupervised learning to identify cases that deviate from established behaviors. The theoretical foundation positions the isolation forest and autoencoder algorithms as methodologies for isolating outliers. This research was conducted with the objective of designing and evaluating an intelligent, scalable, and modular framework tailored for detecting health insurance fraud. By deeply analyzing medical prescription data and extracting diverse behavioral features, this research enables the automated and more accurate identification of various types of fraudulent activities. The absolute ultimate goal is to transition from reactive auditing to a proactive system, thereby enhancing operational efficiency and significantly reducing financial leakage for insurance providers.

METHODOLOGY: The current research focuses on developing a health insurance fraud detection framework based on the Isolation Forest (IF) algorithm, combining feature engineering, unsupervised machine learning, and the analysis of abnormal behaviors. Initially, in collaboration with insurance experts and physicians, various types of fraud and their representative features were identified. Raw data underwent cleaning and structuring within a two-stage data warehouse. Subsequently, 20 key features were extracted and utilized to train the models. These features captured dynamics such as patient-physician interaction frequency, inconsistencies of age and gender with the provided service or diagnosis, abnormal cost growth, and service rarity. The IF algorithm was deployed to cluster

prescriptions and determine fraud probabilities, and the results were visualized using Principal Component Analysis (PCA). Finally, an operational software system was implemented, featuring managerial dashboards based directly on the aforementioned IF algorithm.

FINDINGS: The empirical results demonstrate that anomaly detection algorithms effectively uncover abnormal behavioral patterns with high accuracy. Key indicators such as sudden spikes in physician costs, unusual frequency of services by a single provider, inconsistencies between service types and professional specialties, and abnormal patient spending patterns proved crucial in identifying suspicious prescriptions. Principal Component Analysis (PCA) visualizations revealed that anomalies are primarily distributed at the extreme edges of the data, confirming that unsupervised models possess a robust capability to distinguish these from normal records. Comparatively, the Isolation Forest (IF) algorithm significantly outperformed the Autoencoder (AE) model in both predictive accuracy and computational efficiency. The deployed software, built with a RESTful architecture, MariaDB, and a React-based interface, incorporates 20 specific risk indicators. This system provides managerial dashboards that significantly reduce the time required for experts to analyze high-risk claims. The interactive nature of these tools allows for the ongoing guidance of training processes, parameter tuning, and detailed actor behavior analysis. Furthermore, the study emphasizes that the quality of extracted features is a decisive factor in algorithmic performance; precise selection through expert collaboration leads to significant improvements in results.

CONCLUSION: This research designed and evaluated an intelligent, modular framework for health insurance fraud detection based on unsupervised learning. Addressing the inadequacy of traditional auditing in the face of massive data volumes and complex fraud patterns, the IF-based framework offers an independent, configurable structure that adapts to the dynamic nature of non-conforming behaviors. The system's modular architecture, comprehensive documentation, and seamless integration capabilities make it a practical and scalable solution for insurance companies. Implementing such intelligent systems significantly reduces fraud-related losses and enhances operational efficiency. By analyzing historical data and identifying deviations, these tools enable the proactive detection of suspicious cases, thereby improving decision-making and resource allocation. However, system performance requires ongoing development of features related to actors, transactions, and communication networks to capture network-based risks. Continuous model updates and real-world feedback loops are essential for maintaining adaptability. Despite these strengths, the study acknowledges limitations including data quality issues (noise and incompleteness), reliance on a single insurance company, and legal/privacy constraints that challenge real-world implementation. Future research should focus on deeper expert involvement in defining fraud types, the use of meta-heuristic algorithms for feature engineering, and the exploration of graph-based and deep learning models to identify complex collusive rings. Additionally, integrating Natural Language Processing (NLP) to analyze unstructured medical notes and expanding the framework to other domains, such as auto insurance and credit risk, represent promising paths for future development.

KEYWORDS: Anomaly-Detection; Artificial intelligence; Fraud-Detection; Health insurance; Unsupervised learning.



کشف تقلب بیمه درمان با بهره‌گیری از هوش مصنوعی: چارچوب چندماژوله مبتنی بر تحلیل ناهنجاری

مجتبی فرخ^{۱*}، سپروس شریفی^۲، نسرين حصارمقدم^۳

^۱ استادیار، گروه مدیریت فن آوری اطلاعات، دانشکده مدیریت، دانشگاه خوارزمی، تهران، ایران.

^۲ استادیار، گروه بانک، بیمه و گمرک، دانشکده مدیریت، دانشگاه خوارزمی، تهران، ایران.

^۳ استادیار، گروه آموزشی، پژوهشکده بیمه، تهران، ایران.

چکیده

پیشینه و اهداف: افزایش چشمگیر خسارت‌های بیمه درمان در سال‌های اخیر و رشد انواع رفتارهای متقلبانه از سوی بیمه‌شدگان، پزشکان و سایر ارائه‌دهندگان خدمات درمانی، ضرورت بهره‌گیری از روش‌های هوشمند برای کنترل و شناسایی تقلب را دوجندان کرده است. محدودیت روش‌های سنتی، نبود داده‌های برجسته‌گذاری شده، و پیچیدگی الگوهای تقلب، چالش‌های اصلی شرکت‌های بیمه به‌ویژه در ایران به‌شمار می‌روند. پژوهش حاضر با هدف طراحی چارچوبی هوشمند، ماژولار و قابل‌گسترش برای کشف تقلب در بیمه درمان انجام شده است تا با تحلیل داده‌های نسخه‌های پزشکی و استخراج ویژگی‌های مختلف، امکان شناسایی دقیق‌تر انواع تقلب فراهم شود.

روش‌شناسی: این پژوهش چارچوبی برای کشف تقلب در بیمه درمان بر پایه الگوریتم جنگل ایزوله، همراه با مهندسی ویژگی، یادگیری ماشین بدون نظارت و تحلیل رفتارهای غیرعادی، توسعه داد. در مرحله نخست، انواع رایج تقلب و ویژگی‌های نمایانگر آن‌ها با همکاری کارشناسان بیمه و پزشکان شناسایی شد. در مرحله دوم، داده‌های خام نسخه‌ها و مطالبات در یک انبار داده دوبرحله‌ای پاک‌سازی، یکپارچه و ساختاردهی شدند. در مرحله سوم، ۲۰ ویژگی کلیدی شامل فراوانی تعامل بیمار و پزشک، ناسازگاری سن یا جنسیت با خدمت یا تشخیص ثبت‌شده، رشد غیرعادی هزینه‌های درمانی، نادر بودن برخی خدمات، تکرار نامنظم رویه‌های خاص و عدم انطباق میان تخصص ارائه‌دهنده و خدمت ارائه‌شده استخراج شد. سپس الگوریتم جنگل ایزوله برای شناسایی نسخه‌های پرت و برآورد احتمال تقلب آن‌ها به کار رفت. برای افزایش تفسیرپذیری نتایج، از تحلیل مؤلفه‌های اصلی جهت مصورسازی استفاده شد. در نهایت، یک سامانه نرم‌افزاری عملیاتی همراه با داشبوردهای مدیریتی برای پایش، بررسی و تحلیل تقلب از منظر جغرافیایی، جمعیتی و زمانی پیاده‌سازی شد.

یافته‌ها: نتایج نشان داد الگوریتم‌های شناسایی ناهنجاری با اثربخشی بالا قادر به شناسایی الگوهای رفتاری غیرعادی در داده‌های بیمه درمان هستند. افزایش ناگهانی هزینه‌های مرتبط با پزشک، تکرار غیرمعمول خدمات توسط یک ارائه‌دهنده، ناسازگاری میان نوع خدمت و تخصص پزشکی، و الگوهای غیرعادی هزینه در بیماران، از مهم‌ترین شاخص‌های تمایز نسخه‌های مشکوک از رکوردهای عادی بودند. همچنین، تحلیل‌های مبتنی بر مؤلفه‌های اصلی نشان داد که مشاهدات ناهنجار غالباً در حاشیه ساختار داده قرار می‌گیرند؛ موضوعی که مناسب بودن روش‌های بدون نظارت را برای تفکیک موارد مشکوک از مطالبات عادی تأیید می‌کند. افزون‌براین، الگوریتم جنگل ایزوله از نظر دقت شناسایی و کارایی محاسباتی، عملکردی بهتر از مدل خودرمزگذار داشت. نرم‌افزار توسعه‌یافته نیز امکان تحلیل چندبعدی تقلب در طول زمان، میان مناطق مختلف و در گروه‌های جمعیتی گوناگون را فراهم ساخت.

نتیجه‌گیری: چارچوب پیشنهادی نشان داد که تلفیق مهندسی ویژگی، شناسایی ناهنجاری و معماری نرم‌افزاری پیشرفته می‌تواند راهکاری عملی و مؤثر برای کشف تقلب در بیمه درمان فراهم آورد. دقت مناسب مدل و توانایی آن در عملکرد بدون داده‌های برچسب‌گذاری شده، این چارچوب را در شرایطی که سوابق تأییدشده تقلب محدود یا در دسترس نیستند، ارزشمند می‌سازد. همچنین، ساختار ماژولار آن امکان به‌روزرسانی مستمر، توسعه ویژگی‌ها و انطباق با الگوهای متغیر تقلب را فراهم می‌کند. این سامانه می‌تواند به بهبود مدیریت ریسک، پایش کارآمدتر مطالبات و کاهش خسارت‌های مالی در بیمه درمان کمک کند. پیشنهاد می‌شود مطالعات آینده با بهره‌گیری از منابع داده متنوع‌تر، مدل‌های عمیق‌تر و تحلیل شبکه‌ای، عملکرد و استحکام سامانه‌های کشف تقلب را بیش از پیش ارتقا دهند.

کلمات کلیدی: بیمه درمان؛ کشف ناهنجاری (آنومالی)؛ کشف تقلب؛ هوش مصنوعی؛ یادگیری بدون نظارت.

مقدمه

با توسعه روزافزون فناوری‌های اطلاعات و ارتباطات، صنعت مالی و بیمه به‌طور فزاینده‌ای به سمت بهره‌گیری از ابزارهای نوین حرکت کرده است. یکی از فناوری‌های کلیدی که نقش چشمگیری در تحول این حوزه ایفا کرده، هوش مصنوعی است. کاربرد هوش مصنوعی در صنعت بیمه، علاوه بر ارتقای دقت در پردازش داده‌ها، تحلیل صورت‌حساب‌های مالی و ارزیابی ریسک، امکان شناسایی و پیشگیری از تقلب‌های بیمه‌ای را نیز فراهم ساخته است (Phua et al., 2010; Viaene et al., 2002). تقلب در بیمه‌های درمانی به‌عنوان یکی از مهم‌ترین چالش‌های نظام‌های بیمه و درمان، سالانه موجب هدررفت منابع مالی هنگفت و کاهش اعتماد میان بیمه‌گذاران و بیمه‌گران می‌شود (Kou et al., 2004). در سال‌های اخیر، خسارت پرداختی در بیمه درمان به شدت افزایش یافته، به طوریکه در سال ۱۴۰۲ به حدود ۶۹۵۸۲ میلیارد تومان رسیده است، که سهم تقلب بین ۶۹۵۰ تا ۱۰۴۲۰ میلیارد تومان برآورد شده است. تعداد پرونده‌های خسارتی نیز از حدود ۱۷ میلیون در سال ۱۳۹۳ به بیش از ۸۰ میلیون در سال ۱۴۰۲ رسیده که رشد ۳۷۵ درصدی را نشان می‌دهد (سالنامه آماری صنعت بیمه، ۱۴۰۲). افزایش مستمر تعداد خسارت‌های بیمه درمان، به ویژه در سال‌های اخیر، ضرورت مطالعات دقیق، نظارت مؤثر و مدیریت بهینه را برجسته کرده است. ضرب بالای خسارت و حجم قابل توجه تقلب در ادعاهای بیمه‌ای، اهمیت اتخاذ راهکارهای مؤثر برای کنترل و پیشگیری از تقلب را دوچندان می‌کند. تقلبات بیمه‌ای نه تنها به هدررفت منابع مالی شرکت‌ها منجر می‌شوند، بلکه بر کیفیت خدمات‌رسانی نیز تأثیر منفی دارد (رامندی و همکاران، ۱۳۹۹). علاوه بر این، افزایش پیچیدگی خسارت‌های بیمه‌ای ناشی از تغییرات جمعیتی و شیوع بیماری‌های مزمن، فشار مضاعفی بر شرکت‌های بیمه وارد کرده است. برای مدیریت این چالش‌ها، استفاده از فناوری‌های پیشرفته همچون هوش مصنوعی، یادگیری ماشین و تجزیه و تحلیل داده‌ها ضروری شده است (Lu et al., 2023). این فناوری‌ها به تسریع پردازش خسارت‌ها، شناسایی فعالیت‌های تقلبی و تخصیص بهینه منابع کمک می‌کنند و اثرات مثبتی بر عملکرد بیمه‌گران دارند (Bagekari, 2025). بنابراین، شناسایی الگوهای رفتاری غیرعادی و مقابله با این پدیده به ضرورتی اجتناب‌ناپذیر بدل شده است.

در این میان، هوش مصنوعی با استفاده از الگوریتم‌های پیشرفته تحلیل داده و یادگیری ماشین، توانایی کشف ناهنجاری‌ها و الگوهای پنهان در داده‌های مرتبط با خسارت بیمه‌ای را دارد (معصومی جناقرد و همکاران، ۱۴۰۴). تحلیل داده‌های بیمه درمان با کمک روش‌های یادگیری ماشین، نقش کلیدی در کشف موارد تقلب و جلوگیری از هدررفت منابع بیمه دارد (du Preez et al., 2024; Baesens et al., 2021). با این حال، پیاده‌سازی الگوریتم‌ها در بیمه درمان با چالش‌هایی مانند داده‌های گمشده، بدون برچسب، حجیم و ناهمگون روبه‌روست که نیازمند پاک‌سازی، انتخاب ویژگی و استفاده از روش‌های مناسب است. در این پژوهش، داده‌های جمع‌آوری شده عمدتاً ساختاریافته بوده و شامل اطلاعات بیمار، پزشک، خدمات درمانی، مبلغ نسخه‌های پزشکی و مبلغ خسارت پرداختی هستند. تمرکز بر این داده‌ها امکان تحلیل داده‌محور و به‌کارگیری روش‌های آماری و یادگیری ماشین برای کشف تقلب را فراهم می‌کند، هرچند داده‌های نسخه‌های پزشکی به تنهایی دقت کافی در کشف تقلب ندارند (رحیم خانی و منطقی‌پور، ۱۴۰۰).

هدف اصلی این مطالعه، توسعه یک چارچوب ماژولار و انعطاف‌پذیر مبتنی بر الگوریتم‌های هوش مصنوعی برای کشف تقلب در بیمه درمان است که با تحلیل موارد تقلب، مدیران و سیاست‌گذاران را در اتخاذ تصمیمات مؤثر برای کاهش هزینه‌های مرتبط یاری می‌کند. هدف دوم، طراحی یک نرم‌افزار کاربردی در صنعت بیمه ایران است که بتواند به صورت هوشمند و بدون دخالت مستقیم انسان، فرآیند کشف تقلب را انجام دهد و بدین ترتیب کارایی و دقت عملیات بیمه‌ای را ارتقا دهد. نوآوری این پژوهش در استخراج ویژگی‌های کشف انواع تقلب با همکاری کارشناسان صنعت بیمه و استفاده از الگوریتم‌های بدون ناظر مناسب است که هدف آن ارتقای دقت شناسایی تقلب در بیمه درمان می‌باشد. با توجه به ماهیت ناهمگن و بدون برچسب داده‌ها، از الگوریتم‌های شناسایی ناهنجاری (آنومالی) و شاخص‌های اعتبارسنجی مناسب برای ارزیابی عملکرد این الگوریتم‌ها استفاده می‌شود.

مبانی نظری پژوهش

انجمن ملی پیشگیری از تقلب در مراقبت‌های بهداشتی^۱ به‌طور جداگانه، تقلب و سوءاستفاده در سیستم‌های بهداشت و درمان را به شرح زیر تعریف می‌کند: «تقلب در مراقبت‌های بهداشتی، یک فریب یا گمراهی عمدی است که فرد یا نهادی با علم به اینکه این گمراهی می‌تواند به نفع غیرمجاز فرد یا موجودیت دیگری منجر شود، انجام می‌دهد». شرکت‌های بیمه خصوصی و عمومی به‌عنوان تأمین‌کنندگان اصلی تأمین مالی خدمات درمانی، انگیزه قوی برای جلوگیری از خسارات ناشی از تقلب و سوءاستفاده دارند. با این حال، اغلب هیچ رویکرد سیستماتیکی برای رسیدگی به موضوع کشف تقلب در این شرکت‌ها وجود ندارد (رحیم‌خانی و همکاران، ۱۴۰۱). در این میان انواع تقلب‌ها را از سوی بازیگران مختلف شاهد هستیم. کدگذاری نادرست خدمات در فرم‌های ادعای بیمه می‌تواند منجر به پرداخت‌های بیش از حد و خسارات مالی شود. ارتقا کد، شامل ثبت بیمار در گروه تشخیص با کدهای بالاتر برای دریافت بازپرداخت بیشتر، نمونه‌ای از تقلب است (Luo & Gallagher, 2010). صدور صورتحساب تکراری یا استفاده از چندین کد برای یک خدمت نیز از نشانه‌های رایج کلاهبرداری به‌شمار می‌رود (Bauder et al., 2017; Suroor & Misra, 2024). ثبت ادعاهای نامرتب، غیرواقعی یا بیش از حد، صدور صورتحساب برای خدمات پوشش داده نشده و رویه‌های غیرضروری پزشکی نیز معمولاً در تحقیقات تقلب مورد بحث قرار می‌گیرد (Wynia et al., 2000; Hamid et al., 2024).

از سوی دیگر، رفتار ارائه‌دهندگان می‌تواند نشانه‌های کلاهبرداری را آشکار کند؛ مانند نسبت‌های غیرمعمول ملاقات بیماران و ناسازگاری بین تشخیص و برنامه‌های درمانی، یا برنامه‌های درمانی با تعداد بیماران فراتر از توان مدیریت ارائه‌دهنده (Johnson & Nagarur, 2016). ترونتون و همکاران (Thornton et al., 2015) انواع تقلب بیمه درمان را شناسایی کرده‌اند، از جمله: پرداخت رشوه توسط داروخانه‌ها یا پزشکان برای نسخه‌های خاص (Szalados, 2021; du Preez et al., 2025)، و خرید پزشک برای دریافت نسخه‌های مورد نظر (Suroor & Misra, 2024). کدگذاری نادرست و ارتقای کد شامل صدور صورتحساب برای خدمات با مبالغ بالاتر است (Luo & Gallagher, 2010). همچنین جدا کردن خدمات، ارسال ادعاهای تکراری یا برای خدمات ارائه‌نشده، و ارائه مراقبت‌های غیرضروری یا بیش‌از‌حد، نمونه‌های متداول هستند (Johnson & Nagarur, 2016; Suroor & Misra, 2024). دیگر اشکال تقلب شامل استفاده از تشخیص نادرست، صورتحساب توسط افراد غیرمجاز، دروغ درباره صلاحیت، ادعای خسارت غیرواقعی، تقلب در مراقبت مدیریت شده و چشم‌پوشی از پرداخت‌های مشارکتی است که می‌تواند موجب دریافت غیرقانونی پول یا ارائه مراقبت نادرست شود (Szalados, 2021; Sparrow, 2009; Thornton et al., 2015). تصمیمات پزشکان نقش عمده‌ای در هزینه‌های بهداشت و درمان دارند و حدود ۸۰ درصد از این هزینه‌ها به تجویز خدمات آن‌ها وابسته است. برخی پزشکان ممکن است برای افزایش درآمد خود، خدمات غیرضروری ارائه کنند یا رفتارهای کلاهبردانه مانند تغییر نسخه‌ها، درخواست بازپرداخت برای درمان‌های انجام‌نشده و ایجاد «بیماران سایه‌ای»^۲ داشته باشند (Price & Norris, 2009; Wynia et al., 2000). بیماران نیز ممکن است با ارائه اطلاعات نادرست به شرکت‌های بیمه، رفتار کلاهبردانه نشان دهند. نمونه‌های رایج شامل ارسال ادعا برای افراد غیرمشمول، ثبت ادعا برای خدمات دریافت‌نشده و استفاده از بیمه دیگران است (Suroor & Misra, 2024).

در حال حاضر، بیشتر روش‌های کشف تقلب بیمه درمان به جهت نبود برچسب نسخه‌ها بر مدل‌های غیرنظارت‌شده تمرکز دارند. هدف اصلی یادگیری غیرنظارت‌شده، کشف الگوها و روابط پنهان در داده‌ها بدون نیاز به برچسب‌های از پیش تعیین‌شده است. این امر از طریق روش‌هایی مانند خوشه‌بندی، که داده‌های مشابه را بر اساس شباهت یا فاصله گروه‌بندی می‌کند، یا کاهش ابعاد، که ساختار اصلی داده را در قالبی ساده‌تر حفظ می‌کند، قابل دستیابی است (Muscoloni et al., 2017). الگوریتم‌های جنگل ایزوله^۳ (IF) و خودمرزگذار^۴ (AE) از جمله مهمترین این روش‌ها هستند. این روش‌ها، به الگوریتم‌ها شناسایی ناهنجاری‌ها (آنومالی‌ها/تقلب)، یعنی مواردی که از الگوی کلی انحراف دارند، معروف هستند.

مروری بر پیشینه پژوهش

در سال‌های اخیر، تشخیص تقلب در خدمات درمان و بیمه با تکیه بر روش‌های داده‌محور و الگوریتم‌های یادگیری ماشین توجه بسیاری از پژوهشگران را به خود جلب کرده است. نخستین تلاش‌ها در این زمینه بر استفاده از مدل‌های آماری و تحلیل‌های چندمتغیره متمرکز بود. برای نمونه، **مجر و ریدینگر** (Major & Riedinger, 1992) با طراحی سیستم کشف تقلب الکترونیک، مجموعه‌ای از ویژگی‌های مالی، زمانی و منطقی را برای شناسایی ارائه‌دهندگان مشکوک به کار بردند، در حالی که **ویلیامز** (Williams, 1999) از خوشه‌بندی k-means برای شناسایی نقاط داغ در داده‌های بیمه استفاده کرد. به دنبال آن، **یامانی‌شی و همکاران** (Yamanishi et al., 2000) با ارائه‌ی روش SmartSifter بر پایه مدل‌های آماری تطبیقی، گامی مهم در تحلیل ناهنجاری‌های آنلاین در داده‌های درمان برداشتند. این مطالعات، مسیر استفاده از رویکردهای بدون نظارت را برای کشف الگوهای غیرعادی در داده‌های فاقد برچسب هموار کردند. پس از آن، پژوهش‌ها به سمت به‌کارگیری الگوریتم‌های خوشه‌بندی و قواعد انجمنی برای تحلیل

¹ National Health Care Anti-Fraud Association (NHCAA)

² Ghost Patients

³ Isolation Forest (IF)

⁴ Autoencoder (AE)

الگوهای رفتاری ارائه‌دهندگان خدمات درمانی سوق یافت. **لین و همکاران** (Lin et al., 2020) با استفاده از خوشه‌بندی بدون نظارت بر ویژگی‌های عملکرد پزشکان عمومی، گروه‌های بحرانی را شناسایی کردند و **شان و همکاران** (Shan et al., 2008) با استخراج قوانین همبستگی میان تراکنش‌های تخصصی، میزان تخلف هر ارائه‌دهنده را بر اساس نقض قوانین محاسبه کردند. همچنین، **اکینا و همکاران** (Ekina et al., 2013) با به‌کارگیری خوشه‌بندی بیزی توانستند روابط پنهان میان ارائه‌دهندگان همکار در تقلب را شناسایی کنند. این پژوهش‌ها نشان دادند که روش‌های خوشه‌بندی، به‌ویژه در شرایط نبود داده‌های برچسب‌دار، ابزار مناسبی برای شناسایی تقلب‌های ساختاری در نظام درمان محسوب می‌شوند.

از میانه‌ی دهه ۲۰۱۰، مطالعات جدیدتر با ترکیب روش‌های آماری و الگوریتم‌های یادگیری ماشین، چارچوب‌های هیبریدی برای شناسایی تقلب ارائه کردند. **کوزه و همکاران** (Kose et al., 2015) با استفاده از روش AHP برای وزن‌دهی ویژگی‌ها و خوشه‌بندی داده‌ها، یک مدل یادگیری ماشین خودتعالیمی پیشنهاد کردند. در ادامه، **جانسون و ناگرو** (Johnson & Nagarur, 2016) با اجرای یک چارچوب چندمرحله‌ای شامل پروفایل‌سازی و تحلیل ریسک، موفق به دستیابی به دقتی حدود ۸۶ درصد در شناسایی ادعاهای مشکوک شدند. **ورما و همکاران** (Verma et al., 2017) نیز با ترکیب معیارهای مبتنی بر دوره درمان و نوع بیماری، مدلی توسعه دادند که با خوشه‌بندی داده‌ها و مقایسه آماری، توانایی بالایی در شناسایی رفتارهای غیرعادی نشان داد. در ایران نیز **خانی‌زاده و همکاران** (۱۴۰۱) با استفاده از الگوریتم جنگل ایزوله در داده‌های بیمه خودرو، به کاهش معنادار خطای مثبت کاذب دست یافتند که این یافته می‌تواند در بیمه درمان نیز قابل تعمیم باشد.

در سال‌های اخیر، تمرکز پژوهش‌ها به سمت الگوریتم‌های یادگیری عمیق و مدل‌های ترکیبی با قابلیت تبیین حرکت کرده است. **حدادسلیمانی و همکاران** (Haddad Soleymani et al., 2018) با طراحی مدل سه‌مرحله‌ای برای ارزیابی ریسک نسخه‌های پزشکی، توانستند الگوهای مشکوک را با دقت بالاتری شناسایی کنند. **دیب و همکاران** (Dhieb et al., 2020) با ترکیب بلاک چین و الگوریتم XGBoost، چارچوبی امن و دقیق برای شناسایی تقلب در بیمه خودرو پیشنهاد کردند. **گومز و همکاران** (Gomes et al., 2021) و **سوسرمن و همکاران** (Suesserman et al., 2023) با استفاده از خودرمزگذارهای عمیق و تابع زبان وزنی، توانستند مدل‌های مؤثرتری نسبت به روش‌های مبتنی بر چگالی ارائه دهند. همچنین، **ساماریا و همکاران** (Samariya et al., 2023) نشان دادند که الگوریتم IF عملکرد بهتری نسبت به LOF دارد، در حالی که **تبسم و همکاران** (Tabassum et al., 2024) با ترکیب این دو روش، میزان مثبت کاذب را به شکل محسوسی کاهش دادند. تازه‌ترین مطالعات مانند **دیمولمستر و همکاران** (De Meulemeester et al., 2025) و **مطلوب و همکاران** (Matloob et al., 2025) نیز با معرفی چارچوب‌های چندسطحی مبتنی بر ترنسفورمرها و مدل‌های قابل تبیین، دقت بالاتر و درک‌پذیری بیشتری از الگوهای تقلب در داده‌های بیمه درمان فراهم کرده‌اند. در واقع، مسیر پژوهش در زمینه‌ی کشف تقلب بیمه درمان از روش‌های آماری و خوشه‌بندی کلاسیک به سمت الگوریتم‌های عمیق و ترکیبی پیشرفته‌تر حرکت کرده است. چالش‌هایی مانند نبود داده‌های برچسب‌دار و نرخ بالای مثبت کاذب، انگیزه‌ای برای توسعه‌ی روش‌های بدون نظارت بوده است. مرور مطالعات پیشین نشان می‌دهد که علی‌رغم تلاش‌های صورت‌گرفته در زمینه شناسایی تقلب در بیمه درمان، شکاف‌های قابل‌توجهی همچنان وجود دارد. بیشتر پژوهش‌ها تنها بر یک بازیگر یا فرآیند خاص تمرکز کرده‌اند، در حالی که تقلب در عمل، نتیجه تعامل و همکاری میان بازیگران مختلف مانند بیمه‌شدگان، پزشکان و داروسازان است. همچنین، هیچ مدل جامعی که تمامی انواع ادعاها، رفتارها و فرآیندها را دربرگیرد، ارائه نشده و ارتباط میان بازیگران و خدمات در تحلیل‌ها غالباً نادیده گرفته شده است. افزون بر این، تغییرات پویا در رفتار متقلبان مستلزم الگوریتم‌هایی است که قابلیت به‌روزرسانی و انطباق مداوم داشته باشند. در حوزه کشف تقلب مبتنی بر نسخه‌های درمانی نیز چالش‌هایی چون حجم بالای داده‌ها، وجود ناهنجاری‌ها، ابعاد زیاد ویژگی‌ها و نبود داده‌های برچسب‌گذاری‌شده، دقت الگوریتم‌ها را تحت‌تأثیر قرار می‌دهد. نبود معیارهای ارزیابی روشن برای سنجش عملکرد مدل‌ها نیز ارزیابی دقت آن‌ها را دشوار می‌سازد. این محدودیت‌ها ضرورت توسعه روش‌های نوین و کارآمدتر را برجسته می‌سازند. در پاسخ به این شکاف‌ها، پژوهش حاضر بر استفاده از الگوریتم‌های یادگیری ماشین بدون نظارت IF تمرکز دارد تا با استخراج الگوهای پنهان و ویژگی‌های جدید از داده‌های نسخه‌های درمانی، امکان شناسایی دقیق‌تر رویه‌های غیرموجه و رفتارهای تقلب‌آمیز فراهم شود.

روش‌شناسی پژوهش

چارچوب کشف تقلب پیشنهادی مبتنی بر سه ماژول زیر است:

ماژول ۱ (بخش دانشی)

در ماژول نخست، دانش کارشناسان بیمه و پزشکی برای تدوین فریم‌ورک تقلب و تعیین انواع تخلفات قابل شناسایی به‌کار گرفته شد. در ابتدا در راستای تقویت عملکرد مدل‌های یادگیری ماشین در شناسایی تقلب‌های درمانی، با همکاری کارشناسان بیمه و پزشکان برای هر نوع تقلب یک فریم‌ورک تعریف که شامل بازیگران و مجموعه‌ای از ویژگی‌های ترکیبی و مشتق‌شده از داده‌های اولیه برای کشف آن‌ها است، طراحی و سپس ویژگی‌های مورد نیاز جهت کشف این تقلب‌ها استخراج شده‌اند که رفتارهای غیرعادی ارائه‌دهندگان (پزشک، داروخانه، ...) و بیماران را به‌صورت کمی مدل‌سازی می‌کنند. از طریق جلسات تخصصی، بیست ویژگی کلیدی مرتبط با رفتارهای غیرعادی بازیگران، خدمات و کالاها انتخاب شد. این ویژگی‌ها

مبنای آموزش الگوریتم‌ها قرار گرفتند و نقش اساسی در بهبود عملکرد مدل ایفا کردند. ویژگی‌های تعریف شده در چارچوب پیشنهادی که مربوط به ناسازگاری‌های مختلف بین بازیگر-خدمات/کالا است، در جدول ۱ آورده شده است.

جدول ۱. نمونه‌ای از ویژگی‌های استخراج شده برای کشف تقلب بیمه درمان

Table 1. Examples of extracted features for health insurance fraud detection

ویژگی‌ها	ویژگی استخراج شده	شرح ویژگی
۱	نسبت فراوانی مراجعات یک بیمار به یک پزشک خاص	نسبت تعداد کل ارائه‌دهندگان (نسخه‌ها) به تعداد کل ارائه‌دهندگان منحصر به فرد برای بیمار i در ماه t در هر نسخه
۲	نسبت فراوانی ارائه خدمات یک پزشک به یک بیمار خاص.	نسبت تعداد کل بیماران (نسخه‌ها) به تعداد کل بیماران منحصر به فرد در ماه t برای ارائه‌دهنده j در هر نسخه
۳	نسبت متوسط مبلغ دریافتی پزشک به متوسط مبلغ دریافتی همان پزشک در ماه قبل	درصد کاهش هزینه متوسط نسخه‌های ارائه‌شده توسط ارائه‌دهنده j در ماه t نسبت به متوسط هزینه نسخه‌های همان ارائه‌دهنده در ماه t-1 در هر نسخه
۴	نسبت مبلغ پرداخت شده بیمار به پزشک برای یک خدمت خاص به میانگین مبلغ دریافتی سایر پزشکان در همان خدمت	درصد افزایش هزینه خدمت S در ماه t نسبت به میانگین هزینه همان خدمت در ماه t-1 در هر نسخه
۵	نسبت مبلغ دریافتی پزشک برای یک خدمت خاص به میانگین مبلغ دریافتی همان پزشک در همان خدمت در ماه گذشته	درصد افزایش هزینه متوسط خدمت S در ماه t برای ارائه‌دهنده j در نسخه k نسبت به میانگین هزینه همان خدمت در ماه t-1
۶	نسبت مبلغ دریافتی پزشک برای یک تخصص خاص به میانگین مبلغ دریافتی سایر پزشکان در همان تخصص	درصد افزایش هزینه متوسط تخصص S برای ارائه‌دهنده j در ماه t در نسخه k نسبت به میانگین مبلغ تخصص S در ماه t-1
۸	نسبت مبلغ دریافتی پزشک برای یک تخصص خاص به میانگین مبلغ دریافتی سایر پزشکان در همان تخصص	درصد افزایش هزینه تخصص S برای ارائه‌دهنده j در ماه t در نسخه k نسبت به میانگین مبلغ تخصص S در ماه t-1
۹	نسبت مبلغ نسخه بیمار برای یک خدمت خاص به میانگین مبلغ نسخه همان خدمت برای سایر بیماران	درصد افزایش هزینه متوسط خدمت S در ماه t برای بیمار i در نسخه k نسبت به میانگین مبلغ خدمت S در ماه t-1
...	...	...
۲۰	نسبت فراوانی ارائه یک خدمت خاص توسط یک پزشک به تعداد کل نسخه‌های صادر شده توسط همان پزشک. اگر خدمتی که توسط پزشک در یک نسخه انجام شده نادر باشد و همزمان تعداد نسخه‌های پزشک زیاد، این نسبت کوچک بوده و نشان دهنده مشکوک بودن پزشک است.	نسبت تعداد ارائه خدمت S به تعداد نسخه‌ها برای ارائه‌دهنده j در نسخه k

در این ماژول، بعد از عملیات پاک‌سازی، یکپارچه‌سازی و آماده‌سازی داده‌ها، مقدار ویژگی‌های مختلف در قالب یک انبار داده ذخیره و به عنوان ورودی الگوریتم در ماژول دوم مورد استفاده قرار می‌گیرد.

ماژول ۲ (موتور کشف تقلب)

در ماژول دوم، ابتدا الگوریتم IF برای طبقه‌بندی ادعاهای بدون برچسب به تقلب و نرمال اجرا می‌شود. در تشخیص ناهنجاری‌ها (تقلب) با استفاده از الگوریتم IF، یک مجموعه از درخت‌های ایزوله^۱ را به صورت بازگشتی و با تقسیم مجموعه داده‌های آموزش در طول زمان ساخته و توسعه می‌دهد، تا زمانی که هر نقطه داده در برگ خود ایزوله شود یا حد ارتفاع درخت به آن برسد. این الگوریتم نمره ناهنجاری^۲ برای هر یک از نمونه‌ها محاسبه و سپس براساس یک حد آستانه نمونه‌ها را تفکیک می‌کند. سه هاپیرپارامتر ورودی برای الگوریتم IF وجود دارد. اندازه زیر نمونه ϕ که کنترل‌کننده اندازه داده‌های آموزشی است و معمولاً به صورت پیش‌فرض برابر با ۲۵۶ در نظر گرفته می‌شود. هاپیرپارامتر دوم، تعداد درخت‌های ساخته شده است که کنترل‌کننده اندازه مجموعه است و معمولاً مقدار آن ۱۰۰ تعیین می‌شود، چون طول مسیرها معمولاً زودتر از این مقدار همگرایی خوبی دارند.

¹ Ensemble of Isolation Trees

² anomaly score

هایپرپارامتر سوم، نرخ آلودگی^۱ است که نسبت تقریبی نمونه‌های ناهنجار در داده را مشخص می‌کند و آستانه‌ی جداسازی را تعیین می‌نماید. این نرخ براساس نظر کارشناسی در مازول اول تعیین می‌شود. الگوریتم بر اساس نرخ آلودگی تصمیم می‌گیرد چه نمونه‌هایی با نمره‌ی ناهنجاری بالاتر به‌عنوان ناهنجار/تلقب برچسب‌گذاری شوند. در واقع نرخ آلودگی به‌عنوان یک حدآستانه در این الگوریتم عمل می‌کند (Puggini & McLoone, 2018). الگوریتم ۱ و ۲ در شکل ۱ و ۲ جزئیات مربوط به مرحله آموزش را نشان می‌دهند.

Algorithm 1: IF (D, t, φ)	
Input: $D = (x_1, x_2, \dots, x_n)$ - dataset, t - number of trees, φ - sub-sampling	
Output: a set of t isolation trees	
Initialize Forest	
1:	set height limit $l = \text{ceiling}(\log_2 \varphi)$
2:	for $i = 1$ to t do
3:	$D' \leftarrow \text{sample}(D, \varphi)$
4:	Forest $\leftarrow \text{Forest} \cup \text{iTree}(D', 0, l)$
5:	end for
6:	return Forest

شکل ۱. شبه‌کد الگوریتم IF
Fig. 1. Pseudocode of the IF Algorithm

Algorithm ۲ : iTree (D, h, l)	
Input: $D = (x_1, x_2, \dots, x_n)$ - dataset, e - current tree height, l -height limit	
Output: an iTree t	
Initialize: $t = \emptyset$ (an empty tree)	
1:	if $e \geq l$ or $ X \leq 1$ then
2:	return t
3:	else
4:	let set Q to be a list of features in D
5:	randomly select a feature q_i ($q_i \in Q$)
6:	randomly select a split point $p \in (\min(q_i), \max(q_i))$
7:	Define D_l and D_r contain the sample of D , where q_i is smaller and larger than p
8:	$D_l \leftarrow \text{filter}(D, q_i < p)$
9:	$D_r \leftarrow \text{filter}(D, q_i \geq p)$
10:	repeate iTree ($D_l, h + 1, l$) and link the obtained tree as the left tree of t
11:	repeate iTree ($D_r, h + 1, l$) and link the obtained tree as the right tree of t
12:	end if

شکل ۲. شبه‌کد الگوریتم درخت ایزوله
Fig. 2. Pseudocode of the Isolation Tree Algorithm

لازم به ذکر است در پژوهش حاضر، برچسب‌های ثقلب/نرمال مستقیماً از خروجی مدل IF به دست نیامده‌اند. مدل IF صرفاً یک امتیاز ناهنجاری برای هر رکورد محاسبه کرده است. اما برای تفکیک نهایی و برچسب‌گذاری، از آستانه‌ای مبتنی بر دانش کارشناسی و تحلیل توزیع نمرات استفاده شده است. در این راستا، آستانه شناسایی ناهنجاری‌ها بر اساس نرخ‌های ارائه شده برپایه نظر کارشناسی است و نرخ رسمی و دقیقی در دسترس نیست.

ماژول ۳ (نرم/فزار)

¹ Contamination Rate

نرم افزار کشف تقلب بیمه درمان که در پژوهش حاضر طراحی شده است، «سیستم هوشمند تشخیص تقلب بیمه درمان» را به عنوان یک راه حل جامع و پیشرفته برای شرکت های بیمه معرفی می کند که مبتنی بر ترکیب یادگیری ماشین و تحلیل داده ها است. طراحی این نرم افزار از دو بخش اصلی زیر تشکیل شده است:

1- API Backend (Flask)

2- Portal Frontend (React)

بخش بک اند از زبان برنامه نویسی Python 3.x با فریم ورک Flask استفاده و معماری آن بر پایه RESTful API با Blueprint Pattern بنا شده است تا یک API قدرتمند و گسترش پذیر فراهم کند. بخش فرانت اند نیز از زبان برنامه نویسی TypeScript با فریم ورک React ۱۹ استفاده کرده است.

بخش Portal Frontend با استفاده از React و TypeScript طراحی شده است تا رابط کاربری کاربر پسند و واکنش گرا ارائه دهد. این مولفه از ابزارهای مدرن مانند Vite برای بیلد، React Router DOM برای مسیریابی، Axios برای مدیریت درخواست ها و Recharts برای نمودارها به طور یکپارچه استفاده می کند تا تجربه کاربری بهینه ای ارائه شود. با تمرکز بر مقیاس پذیری بالا و کارایی، این سیستم نه تنها نمایشگرهای گرافیکی دقیق از نتایج تحلیل ها را فراهم می کند بلکه امکان اتصال امن و کارآمد به Backend و مدیریت داده های بیمه درمان را نیز فراهم می آورد. برای ذخیره سازی داده ها از MariaDB استفاده شده و مدل تشخیص تقلب با الگوریتم IF پیاده سازی شده است تا بتواند به طور هوشمند و کارآمد الگوهای تقلبی را استخراج و تشخیص دهد. همچنین این سیستم از ۲۰ شاخص ریسک پیشرفته برای ارزیابی نسخه های پزشکی، و از مجموعه ای از کنترل های امنیتی، اعتبارسنجی و پاک سازی داده ها بهره می برد.

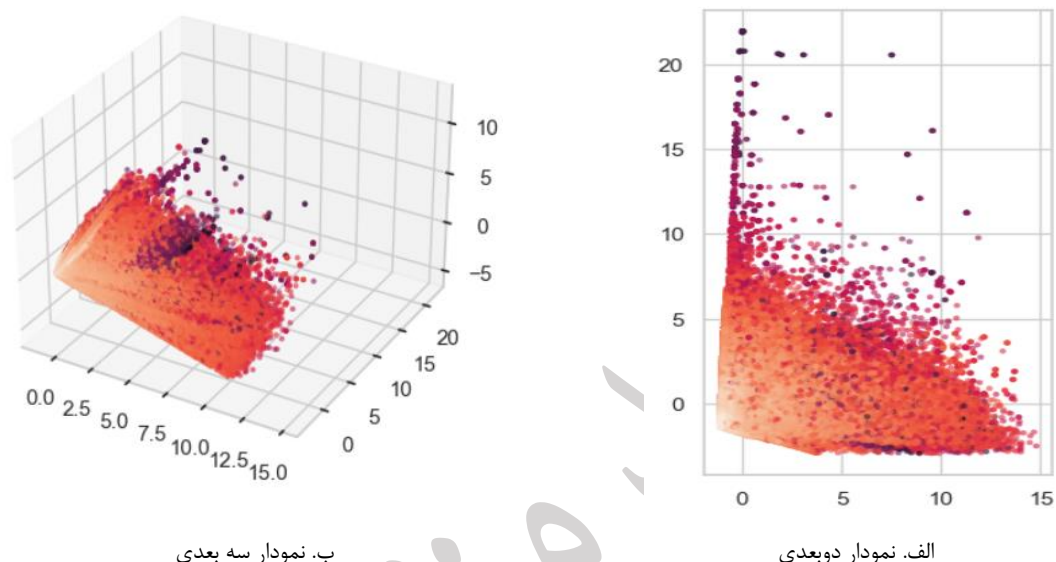
نتایج و بحث

در این مطالعه، از مجموعه داده یک شرکت بیمه مستقر در ایران جهت کشف تقلب بیمه درمان استفاده شده است. در این قسمت، این داده ها ابتدا پیش پردازش و سپس مقادیر ویژگی های تعریف شده در ماژول ۱ محاسبه خواهند شد. مجموعه داده های مورد استفاده مربوط به بخشی از پایگاه داده این شرکت است که حاوی اطلاعات مربوط به نسخه های پزشکی، بیماران، خدمات درمانی و مشخصات ارائه دهندگان مراقبت های بهداشتی (مانند پزشک، مرکز درمانی، داروخانه و غیره) می باشد. هر رکورد در این جدول نمایانگر یک نسخه و یک ادعای خسارت است، و در مجموع شامل ۷۹۳،۵۸۹ رکورد است که بازه زمانی بین اسفند ۱۴۰۲ تا اسفند ۱۴۰۳ را در بر می گیرد. این داده ها شامل نسخه های پزشکی بدون برچسب بوده که براساس محتوی آن ها مجموعه ویژگی های تعریف شده در جدول ۱ محاسبه و به عنوان ورودی الگوریتم پیشنهادی در ماژول ۲ جهت شناسایی تقلب تعریف شده است. در این مجموعه داده، متغیرهای مختلفی وجود دارد که هر یک جنبه ای از فعالیت ها و عملکرد ارائه دهندگان خدمات درمان را نشان می دهد. برخی از این متغیرها عبارتند از: کد شناسایی بیمار (ID)، تاریخ تولد، جنسیت، تاریخ پذیرش نسخه، تاریخ حواله، نوع پوشش بیمه، نوع خدمات دریافت شده (مانند ویزیت ها، عمل جراحی، خدمات اورژانس، دارو و غیره)، نام و نوع مرجع درمان، نام و تخصص پزشک یا ارائه دهنده خدمات، استان محل ارائه خدمات، کل مبلغ خسارت، و مبلغ تایید شده خسارت. مبلغ خسارت به میزان هزینه های پرداخت شده برای هر نسخه درمانی اشاره دارد. این متغیرها به مشخصات بیماران، ارائه دهندگان خدمات درمانی و خدمات ارائه شده مربوط می شود و می تواند در شناسایی الگوهای مشکوک و تقلب در خدمات درمانی مؤثر باشند.

تحلیل عملکرد چارچوب پیشنهادی

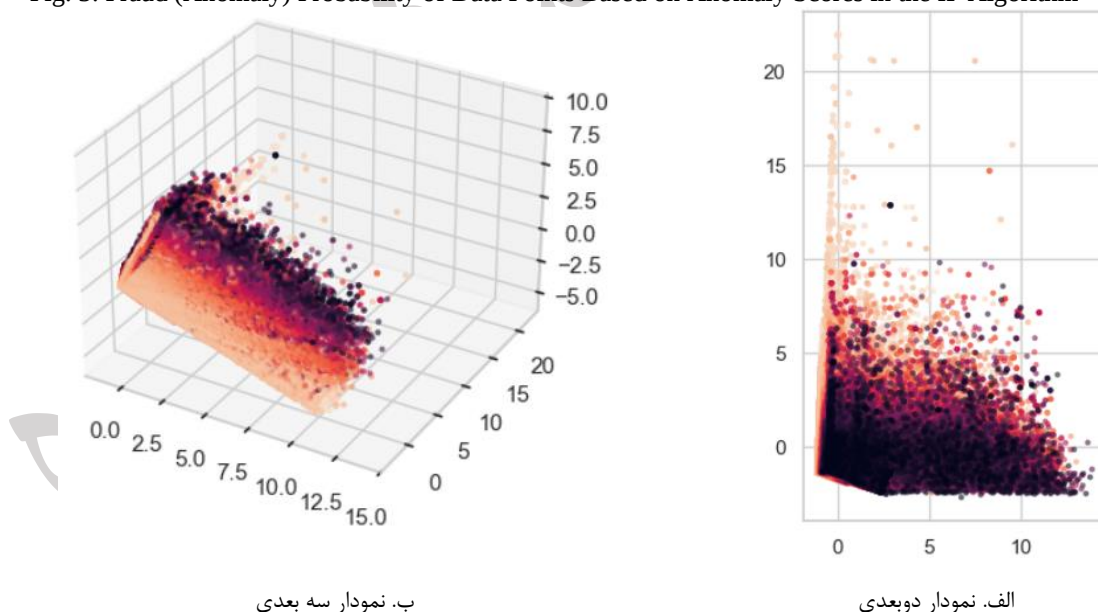
در تحلیل عملکرد الگوریتم IF در مقایسه با الگوریتم AE برای شناسایی نقاط آنومالی (تقلب)، نخست در الگوریتم IF، شاخص نمره ناهنجاری/آنومالی به هر نمونه (نسخه پزشکی) اختصاص می یابد که بیانگر احتمال تقلب در آن نمونه نسبت به سایر داده ها است؛ مقادیری بالاتر برای نمره آنومالی به معنای احتمال بیشتری برای تقلب است. این دو الگوریتم به طور مؤثری لبه ها و نقاط کم تعداد را تشخیص می دهد. دوم، در الگوریتم AE، شاخص loss یا مقدار بازسازی نامطلوب بین ورودی و بازتولید خروجی، به عنوان نشانگری از انحراف از الگوی غالب استفاده می شود؛ مقادیر بالای loss نشان می دهد که نمونه ها از الگوی بازسازی شده توسط AE فاصله زیادی دارند؛ به عبارت دیگر به خاطر تفاوت های ساختاری از مدل گرایش خارج می شوند و احتمال وجود ناهنجاری در این نمونه ها بیشتر است. به طور کلی، این دو شاخص از دو دیدگاه متفاوت برای تشخیص آنومالی استفاده می کنند: یکی مبتنی بر فاصله نمونه از گروه های عادی با توجه به ساختار داده، و دیگری مبتنی بر میزان بازتولید و تفاوت بین ورودی و خروجی در یک شبکه بازسازی. با توجه به نظر کارشناسی در ماژول اول مبنی بر اینکه حدود ۲۵ درصد نسخه ها شامل تقلب هستند، نرخ آلودگی در هر دو الگوریتم ۲۵ درصد در نظر گرفته شد. لذا، حد آستانه برای نمرات آنومالی و شاخص loss برای تفکیک نسخه های پزشکی به نرمال و تقلب در هر دو الگوریتم ۲۵ درصد خواهد بود.

برای بررسی قدرت این شاخص‌ها در شناسایی نقاط آنومالی (تقلب)، با استفاده از رویکرد تحلیل مولفه اصلی (PCA) داده‌های اصلی با ابعاد ۲۰ ویژگی به طور تقریبی به ۲ تا ۳ بعد فشرده شده‌اند تا نمودار نتایج الگوریتم IF و AE را به صورت بصری نشان داد که در شکل‌های ۳ و ۴ مشاهده می‌شوند. در این شکل‌ها، هر چه رنگ نقاط تیره‌تر باشد، احتمال وجود آنومالی/تقلب بیشتر است. می‌توان مشاهده کرد الگوریتم IF در یافتن لبه‌های داده و خوشه‌های بسیار کوچک عملکرد بهتری دارد. در الگوریتم IF، تفاوت مقادیر نسبتاً کم است، و می‌توان نقاط بیشتری را در لبه‌ها پیدا کرد که احتمال وجود آنومالی در آن‌ها بالا است. درحالی که الگوریتم AE نسبت به خوشه‌بندی‌هایی که از اکثریت دور هستند، حساسیت بالایی نشان می‌دهد؛ یعنی نقاطی که خارج از دسته‌بندی غالب داده‌ها واقع می‌شوند (یعنی نقاطی که نماینده کم‌تعداد اما با ویژگی‌های متمایز هستند) برای این مدل احتمالاً به عنوان آنومالی/تقلب تشخیص داده می‌شوند. در این نمودار، الگوریتم AE به نقاط دور از کل اکثریت داده‌ها حساس است، به عبارت دیگر تمایز میان الگوهای نادر یا خوشه‌های کم‌تعداد و خوشه‌های غالب را به طور نسبتاً قوی‌تری می‌تواند انجام بدهد.



شکل ۳. احتمال تقلب (آنومالی) نقاط براساس نمره آنومالی در الگوریتم IF

Fig. 3. Fraud (Anomaly) Probability of Data Points Based on Anomaly Scores in the IF Algorithm



شکل ۴. احتمال تقلب (آنومالی) نقاط براساس نمره **loss** در الگوریتم AE

Fig. 4. Fraud (Anomaly) Probability of Data Points Based on Loss Scores in the AE Algorithm

برای ارزیابی کارایی تشخیص تقلب در مطالعه حاضر، ابتدا از آنجا که هیچ برچسب حقیقی برای مجموعه داده در دسترس نبود، با استفاده از الگوریتم‌های IF و AE، برچسب‌های تقلب/نرمال نسخه‌های پزشکی پیش‌بینی می‌شود. دلیل مقایسه نتایج پیش‌بینی و تشخیص تقلب در دو الگوریتم

IF با الگوریتم AE، این است که این دو روش ویژگی‌های مختلفی از داده‌ها را لحاظ می‌کنند: رویکردهای مبتنی بر IF یک روش مبتنی بر نمونه-محوری و انسجام‌یابی بدنه داده‌ها است که بر پایه شدت انحراف از رفتار عادی و عدم نیاز به بازسازی داده‌ها عمل می‌کند، در حالی که AE با یادگیری بازسازی داده‌ها از ورودی‌های نرمال، خطاهای بازسازی را به عنوان نشانگرهای غیرعادی اندازه‌گیری می‌کند؛ چنین تفاوتی امکان مقایسه بین دو منطق تشخیصی و ناتوانی نسبی هر کدام را در تشخیص تقلب در یک مجموعه داده با ابعاد بالا و پیچیدگی ذاتی آن‌ها و بدون برچسب‌های کامل نشان می‌دهد.

ارزیابی الگوریتم IF

داده‌های شرکت بیمه‌ای که در قسمت پیاده‌سازی الگوریتم پیشنهادی مورد استفاده قرار گرفت، فاقد برچسب بودند و به همین دلیل امکان مقایسه مستقیم الگوریتم پیشنهادی با سایر الگوریتم‌های بدون ناظر مانند AE براساس شاخص‌های مختلف فراهم نبود. در این قسمت از یک مجموعه داده خارجی که از پروژه‌ای تحت عنوان «کشف تقلب ارائه‌دهندگان پزشکی» از آدرس (<https://www.kaggle.com/code/rohitrox/medical-provider-fraud-detection>) گرفته شده است، برای این منظور استفاده می‌شود. با توجه به حجم مناسب داده‌ها و برچسب‌دار بودن ارائه‌دهندگان، این مجموعه داده که دارای ۱۵۷ ویژگی تعریف شده است، برای ارزیابی عملکرد الگوریتم IF به کار گرفته شد. با توجه به نرخ ۲۰ درصدی تقلب در این داده‌ها حد‌آستانه در الگوریتم‌های IF و AE، ۲۰ درصد در نظر گرفته شد. این دو الگوریتم براساس معیارهای مختلف شامل دقت (Accuracy)، f1-score، صحت (Precision)، فراخوانی (Recall)، و AUC، ارزیابی و نتایج در جدول ۲ نشان داده شده‌اند.

جدول ۲. تحلیل مقایسه عملکرد مدل‌های پیش‌بینی

Table 2. Comparative analysis of the predictive models' performance

مدل	Accuracy	Precision	Recall	F1-score	AUC	زمان (ثانیه)
IF	۰.۹۵	۰.۷۷	۰.۷۷	۰.۷۷	۰.۸۷	۵۲۴
AE	۰.۷۶	۰.۶۲	۰.۳۹	۰.۴۰	۰.۵۳	۴۵۶۶

جدول مقایسه عملکرد الگوریتم‌های IF و AE در کشف تقلب بیمه درمان نشان می‌دهد که مدل IF به‌طور قابل‌توجهی عملکرد بهتری نسبت به AE دارد. دقت کلی IF برابر ۰.۹۵ بوده و در شاخص‌های Precision، Recall و F1-score نیز مقدار ۰.۷۷ را کسب کرده است، در حالی که AE تنها به دقت ۰.۷۶ و مقادیر پایین‌تری در سایر شاخص‌ها (Precision=0.62، Recall=0.39 و F1-score=0.40) دست یافته است. همچنین، مقدار AUC برای IF برابر با ۰.۸۷ است که نشان‌دهنده توان بالای آن در تفکیک نمونه‌های تقلبی از غیرتقلبی است، در حالی که AE با AUC برابر با ۰.۵۳ عملکرد ضعیف‌تری دارد. از نظر زمان اجرا نیز IF با ۵۲۴ ثانیه بسیار سریع‌تر از AE با ۴۵۶۶ ثانیه عمل کرده است. این نتایج بیانگر برتری الگوریتم جنگل ایزوله در دقت، کارایی و سرعت پردازش در زمینه کشف تقلب بیمه درمان می‌باشد.

نتایج اجرای نرم‌افزار

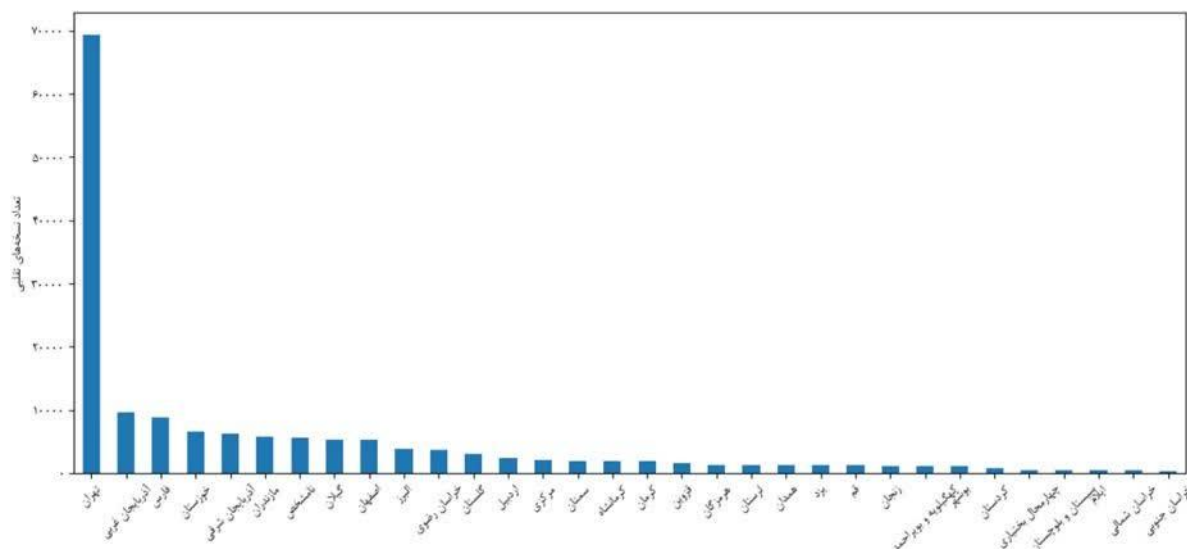
پس از تایید عملکرد الگوریتم پیشنهادی، در مرحله نهایی اقدام به تعریف نرم‌افزار کشف تقلب طبق مشخصات مازول ۳ شد. هم‌اکنون امکان بهره‌برداری از نرم‌افزار بر روی سرور اختصاصی فراهم شده و یک دامنه اینترنتی نیز برای دسترسی به آن در نظر گرفته شده است.

گزارشات نرم‌افزار

پس از آموزش الگوریتم IF و پیش‌بینی برچسب نسخه‌های پزشکی مربوط به داده‌های آموزش، امکان ارائه گزارش‌های مفیدی در نرم‌افزار جهت تحلیل نسبت نسخه‌های تقلبی بر اساس متغیرهای مختلف مانند استان، جنسیت، گروه‌های سنی، زمان پذیرش نسخه‌های پزشکی، نوع پوشش بیمه و نوع صورتحساب فراهم می‌شود. با ارائه داده‌های تفکیکی و نمودارهای بصری، این تحلیل‌ها امکان شناسایی الگوهای جغرافیایی و جمعیتی تقلب را فراهم می‌کند و نقاط داغی را که در آن‌ها احتمال تقلب بالاتر است نمایان می‌سازد. به علاوه، امکان نظارت بر تغییرات زمانی و مقایسه اثرات سیاست‌های بیمه‌ای مختلف بر نرخ تقلب وجود دارد، که می‌تواند به تصمیم‌گیری‌های مدیریتی و بهبود کنترل‌های داخلی منتهی شود. در ادامه نمودارهای بصری مربوط به این نسبت‌ها آورده شده است.

شکل ۵ و جدول ۳ تعداد نسخه‌های تقلبی بیمه درمانی به تفکیک استان‌های کشور است. نتایج نشان می‌دهد تهران با اختلاف قابل توجهی در صدر قرار دارد، پس از آن استان‌های آذربایجان غربی و فارس در رتبه‌های بعدی هستند، درحالی که خراسان جنوبی و شمالی کمترین

تعداد نسخه‌های تقلبی را گزارش کرده‌اند. با این حال، این توزیع متأثر از متغیرهای مانند تعداد جمعیت، حجم خدمات درمانی، و سطح نظارت اداری در هر استان است، لذا تفسیر قطعی نیازمند تعدیل داده‌ها براساس معیارهای نسبی است.



شکل ۵. تعداد نسخه‌های تقلبی در هر استان

Fig. 6. Number of Fraudulent Prescriptions in Each Province

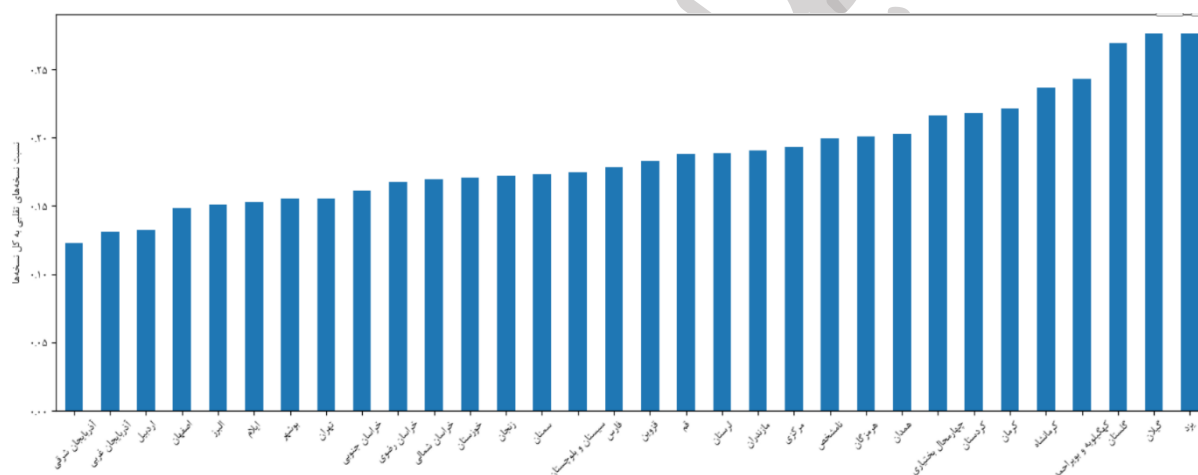
جدول ۳. داده‌های مربوط به تعداد نسخه‌های تقلبی در هر استان

Table 3. Number of Fraudulent Prescriptions in Each Province

نام استان	تعداد نسخه های تقلبی
تهران	۶۹۵۰۰
آذربایجان غربی	۹۳۶۰
فارس	۹۰۹۱
خوزستان	۶۸۱۸
آذربایجان شرقی	۵۹۰۹
مازندران	۵۹۰۰
نامشخص	۵۷۲۵
گیلان	۵۰۹۱
اصفهان	۴۹۱۰
البرز	۳۸۱۸
خراسان رضوی	۳۶۳۶
گلستان	۳۲۷۳
اردبیل	۲۳۶۳
مرکزی	۲۰۰۹
سمنان	۱۸۱۸
کرمانشاه	۱۷۲۸
کرمان	۱۶۳۶
قزوین	۱۵۴۶
هرمزگان	۱۴۵۵
لرستان	۱۳۶۲
همدان	۱۲۷۵
یزد	۱۲۴۵

۱۰۹۰	قم
۱۰۶۳	زنجان
۱۰۴۱	کهگیلویه و بویراحمد
۱۰۰۹	بوشهر
۹۰۹	کردستان
۸۲۰	چهارمحال بختیاری
۷۸۲	سیستان و بلوچستان
۶۳۱	ایلام
۵۴۳	خراسان شمالی
۴۹۱	خراسان جنوبی

شکل ۶ و جدول ۴ نتیجه تحلیل نسبت نسخه‌های تقلبی به کل نسخه‌های بیمه‌ای نشان می‌دهد استان‌های یزد، گیلان و گلستان به ترتیب در سه رتبه اول با بالاترین میزان تخلف را بوده‌اند، درحالی که آذربایجان شرقی آذربایجان غربی و اردبیل به ترتیب کمترین نسبت را به خود اختصاص داده‌اند. این الگوی جغرافیایی احتمالاً بازتابی از تفاوت‌های کارایی نظام‌های نظارتی، سطح فرهنگ بیمه‌ای و یا وجود شبکه‌های سازمان‌یافته تقلب در مناطق مختلف است. یافته‌ها بر ضرورت طراحی راهکارهای نظارتی متناسب با ویژگی‌های بومی هر استان تأکید دارد و لزوم مطالعات عمیق‌تر برای شناسایی دقیق عوامل این تفاوت‌های منطقه‌ای را نشان می‌دهد.



شکل ۶. نسبت نسخه‌های تقلبی به کل نسخه‌ها در هر استان

Fig. 7. Ratio of Fraudulent Prescriptions to Total Prescriptions in Each Province

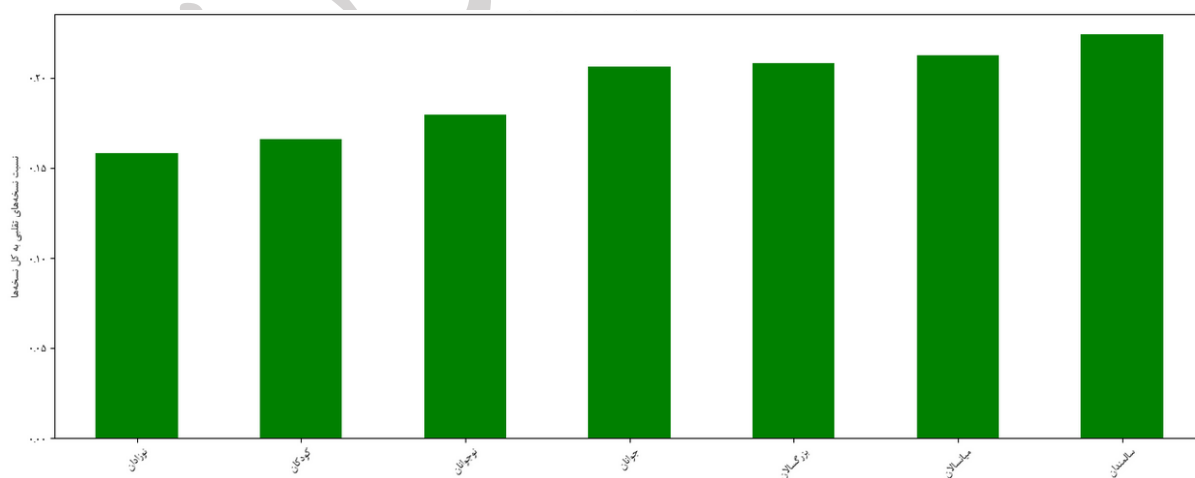
جدول ۴. داده‌های مربوط به نسبت نسخه‌های تقلبی به کل نسخه‌ها در هر استان

Table 4. Ratio of Fraudulent Prescriptions to Total Prescriptions in Each Province

نسبت نسخه‌های تقلبی به کل نسخه‌ها (%)	نام استان
۲۸.۵	یزد
۲۸.۴۲	گیلان
۲۷.۸	گلستان
۲۴.۹	کهگیلویه و بویراحمد
۲۳.۳	کرمانشاه
۲۲.۸۳	کرمان
۲۲.۲۱	کردستان
۲۱.۸۵	چهارمحال بختیاری
۲۱	همدان
۲۰.۷۱	هرمزگان

نامشخص	۲۰.۶۴
مرکزی	۱۹.۷۳
مازندران	۱۹.۵۴
لرستان	۱۹.۴۱
قم	۱۹.۲۹
قزوین	۱۸.۸۶
فارس	۱۸.۴۳
سیستان و بلوچستان	۱۸.۰۳
سمنان	۱۷.۸۳
زنجان	۱۷.۷۵
خوزستان	۱۷.۷۱
خراسان شمالی	۱۷.۵۸
خراسان رضوی	۱۷.۴۶
خراسان جنوبی	۱۶.۷۵
تهران	۱۵.۸۷
بوشهر	۱۵.۷۳
ایلام	۱۵
البرز	۱۴.۵۶
اصفهان	۱۳.۹۷
اردبیل	۱۲.۸۲
آذربایجان غربی	۱۲.۴۳
آذربایجان شرقی	۱۱.۵۶

در شکل ۷ و جدول ۵ نسبت نسخه‌های تقلبی به کل نسخه‌ها برحسب گروه‌های سنی ارایه شده است. نتایج نشان می‌دهد که سالمندان و میانسالان سهم بیشتری از نسخه‌های تقلبی را به خود اختصاص داده‌اند، در حالی که کمترین میزان تخلف مربوط به گروه سنی نوزادان و سپس کودکان بوده است. این یافته‌ها اهمیت توجه ویژه به نظارت بر نسخه‌های پزشکی مربوط به گروه‌های سنی پرخطر و همچنین طراحی مکانیزم‌های کنترلی متناسب با هر گروه سنی را برجسته می‌سازد.



شکل ۷. نسبت نسخه‌های تقلبی برحسب گروه‌های سنی
Fig. 8. Ratio of Fraudulent Prescriptions by Age Groups

Table 5. Ratio of Fraudulent Prescriptions by Age Groups

نسبت نسخه‌های تقلبی به کل نسخه‌ها (%)	گروه سنی
۲۸.۵	سال‌مندان
۲۷	میانسالان
۲۶.۶	بزرگسالان
۲۶.۲۵	جوانان
۲۲.۸	نوجوانان
۲۱.۳	کودکان
۲۰.۲۶	نوزادان

[illegible]

Fig. 9. Fraudulent Claims Ratio by Coverage Type

Table 6. Fraudulent Claims Ratio by Coverage Type

15

۲۷.۵	زایمان
۲۵	دندانپزشکی
۲۴.۶	درمان بیماری های خاص
۲۴	درمان بیماریهای اعصاب و روان
۲۳.۸۱	دارو و ملزومات دارویی
۲۳.۶	دارو داخلی
۲۳	دارو بیماران خاص و صعب العلاج
۲۱.۲۵	خدمات توانبخشی
۲۰.۷۵	خدمات اورژانس
۱۷	خدمات آزمایشگاهی
۱۵	جراحی های مجاز سرپایی
۱۳	جراحی تخصصی
۳	بیماری ها و ناهنجاری های جنین
۰	انواع ویزیت
۰	آمبولانس

جمع‌بندی و پیشنهادها

این پژوهش به طراحی و ارزیابی یک چارچوب هوشمند و ماژولار برای کشف تقلب در بیمه درمان بر پایه الگوریتم‌های یادگیری بدون نظارت می‌پردازد. با توجه به ناکارآمدی روش‌های سنتی حسابرسی در مواجهه با حجم عظیم داده‌ها و الگوهای پیچیده تقلب، در چارچوب پیشنهادی از الگوریتم IF به کار گرفته شده است و ساختاری مستقل، قابل‌پیکربندی و منطبق با ماهیت پویا و پراکنده رفتارهای غیرعادی ارائه می‌دهد. این چارچوب، علاوه بر موتور کشف تقلب، شامل ابزارهای تجسم و داشبوردهای مدیریتی است که زمان تحلیل ادعاهای پرریسک را برای کارشناسان به‌طور قابل توجهی کاهش می‌دهد. نتایج نشان داد این الگوریتم نسبت به روش AE عملکرد بهتری دارد. تحلیل‌های PCA نیز نشان داد IF قادر است نقاط بیشتری را در لبه‌ها به‌عنوان آنومالی شناسایی کند. در ماژول سوم، ابزارهای تجسم داده و داشبوردهای تعاملی امکان هدایت فرایند آموزش، تنظیم پارامترها و تحلیل رفتار بازیگران را فراهم کردند. این قابلیت‌ها باعث بهبود تدریجی مدل و انطباق پویا با الگوهای جدید تقلب می‌شود. نتایج پژوهش تأکید می‌کند که کیفیت ویژگی‌های استخراج‌شده نقش تعیین‌کننده‌ای در عملکرد الگوریتم‌ها دارد و انتخاب دقیق آنها با مشارکت کارشناسان، موجب بهبود معنادار عملکرد الگوریتم دارد. در نهایت، چارچوب پیشنهادی به‌صورت یک بسته نرم‌افزاری با معماری RESTful، پایگاه داده MariaDB و رابط کاربری مدرن مبتنی بر React توسعه داده شد. این سیستم با بهره‌گیری از ۲۰ شاخص ریسک، پاک‌سازی و اعتبارسنجی داده‌ها و موتور هوشمند IF، ابزاری قدرتمند برای شناسایی تقلب در بیمه درمان ارائه می‌دهد. معماری ماژولار، مستندات کامل و قابلیت ادغام با سامانه‌های موجود، آن را به یک راهکار عملی و قابل‌توسعه برای شرکت‌های بیمه تبدیل می‌کند.

پیشنهادهای پژوهش

پیاده‌سازی سامانه‌های هوشمند کشف تقلب مبتنی بر تحلیل الگوهای غیرعادی، به‌طور معناداری خسارت‌های ناشی از تقلب را کاهش داده و کارایی عملیاتی شرکت‌های بیمه در حوزه بیمه‌های درمان را ارتقا می‌دهد. این سامانه‌ها با تحلیل داده‌های تاریخی و شناسایی انحراف‌ها، امکان تشخیص پیش‌دستانه پرونده‌های مشکوک را فراهم کرده و دقت تصمیم‌گیری و تخصیص منابع را افزایش می‌دهند. استفاده از داشبوردهای مدیریتی نیز سرعت واکنش و اثربخشی کنترل‌های پیشگیرانه را بهبود می‌بخشد. ارتقای عملکرد سیستم مستلزم توسعه ویژگی‌های مرتبط با بازیگران، تراکنش‌ها و شبکه‌های ارتباطی آنهاست تا ریسک شبکه‌ای در تحلیل‌ها منعکس شود. به‌روزرسانی مستمر مدل‌ها و دریافت بازخورد از عملیات واقعی نیز سازگاری سیستم با الگوهای جدید تقلب را تقویت می‌کند. ترکیب داده‌های جمعیتی، مالی و خدماتی همراه با ثبت دقیق جزئیات تراکنش‌ها، پایه‌ای ضروری برای آموزش مدل‌های دقیق و قابل‌اعتماد فراهم می‌آورد. استقرار این سامانه در قالب مرکز تحلیل تقلب و ادغام آن با زیرساخت‌های موجود، با همکاری واحدهای تخصصی و بهره‌گیری از الگوریتم‌های یادگیری ماشین، شناسایی سریع و مقیاس‌پذیر تقلب را ممکن می‌سازد. مستندسازی، امنیت داده و قابلیت توسعه نیز این سیستم را به یک سرمایه راهبردی برای افزایش شفافیت و کاهش هزینه‌ها در صنعت بیمه درمان تبدیل می‌کند.

در تفسیر یافته‌های این پژوهش لازم است به چند محدودیت توجه شود. نخست آنکه کیفیت محدود داده‌های بیمه‌ای، وجود داده‌های ناقص یا نویزی، و اتکای مطالعه بر داده‌های یک شرکت بیمه در ایران، قابلیت تعمیم نتایج را کاهش می‌دهد. همچنین، محدودیت‌های قانونی و الزامات مربوط به حریم خصوصی، دسترسی به داده‌های درمانی را محدود ساخته و امکان پیاده‌سازی عملی مدل‌ها را با چالش مواجه می‌کند. کمبود داده‌های برجسب‌گذاری‌شده و تعامل محدود با کارشناسان بیمه نیز می‌تواند بر دقت الگوریتم‌ها اثرگذار باشد. علاوه بر این، پویایی الگوهای خدمات درمانی و تغییر در رفتارهای تقلب‌آمیز، ضرورت بازنگری و به‌روزرسانی مستمر مدل‌ها را برجسته می‌سازد.

در حوزه پژوهش‌های آتی، بررسی عمیق‌تر نقش کارشناسان در تعریف انواع تقلب و تعیین ویژگی‌های مؤثر می‌تواند به تقویت فرایند تصمیم‌گیری و ساختار یادگیری چارچوب پیشنهادی منجر شود. به‌کارگیری روش‌های پیشرفته مهندسی ویژگی، شامل الگوریتم‌های فراابتکاری، تحلیل‌های آماری و تکنیک‌های کاهش ابعاد، می‌تواند به انتخاب بهینه ویژگی‌ها کمک کند. همچنین توسعه و ارزیابی مدل‌های پیشرفته‌تر یادگیری ماشین - از جمله الگوریتم‌های یادگیری عمیق، ترکیبی و مبتنی بر گراف - ظرفیت شناسایی دقیق‌تر تقلب‌های فردی و گروهی را افزایش می‌دهد. بهره‌گیری از پردازش زبان طبیعی برای تحلیل داده‌های غیرساختاریافته همچون یادداشت‌های پزشکی نیز می‌تواند دامنه کاربرد چارچوب را گسترش دهد. در نهایت، بررسی روش‌های بهینه‌سازی محاسباتی و ارزیابی کارایی چارچوب در سایر حوزه‌های بیمه‌ای مانند بیمه خودرو و ریسک اعتباری، مسیر مناسبی برای توسعه و تعمیم نتایج این پژوهش فراهم می‌آورد.

- بیمه مرکزی جمهوری اسلامی. (۱۴۰۲). *سالنامه آماری صنعت بیمه*. <https://www.centinsur.ir/fa-IR/Portal/5493>
- خانی‌زاده، ف.، خامسیان، ف.، و اثنی‌عشری، م. (۱۴۰۱). کاربرد یادگیری بدون نظارت در کشف تقلبات بیمه اتومبیل (الگوریتم جنگل ایزوله). *حسابداری مدیریت*، ۵ (۵۳)، ۱۴۱-۱۵۳. <https://www.sid.ir/paper/1063243/fa>
- رحیم خانی، پ.، و منطقی پور، م. (۱۴۰۰). کلان داده‌ها و شناسایی تقلبات بیمه‌ای در رشته بیمه درمان: مرور سیستماتیک [مقاله کنفرانسی]. *بیست و هشتمین همایش ملی و نهمین همایش بین‌المللی بیمه و توسعه*، تهران، ایران. <https://civilica.com/doc/1390780>
- رحیم خانی، پ.، منطقی پور، م.، و نورانی، و. (۱۴۰۱). ساختارهای کشف تقلب بیمه‌ای در کشورهای منتخب [طرح پژوهشی، شماره ۱۶۶]. *پژوهشکده بیمه*، طرح پژوهشی، شماره ۱۶۶. <https://www.irc.ac.ir/fa-IR/Irc/4946/Articles/view/14643/1576>
- رامندی، س.، عباسی، م.، و عاشقی، ه. (۱۳۹۹). مطالعه و بررسی تقلب در بیمه‌های درمان تکمیلی و راه‌های مقابله با آن. *پژوهشکده بیمه* [طرح پژوهشی، شماره ۱۳۱]. *پژوهشکده بیمه*. <https://www.irc.ac.ir/fa-IR/Irc/4944/Articles/view/16040/1292>
- معصومی‌جناقرد، ع.، زارعی، ق.، و رحیمی‌کلور، ح. (۱۴۰۴). پیشایندها، مؤلفه‌ها و پیامدهای به‌کارگیری هوش مصنوعی در صنعت بیمه. *پژوهشنامه بیمه*. ۱۵ (۲)، ۹۵-۱۱۸. https://ijir.irc.ac.ir/article_160363.html
- Baesens, B., Höppner, S., & Verdonck, T. (2021). Data Engineering for Fraud Detection. *Decision Support Systems*, 150, 113492. <https://doi.org/10.1016/j.dss.2021.113492>
- Bagekari, A. (2025). *Claims Management Software Market Analysis (Global edition)* [Market report]. Cognitive Market Research. https://www.cognitivemarketresearch.com/claims-management-software-market-report#author_details
- Bauder, R. A., Khoshgoftaar, T. M., & Seliya, N. (2017). A Survey on the State of Healthcare Upcoding Fraud Analysis and Detection. *Health Services and Outcomes Research Methodology*, 17(1), 31–55. <https://doi.org/10.1007/s10742-016-0154-x>
- Central Insurance of the Islamic Republic of Iran. (2024). *Statistical Yearbook*. <https://www.centinsur.ir/fa-IR/Porta''4l666/5493/page> [In Persian].
- De Meulemeester, H., De Smet, F., van Dorst, J., Derroitte, E., & De Moor, B. (2025). Explainable Unsupervised Anomaly Detection for Healthcare Insurance Data. *BMC Medical Informatics and Decision Making*, 25(1), 14. <https://doi.org/10.1186/s12911-024-02602-3>
- Dhieb, N., Ghazzai, H., Besbes, H., & Massoud, Y. (2020). A Secure AI-Driven Architecture for Automated Insurance Systems: Fraud Detection and Risk Measurement. *IEEE Access*, 8, 58546–58558. <https://doi.org/10.1109/ACCESS.2020.2982410>
- du Preez, A., Bhattacharya, S., Beling, P., & Bowen, E. (2025). Fraud Detection in Healthcare Claims Using Machine Learning: A Systematic Review. *Artificial Intelligence in Medicine*, 160, 103061. <https://doi.org/10.1016/j.artmed.2024.103061>
- Ekina, T., Leva, F., Ruggeri, F., & Soyer, R. (2013). Application of Bayesian Methods in Detection of Healthcare Fraud. *Chemical Engineering Transactions*, 33, 1291–1296. <https://doi.org/10.3303/CET1333216>

- Gomes, C., Jin, Z., & Yang, H. (2021). Insurance Fraud Detection with Unsupervised Deep Learning. *Journal of Risk and Insurance*, 88(3), 591-624. <https://doi.org/10.1111/jori.12359>
- Haddad Soleymani, M., Yaseri, M., Farzadfar, F., Mohammadpour, A., Sharifi, F., & Kabir, M. J. (2018). Detecting Medical Prescriptions Suspected of Fraud Using an Unsupervised Data Mining Algorithm. *DARU Journal of Pharmaceutical Sciences*, 26(2), 209–214. <https://doi.org/10.1007/s40199-018-0225-1>
- Hamid, Z., Khalique, F., Mahmood, S., Daud, A., Bukhari, A., & Alshemaimri, B. (2024). Healthcare Insurance Fraud Detection Using Data Mining. *BMC Medical Informatics and Decision Making*, 24(1), 112. <https://doi.org/10.1186/s12911-024-02515-1>
- Johnson, M. E., & Nagarur, N. (2016). Multi-Stage Methodology to Detect Health Insurance Claim Fraud. *Health Care Management Science*, 19(3), 249–260. <https://doi.org/10.1007/s10729-015-9319-1>
- Khanizadeh, F., Khamsian, F., & Asni Ashari, M. (2022). Application of Unsupervised Learning in Detecting Auto Insurance Fraud (Isolation Forest algorithm). *Journal of Management Accounting*, 15(53), 141–153 . [SID. https://sid.ir/paper/1063243/en](https://sid.ir/paper/1063243/en) [In Persian].
- Kose, I., Gokturk, M., & Kilic, K. (2015). An Interactive Machine-Learning-Based Electronic Fraud and Abuse Detection System in Healthcare Insurance. *Applied Soft Computing*, 36, 283–299. <https://doi.org/10.1016/j.asoc.2015.07.018>
- Kou, Y., Lu, C.-T., Sirwongwattana, S., & Huang, Y.-P. (2004). Survey of Fraud Detection Techniques. In *Proceedings of the 2004 IEEE International Conference on Networking, Sensing and Control* (Vol. 2, pp. 749–754). IEEE. <https://doi.org/10.1109/ICNSC.2004.1297070>
- Leevy, J. L., Salekshahrezaee, Z., & Khoshgoftaar, T. M. (2024). A Review of Unsupervised Anomaly Detection Techniques for Health Insurance Fraud. In *2024 IEEE 10th International Conference on Big Data Computing Service and Machine Learning Applications (BigDataService)* (pp. 141–149). IEEE. <https://doi.org/10.1109/BigDataService63001.2024.00027>
- Lin, E., Chen, Q., & Qi, X. (2020). Deep Reinforcement Learning for Imbalanced Classification. *Applied Intelligence*, 50(8), 2488–2502. <https://doi.org/10.1007/s10489-020-01637-x>
- Lu, J., Lin, K., Chen, R., Lin, M., Chen, X., & Lu, P. (2023). Health Insurance Fraud Detection by Using an Attributed Heterogeneous Information Network with a Hierarchical Attention Mechanism. *BMC Medical Informatics and Decision Making*, 23(1), 62. <https://doi.org/10.1186/s12911-023-02156-w>
- Luo, W., & Gallagher, M. (2010). Unsupervised DRG Upcoding Detection in Healthcare Databases. In *2010 IEEE International Conference on Data Mining Workshops* (pp. 600–605). IEEE. <https://doi.org/10.1109/ICDMW.2010.56>
- Major, J. A., & Riedinger, D. R. (1992). EFD: A Hybrid Knowledge/Statistical-Based System for the Detection of Fraud. *International Journal of Intelligent Systems*, 7(7), 687–703. <https://doi.org/10.1002/int.4550070707>
- Masoumi Janaghari, I., Zarei G., and Rahimi Kluor, H. (2025). Antecedents, Components, and Consequences of Applying Artificial Intelligence in the Insurance Industry. *Iranian Journal of Insurance Research*. https://ijir.irc.ac.ir/article_160363.html [In Persian].
- Matloob, I., Khan, S., Rukaiya, R., Alfrahi, H., & Ali Khan, J. (2025). Healthcare Fraud Detection Using Adaptive Learning and Deep Learning Techniques. *Evolving Systems*, 16(2), 72. <https://doi.org/10.1007/s12530-024-09591-8>

- Muscoloni, A., Thomas, J. M., Ciucci, S., Bianconi, G., & Cannistraci, C. V. (2017). Machine Learning Meets Complex Networks Via Coalescent Embedding in the Hyperbolic Space. *Nature Communications*, 8(1), 1615. <https://doi.org/10.1038/s41467-017-01825-5>
- Phua, C., Lee, V., Smith, K., & Gayler, R. W. (2010). A Comprehensive Survey of Data Mining-Based Fraud Detection Research (arXiv:1009.6119) [Preprint]. arXiv. <https://doi.org/10.48550/arXiv.1009.6119>
- Price, M., & Norris, D. M. (2009). Health Care Fraud: Physicians as White Collar Criminals? *Journal of the American Academy of Psychiatry and the Law Online*, 37(3), 286–289. <https://jaapl.org/content/37/3/286>
- Puggini, L., & McLoone, S. (2018). An Enhanced Variable Selection and Isolation Forest-Based Methodology for Anomaly Detection with OES Data. *Engineering Applications of Artificial Intelligence*, 67, 126–135. <https://doi.org/10.1016/j.engappai.2017.09.014>
- Rahimkhani, P., & Manteghipour, M. (2021). A Systematic Review of Big Data Applications for Detecting Fraud in Health Insurance. *The 28th National and 9th International Conference on Insurance and Development*, Tehran, Iran. <https://civilica.com/doc/1390780> [In Persian].
- Rahimkhani, P., Manteghipour, M., & Nourani, V. (2022). Insurance Fraud Detection Frameworks in Selected Countries. Insurance Research Center, Research Project, No. 166. <https://www.irc.ac.ir/fa-IR/Irc/4946/Articles/view/14643/1576> [In Persian].
- Ramandi, S., Abbasi, M., & Asheghi, H. (2020). Study and Analysis of Fraud in Supplementary Health Insurance and Strategies to Combat it. Insurance Research Center, Research Project No. 131. <https://www.irc.ac.ir/fa-IR/Irc/4944/Articles/view/16040/1292> [In Persian].
- Samariya, D., Ma, J., Aryal, S., & Zhao, X. (2023). Detection and Explanation of Anomalies in Healthcare Data. *Health Information Science and Systems*, 11(1), 20. <https://doi.org/10.1007/s13755-023-00222-1>
- Shan, Y., Jeacocke, D., Murray, D. W., & Sutinen, A. (2008, November). Mining Medical Specialist Billing Patterns for Health Service Management. In *Proceedings of the 7th Australasian Data Mining Conference*-Volume 87 (pp. 105-110). Australian Computer Society, Inc. <https://dl.acm.org/doi/10.5555/2449288.2449306>
- Sparrow, M. K. (2009). Fraud in the U.S. health-care system: Exposing the vulnerabilities of automated payments systems. *Social Research: An International Quarterly*, 75(4), 1151–1180. <https://doi.org/10.1353/sor.2008.0015>
- Suesserman, M., Gorny, S., Lasaga, D., Helms, J., Olson, D., Bowen, E., & Bhattacharya, S. (2023). Procedure Code Overutilization Detection From Healthcare Claims Using Unsupervised Deep Learning Methods. *BMC Medical Informatics and Decision Making*, 23(1), 196. <https://doi.org/10.1186/s12911-023-02299-w>
- Suroor, N., & Misra, T. (2024). Medical Insurance Fraud Detection. In *Deep Learning in Internet of Things for Next Generation Healthcare* (pp. 182–193). Chapman and Hall/CRC. <https://www.taylorfrancis.com/chapters/edit/10.1201/9781003451846-15>
- Szalados, J. E. (2021). Regulations and Regulatory Compliance: False Claims Act, Kickback and Stark Laws, and HIPAA. In *The Medical-Legal Aspects of Acute Care Medicine: A Resource for Clinicians, Administrators, and Risk Managers* (pp. 277–313). Springer International Publishing. https://doi.org/10.1007/978-3-030-68570-6_16

- Tabassum, M., Mahmood, S., Bukhari, A., Alshemaimri, B., Daud, A., & Khalique, F. (2024). Anomaly-Based Threat Detection in Smart Health Using Machine Learning. *BMC Medical Informatics and Decision Making*, 24(1), 347. <https://doi.org/10.1186/s12911-024-02658-1>
- Thornton, D., Brinkhuis, M., Amrit, C., & Aly, R. (2015). Categorizing and Describing the Types of Fraud in Healthcare. *Procedia Computer Science*, 64, 713–720. <https://doi.org/10.1016/j.procs.2015.08.598>
- Verma, A., Taneja, A., & Arora, A. (2017). Fraud Detection and Frequent Pattern Matching in Insurance Claims Using Data Mining Techniques. In *2017 Tenth International Conference on Contemporary Computing (IC3)* (pp. 1–7). IEEE. <https://doi.org/10.1109/IC3.2017.8284343>
- Viaene, S., Derrig, R. A., Baesens, B., & Dedene, G. (2002). A Comparison of State-of-the-Art Classification Techniques for Expert Automobile Insurance Claim Fraud Detection. *Journal of Risk and Insurance*, 69(3), 373–421. <https://doi.org/10.1111/1539-6975.00023>
- Williams, G. J. (1999). Evolutionary Hot Spots Data Mining: An Architecture for Exploring for Interesting Discoveries. In *Pacific-Asia Conference on Knowledge Discovery and Data Mining* (pp. 184–193). Springer Berlin Heidelberg. <https://www.researchgate.net/publication/220894448>
- Wynia, M. K., Cummins, D. S., VanGeest, J. B., & Wilson, I. B. (2000). Physician Manipulation of Reimbursement Rules for Patients: Between a Rock and a Hard Place. *JAMA*, 283(14), 1858–1865. <https://doi.org/10.1001/jama.283.14.1858>
- Yamanishi, K., Takeuchi, J. I., Williams, G., & Milne, P. (2000). On-Line Unsupervised Outlier Detection Using Finite Mixtures with Discounting Learning Algorithms. In *Proceedings of the Sixth ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 320–324). ACM. <https://doi.org/10.1145/347090.347168>