

Paper Type: Original Article



Detecting Suspicious Transactions in Banking Cyber Physical Systems Using Artificial Intelligence and Machine Learning

Seyyed Ruhollah MirHosseini^{1,*}, Behrouz Minaei Bidgoli² , Bahareh Shaker Ardakani³

¹ Department of Computer Engineering, Artificial Intelligence and Robotics, North Tehran Branch, Islamic Azad University, Tehran, Iran; mirhoseini@alumni.iust.ac.ir.

² Department of Computer Engineering, University of Science and Technology, Tehran, Iran; b_minaei@iust.ac.ir.

³ Specialist at Behsazan Mellat Company, Iran; bshaker@chmail.ir.

Citation:

MirHosseini, S. R., Minaei Bidgoli, B., & Shaker Ardakani, B. (2025). Detecting suspicious transactions in banking cyber physical systems using artificial intelligence and machine learning. *Management sciences and decision analysis*, 3(1), 24-32.

Received: 12/03/2024

Reviewed: 17/15/2024

Revised: 27/06/2024

Accepted: 22/07/2024

Abstract

Purpose: The purpose of this study is to detect suspicious transactions in banking cyber-physical systems by employing artificial intelligence and machine learning techniques. Given the growth of digital transactions and the increasing complexity of fraud, the study aims to develop an efficient and reliable anomaly detection algorithm for real-time monitoring.

Methodology: This research utilizes both supervised (e.g., decision trees, SVM, neural networks) and unsupervised learning methods (e.g., clustering and anomaly detection). Evaluation metrics such as precision, recall, true positive rate, and AUC-ROC were used to assess model performance. The proposed system was simulated using real banking datasets and represented through algorithmic flowcharts.

Findings: Results indicate that AI-driven hybrid algorithms can accurately detect suspicious banking transactions. The model demonstrated high accuracy, effective detection of contextual, point, and collective anomalies, and strong performance in real-time data analysis. Integration of NLP techniques further improved the predictive accuracy of the system.

Originality / Value: This study introduces an innovative AI-based framework capable of analyzing large-scale banking data swiftly and accurately. Compared to traditional fraud detection methods, the proposed approach enhances security in banking cyber-physical systems by being more efficient and adaptable. Additionally, the research contributes to the intersection of data governance and cybersecurity.

Keywords: Cyber-physical systems, Anomaly, Suspicious transactions, Artificial intelligence, Machine learning.



Corresponding Author: mirhoseini@alumni.iust.ac.ir



10.22105/msda.v3i1.78



Licensee. **Management Sciences and Decision Analysis**. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0>).



شناسایی تراکنش‌های مشکوک در سیستم‌های فیزیکی سایبری بانکی با استفاده از هوش مصنوعی و یادگیری ماشین

سید روح‌اله میرحیسنی^۱، بهروز مینایی بیدگلی^۲، بهاره شاکر اردکانی^۳

^۱گروه مهندسی کامپیوتر، گرایش هوش مصنوعی و رباتیک، واحد تهران شمال، دانشگاه آزاد اسلامی، تهران، ایران.

^۲گروه مهندسی کامپیوتر، دانشگاه علم و صنعت، تهران، ایران.

^۳شرکت به‌سازان ملت، ایران.

چکیده

هدف: این پژوهش با هدف شناسایی تراکنش‌های مشکوک در سیستم‌های سایبری فیزیکی بانکی با بهره‌گیری از روش‌های هوش مصنوعی و یادگیری ماشین انجام شده است. با توجه به گسترش تراکنش‌های دیجیتال و افزایش پیچیدگی حملات کلاهبرداری، هدف اصلی، ارائه الگوریتمی کارآمد و قابل اعتماد برای تشخیص ناهنجاری‌های بانکی در زمان واقعی است.

روش‌شناسی پژوهش: مطالعه حاضر از روش‌های یادگیری ماشین، از جمله الگوریتم‌های یادگیری با نظارت (مانند درخت تصمیم، *SVM* و شبکه‌های عصبی و بدون نظارت (مانند خوشه‌بندی و تشخیص ناهنجاری) استفاده کرده است. از معیارهایی مانند دقت، حساسیت، نرخ مثبت واقعی و سطح زیر منحنی *ROC* برای ارزیابی عملکرد مدل‌ها بهره گرفته شد. سیستم پیشنهادی با داده‌های بانکی واقعی شبیه‌سازی شده و در قالب نمودارهای الگوریتمی ارائه گردید.

یافته‌ها: یافته‌ها نشان دادند که الگوریتم‌های ترکیبی بر پایه یادگیری ماشین قادر به تشخیص دقیق تراکنش‌های مشکوک در سیستم‌های سایبری فیزیکی بانکی هستند. مدل ارائه‌شده دارای دقت بالا، نرخ تشخیص مثبت مناسب و قابلیت شناسایی ناهنجاری‌های متنی، نقطه‌ای و جمعی می‌باشد. همچنین نتایج حاصل نشان می‌دهد که استفاده از تحلیل برخط و پردازش زبان طبیعی، نقش مهمی در ارتقای دقت مدل‌های پیش‌بینی دارد.

اصالت / ارزش افزوده علمی: نوآوری اصلی این تحقیق، طراحی و پیاده‌سازی چارچوبی مبتنی بر هوش مصنوعی است که قابلیت تحلیل سریع و دقیق داده‌های حجیم بانکی را دارد. این چارچوب، با توجه به شرایط خاص سیستم‌های سایبری فیزیکی بانکی، امنیت اطلاعات مالی را ارتقا می‌بخشد و نسبت به روش‌های سنتی تشخیص تقلب، کارآمدتر و منعطف‌تر عمل می‌کند. همچنین، تحقیق با نگاهی بین‌رشته‌ای به مقوله حکمرانی داده و امنیت سایبری می‌پردازد.

کلیدواژه‌ها: سیستم‌های سایبری فیزیکی، ناهنجاری، تراکنش‌های مشکوک، هوش مصنوعی، یادگیری ماشین.

۱- مقدمه

پیشگیری از تقلب چالشی حیاتی برای موسسات مالی، مشاغل و دولت‌ها در سراسر جهان است. ظهور تراکنش‌های دیجیتال و سیستم‌های مالی پیچیده منجر به فعالیت‌های کلاهبرداری پیچیده‌تر شده است. هوش مصنوعی راه‌حل‌های نوآورانه‌ای را برای این مشکل رو به رشد ارائه می‌دهد و از توانایی خود برای تجزیه و تحلیل حجم وسیعی از داده‌ها، شناسایی الگوها و پیش‌بینی رفتارهای متقلبانه با دقت بالا استفاده می‌کند. تکنیک‌های

هوش مصنوعی مانند یادگیری ماشینی^۱، یادگیری عمیق و پردازش زبان طبیعی، انقلابی در تشخیص و پیشگیری از تقلب به وجود آورده‌اند. الگوریتم‌های یادگیری ماشین، به‌ویژه مدل‌های یادگیری تحت نظارت مانند درخت‌های تصمیم‌گیری و شبکه‌های عصبی، به‌طور گسترده برای شناسایی تراکنش‌های جعلی با یادگیری از داده‌های تاریخی استفاده می‌شوند. این مدل‌ها می‌توانند با شناسایی الگوهای ظریفی که ممکن است توسط سیستم‌های مبتنی بر قوانین سنتی نادیده گرفته شوند، بین تراکنش‌های مشروع و جعلی تمایز قایل شوند. روش‌های یادگیری بدون نظارت، از جمله خوشه‌بندی و تشخیص ناهنجاری، برای شناسایی طرح‌های کلاهبرداری جدید با شناسایی موارد پرت در داده‌های تراکنش که با رفتار مورد انتظار مطابقت ندارند، استفاده می‌شوند. یادگیری عمیق، زیرمجموعه‌ای از یادگیری ماشینی، به دلیل توانایی آن در پردازش و تجزیه و تحلیل داده‌های بدون ساختار مانند تصاویر، متن و صدا، نوید استثنایی در تشخیص تقلب نشان داده است. تکنیک‌هایی یادگیری ماشینی در برنامه‌های کاربردی مختلف از شناسایی تقلب در کارت اعتباری تا تلاش‌های ضد پول‌شویی مورد استفاده قرار می‌گیرند. پردازش زبان طبیعی با تجزیه و تحلیل داده‌های متنی، مانند ایمیل‌ها و توضیحات تراکنش‌ها، برای شناسایی زبان و الگوهای مشکوک به شناسایی فعالیت‌های متقلبانه کمک می‌کند. کاربرد هوش مصنوعی در پیشگیری از تقلب فراتر از تشخیص و اقدامات پیشگیرانه است. تجزیه و تحلیل‌های پیش‌بینی‌شده توسط هوش مصنوعی می‌تواند نقاط بالقوه کلاهبرداری را پیش‌بینی کند و به سازمان‌ها اجازه می‌دهد تا استراتژی‌های پیشگیرانه را پیاده‌سازی کنند. سیستم‌های پیش‌بینی درنگ که توسط هوش مصنوعی تقویت شده‌اند، هشدارهای آنی را برای فعالیت‌های مشکوک ارایه می‌دهند و اقدام سریع برای کاهش تقلب را ممکن می‌سازند. با تکامل فناوری‌های هوش مصنوعی، نقش آن‌ها در حفاظت از سیستم‌های مالی و کاهش خسارات ناشی از تقلب بیشتر خواهد شد و اهمیت ادامه نوآوری و تحقیق در این زمینه را نشان می‌دهد.

در عصر حاضر، گسترش تراکنش‌های برخط، تجارت الکترونیک و بانکداری دیجیتال فرصت‌های جدیدی را برای کلاهبرداران ایجاد کرده است تا از آسیب‌پذیری‌های سیستم‌های مالی سو استفاده کنند. جرایم سایبری با طرح‌های پیچیده‌تر و فزاینده‌ای افراد، مشاغل و دولت‌ها را هدف قرار می‌دهند و در حال افزایش است. این طرح‌ها از حملات فیشینگ و سرقت هویت گرفته تا اشکال پیچیده‌تر کلاهبرداری مالی مانند تصاحب حساب‌ها و پول‌شویی را شامل می‌شود. تکامل سریع این فعالیت‌های متقلبانه چالش‌های مهمی را برای روش‌های سنتی کشف و پیشگیری از تقلب ایجاد می‌کند که اغلب برای همگام شدن با چابکی و نبوغ مجرمان سایبری مدرن تلاش می‌کنند. استراتژی‌های موثر پیشگیری از تقلب برای حفاظت از سیستم‌های مالی، حفاظت از اعتماد مصرف‌کننده و تضمین ثبات فعالیت‌های اقتصادی بسیار مهم هستند. زیان‌های مالی مرتبط با تقلب می‌تواند هم برای افراد و هم برای سازمان‌ها ویرانگر باشد و منجر به تاثیر اقتصادی قابل توجه و آسیب به اعتبار شود [1]، [2]. علاوه بر این، نهادهای نظارتی به‌طور فزاینده‌ای بر نیاز به مکانیسم‌های قوی پیشگیری از تقلب برای رعایت الزامات قانونی سختگیرانه تأکید می‌کنند. اجرای استراتژی‌های موثر پیشگیری از تقلب نه تنها به کاهش خسارات مالی کمک می‌کند، بلکه انعطاف‌پذیری موسسات مالی را در برابر حملات احتمالی تقویت می‌کند. این یک محیط دیجیتال امن را تضمین می‌کند و اعتماد و اطمینان را در بین مشتریان و ذینفعان تقویت می‌کند [3].

هوش مصنوعی به‌عنوان یک تغییردهنده بازی در حوزه تشخیص و پیشگیری از تقلب ظاهر شده است. هوش مصنوعی با استفاده از قدرت یادگیری ماشین، تجزیه و تحلیل داده‌ها و مدل‌سازی، ابزارهای پیچیده‌ای را برای شناسایی و کاهش فعالیت‌های متقلبانه در زمان واقعی ارایه می‌کند [4-6]. سیستم‌های مبتنی بر هوش مصنوعی می‌توانند حجم وسیعی از داده‌ها را با سرعتی بی‌سابقه تجزیه و تحلیل کنند و الگوهای پنهان و ناهنجاری‌هایی را که روش‌های سنتی ممکن است نادیده بگیرند، آشکار کنند. تکنیک‌هایی مانند یادگیری تحت نظارت و بدون نظارت، شبکه‌های عصبی و پردازش زبان طبیعی امکان توسعه مدل‌های پیشرفته تشخیص تقلب را فراهم می‌کند که به‌طور مداوم یاد می‌گیرند و با تهدیدهای نوظهور سازگار می‌شوند. با خودکارسازی و افزایش دقت فرآیندهای کشف تقلب، هوش مصنوعی به سازمان‌ها کمک می‌کند تا یک قدم جلوتر از کلاهبرداران باقی بمانند و اقدامات موثرتر و کارآمدتر پیشگیری از تقلب را تضمین کنند [3].

۲- مفاهیم پایه

از آنجایی که در این تحقیق با نگاه حاکمیت داده قصد استفاده از روش‌های هوش مصنوعی و یادگیری ماشین برای پیش‌بینی و پیشگیری ناهنجاری در سیستم بانکی استفاده نمایم، ابتدا این مفاهیم پایه‌ای معرفی می‌گردد.

¹ Machine Learning (ML)

۲-۱- سیستم‌های سایبر فیزیکی بانکی

سیستم‌های سایبری فیزیکی^۱، سیستم‌های ناهمگن، فدرال، پراکنده جغرافیایی و در مقیاس بزرگ هستند که شامل شبکه و اجزای کنترل‌کننده هستند. چنین سیستم‌هایی دارای بخش‌های شبکه بی‌سیم، اجزای قدیمی، ترافیک شبکه قابل پیش‌بینی، نیازمندی‌های پیچیده شبکه و حلقه‌های کنترل مختلف هستند. این سیستم ترکیب دو حوزه فیزیکی (محرک‌ها و حسگرها) و سایبری (سرورها و اجزای شبکه) می‌باشد. در دوران معاصر، همکاری بین سیستم‌های فیزیکی و محاسبات تکنولوژیکی در حال گسترش می‌باشد. سیستم‌های سایبری فیزیکی از مجموعه‌ای از دستگاه‌هایی ساخته شده است که با یکدیگر ارتباط برقرار می‌کنند و با دنیای فیزیکی نیز در تعامل هستند. این سیستم‌ها از دیدگاه ارتباط و محاسبات (با نظارت و رویه‌های کنترل) هماهنگ می‌شوند. سیستم‌های سایبری فیزیکی دارای سه مولفه اصلی ارتباطات و محاسبات^۲، کنترل^۳ و دست‌کاری و نظارت^۴ می‌باشد. سیستم‌های سایبری فیزیکی با چارچوب‌های مختلف مانند مراکز کنترل و اجزای سطح پایین که از کانال بی‌سیم برای ارتباط استفاده می‌کنند، هماهنگ می‌شوند [7].

در کنار مزایای مختلف سیستم‌های سایبری فیزیکی، چالش‌ها و مسایل امنیتی متعددی وجود دارد که باید با آن‌ها کنار آمد. این روزها بخش زیادی از مبادلات تجاری از طریق کارت‌های اعتباری انجام می‌شود. چنین چارچوب‌هایی برای تکنیک‌ها و حملات جدید کلاهبرداری در سطحی ترسناک باز هستند. کشف تقلب، به‌ویژه در جامعه مالی، به‌طور مداوم با چند مشکل مواجه است و حل آن‌ها یک کار پیچیده است. با این حال، تحقیقات مختلفی علیه کشف تقلب وجود دارد که هنوز قابل اعتماد، سریع و ایمن نیست. به همین ترتیب، شناسایی تراکنش‌های مشکوک در سیستم‌های سایبری فیزیکی بانکی^۵ یکی از چشم‌اندازهای حیاتی این روزها است [7]. کلاهبرداری در سیستم‌های پولی برای منافع مالی یک انجمن یا فرد به‌عنوان سو استفاده از سیستم تلقی می‌شود. در این دوره رقابتی، ارایه نادرست یا جعل ممکن است منجر به مسایل مخرب کسب‌وکار شود و یک موضوع بسیار جدی می‌باشد.

این کلاهبرداری‌های اعتباری به‌طور قابل توجهی گسترش یافته است و سالانه ضررهای میلیارد دلاری را برای صاحبان بانک‌ها و سازمان‌های مالی وارد می‌سازد. سیستم‌های سایبری فیزیکی بانکی یک بخش مهم از حوزه مالی فعلی است. جایی که هر فرد باید به‌صورت آنلاین یا فیزیکی با آن تعامل داشته باشد. کارایی، سودآوری و بهره‌وری بخش خصوصی و دولتی به دلیل سیستم‌های سایبری فیزیکی بانکی به‌شدت افزایش یافته است. این روزها بخش بزرگی از مبادلات تجاری از طریق کارت‌های اعتباری انجام می‌شود. پس شرایط را برای کلاهبرداری با تکنیک‌های جدید و حملات ترسناک ممکن می‌سازد. این روزها شناسایی تراکنش‌های مشکوک در سیستم‌های سایبری فیزیکی بانکی یک چشم‌انداز بسیار مهم است، زیرا یکی از بخش‌های مالی مهم زندگی روزمره، بخش بانکی است. با افزایش تهدیدات امنیتی، نیاز به بررسی و ادغام رویکردهای هوش مصنوعی کارآمد در سیستم‌های سایبری فیزیکی بانکی و ایجاد یک رویکرد قابل اعتمادتر، کارآمدتر و سریع‌تر در برابر شناسایی تراکنش‌های مشکوک وجود دارد. شناسایی تراکنش مشکوک در سیستم‌های سایبری فیزیکی بانکی برای افزایش ایمنی و قابلیت اطمینان برای اهداف راستی‌آزمایی و کشف تقلب مورد نیاز است [7].

۲-۲- ناهنجاری و انواع آن

بسیاری از نویسندگان تعاریف مختلفی را برای ناهنجاری ارایه کرده‌اند و تعریف جامعی وجود ندارد. تعاریف دقیق یک ناهنجاری به فرضیات مربوط به ساختار داده‌ها و کاربرد مورد نظر بستگی دارد. با این حال، تعاریفی وجود دارد که صرف نظر از تنظیمات یا کاربرد، برای اکثر موارد و نه برای همه موارد کلی در نظر گرفته می‌شود. از میان این تعاریف، شناخته‌شده‌ترین تعاریف توسط هاوکنینز است که مفهوم ناهنجاری یا پرت را در این مورد چنین تعریف می‌کند: «مشاهده پرت، مشاهده‌ای است که به قدری از مشاهدات دیگر منحرف می‌شود که باعث ایجاد سوطن می‌شود که توسط مکانیسم متفاوتی ایجاد شده است» [8]. این تعریف به داده‌هایی از یک شهود مبتنی بر آمار اشاره دارد که در آن داده‌های عادی از یک مکانیسم

¹ Cyber Physical Systems (CPS)² Communication and computation³ Control⁴ Manipulation and monitoring⁵ Banking Cyber Physical Systems (BCPS)

تولید پیروی می‌کنند و ناهنجاری‌ها نمونه‌ها یا نمونه‌هایی هستند که از این مکانیسم منحرف می‌شوند؛ بنابراین، ناهنجاری‌ها اغلب اطلاعات مفیدی را در مورد ویژگی‌های غیرعادی یک سیستم که بر مکانیسم تولید تاثیر می‌گذارند، منتقل می‌کنند [9].

تشخیص ناهنجاری شباهت‌هایی به حذف نویز دارد که با نویز ناخواسته در داده‌ها سروکار دارد، اما این دو از یکدیگر متمایز هستند [10]. در برنامه‌های کاربردی دنیای واقعی، داده‌ها ممکن است تحت تاثیر مقدار قابل‌توجهی از نویز قرار گیرند که ممکن است برای تحلیلگر جالب نباشد، اما در مرحله تجزیه و تحلیل داده‌ها به‌عنوان یک مانع عمل می‌کند. معمولاً فقط انحرافات قابل‌توجه جالب توجه هستند [9]. اثر مخرب نویز بر تجزیه و تحلیل داده‌ها، نیاز به حذف نویز را تحریک می‌کند، زیرا هر گونه اشیاء ناخواسته را قبل از تجزیه و تحلیل داده‌ها حذف می‌کند [11]. تشخیص تازگی موضوع دیگری است که با تشخیص ناهنجاری مرتبط است اما متمایز از آن است. در عوض، این نوع تکنیک‌ها الگوهای مشاهده نشده یا جدید را در داده‌ها شناسایی می‌کنند و به‌طور کلی، تمایز اصلی از تشخیص ناهنجاری این است که الگوهای جدید معمولاً پس از تشخیص در یک مدل گنجانده می‌شوند [12].

ناهنجاری‌ها به دلایل مختلفی مانند خطای انسانی، خطاهای ابزار، انحرافات طبیعی در جمعیت‌ها، فعالیت‌های متقلبان، تغییرات رفتاری سیستم‌ها یا خطاهای درون یک سیستم رخ می‌دهد [13]. نحوه برخورد سیستم تشخیص ناهنجاری با ناهنجاری بستگی به نوع کاربرد دارد [12]. ناهنجاری‌ها را می‌توان به سه دسته مختلف طبقه‌بندی کرد و نوع ناهنجاری‌ای که با آن برخورد می‌شود، جنبه مهمی از توجه برای هر تکنیک تشخیص ناهنجاری است. ناهنجاری‌های نقطه‌ای ساده‌ترین نوع ناهنجاری هستند و تمرکز اکثر تحقیقات در مورد استفاده از تکنیک‌های تشخیص ناهنجاری هستند. آن‌ها به‌عنوان نمونه‌های فردی از داده‌ها یا رویدادها به‌طور قابل‌توجهی متفاوت از بقیه داده‌ها مشخص می‌شوند. به‌طور کلی، آن‌ها نقاطی هستند که در محدوده عادی یا جداکننده تعیین شده قرار نمی‌گیرند [12]. ناهنجاری‌های متنی نمونه‌هایی از داده‌ها هستند که در یک زمینه خاص غیرعادی هستند و در غیر این صورت نیستند. ساختار مجموعه داده مفهوم زمینه این نوع ناهنجاری‌ها را القا می‌کند و فرمول‌بندی مساله به‌شدت نیازمند مشخصات آن است. دو ویژگی برای تعریف هر نمونه داده استفاده می‌شود، ویژگی‌های زمینه‌ای و رفتاری. ویژگی‌های متنی، زمینه را برای آن نمونه تعیین می‌کنند، مانند زمان در داده‌های سری زمانی. ویژگی‌های رفتاری ویژگی‌های غیر متنی، مانند مقدار خرید در مجموعه داده‌های کارت اعتباری را تعریف می‌کنند. رفتار غیرعادی نیز به نوبه خود با استفاده از ویژگی‌های رفتاری در یک زمینه خاص تعیین می‌شود. به‌عنوان مثال، یک ویژگی متنی در کلاهبرداری از کارت اعتباری، زمان خرید است. با فرض اینکه فردی یک صورتحساب خرج هفتگی صد دلاری داشته باشد، به‌جز هفته کریسمس که هزار دلار است. خرید هزار دلار در یک هفته متوسط در ماه مه یک ناهنجاری زمینه‌ای در نظر گرفته می‌شود، زیرا با رفتار معمولی در زمینه زمان مطابقت ندارد، حتی اگر خرج کردن همان مبلغ در هفته کریسمس طبیعی تلقی شود [10].

در نهایت، ناهنجاری‌های جمعی مجموعه‌ای از نمونه‌های مرتبط هستند که با توجه به کل مجموعه داده‌ها غیرعادی هستند. رویدادهای فردی یا نمونه‌های داده در یک ناهنجاری جمعی ممکن است لزوماً به‌خودی‌خود ناهنجاری نباشند. با این حال، وقوع آن‌ها به‌عنوان یک مجموعه غیرعادی است. توجه به این نکته مهم است که ناهنجاری‌های نقطه‌ای می‌توانند در هر مجموعه داده‌ای رخ دهند و ناهنجاری‌های جمعی تنها در مجموعه‌های داده‌ای که رابطه‌ای بین نمونه‌های داده وجود دارد، رخ می‌دهند. ویژگی‌های متنی تنها زمانی می‌توانند رخ دهند که ویژگی‌های متنی در داده‌ها وجود داشته باشد. ناهنجاری‌های نقطه‌ای یا جمعی حتی می‌توانند ناهنجاری‌های زمینه‌ای باشند اگر با توجه به یک زمینه مشاهده شوند [10].

۳-۲- هوش مصنوعی

هوش مصنوعی به ماشین‌ها اجازه می‌دهد تا توانایی‌های ذهن انسان را مدل‌سازی کنند و حتی آن‌ها را بهبود ببخشند [14]. از توسعه خودروهای خودران گرفته تا گسترش دستیارهای هوشمند مانند سیری و الکسا، هوش مصنوعی بخش رو به رشدی از زندگی روزمره است. در نتیجه، بسیاری از شرکت‌های فناوری در صنایع مختلف در حال سرمایه‌گذاری در فناوری‌های هوشمند مصنوعی هستند [15].

کم‌تر از یک دهه پس از کمک به نیروهای متفقین برای پیروزی در جنگ جهانی دوم با شکستن ماشین رمزگذاری نازی‌ها، آلن تورینگ ریاضیدان معروف با یک سوال ساده تاریخ را برای بار دوم تغییر داد: «آیا ماشین‌ها می‌توانند فکر کنند؟» مقاله تورینگ در سال ۱۹۵۰ «ماشین محاسباتی و هوش» و آزمون تورینگ متعاقب آن هدف و چشم‌انداز اساسی هوش مصنوعی را مشخص کرد [15]. هوش مصنوعی شاخه‌ای از علوم کامپیوتر است که هدف آن پاسخ مثبت به سوال تورینگ است. این تلاش برای تکرار یا شبیه‌سازی هوش انسانی در ماشین‌ها است. هدف گسترده هوش

مصنوعی سوالات و بحث‌های زیادی را به وجود آورده است. به حدی که هیچ تعریف واحدی از این رشته به‌طور کلی پذیرفته نشده است [15]. محدودیت اصلی در تعریف هوش مصنوعی به‌عنوان «ساخت ماشین‌هایی که هوشمند هستند» این است که در واقع توضیح نمی‌دهد که هوش مصنوعی چیست و چه چیزی یک ماشین را هوشمند می‌کند. هوش مصنوعی یک علم میان رشته‌ای با رویکردهای متعدد است، اما پیشرفت‌ها در یادگیری ماشینی و یادگیری عمیق تقریباً در هر بخش از صنعت فناوری یک تغییر پارادایم ایجاد می‌کند.

با این حال، اخیراً آزمایش‌های جدید مختلفی پیشنهاد شده‌اند که تا حد زیادی مورد استقبال قرار گرفته‌اند، از جمله یک مقاله تحقیقاتی در سال ۲۰۱۹ با عنوان «معیار اندازه‌گیری هوش». در این مقاله، فرانسوا شولت، محقق کهنه‌کار یادگیری عمیق و مهندس گوگل، استدلال می‌کند که هوش عبارت است از «میزانی که در آن یک یادگیرنده تجربیات و پیشینه‌های خود را به مهارت‌های جدید در کارهای ارزشمندی تبدیل می‌کند که شامل عدم قطعیت و سازگاری است»؛ به عبارت دیگر: هوشمندترین سیستم‌ها می‌توانند تنها مقدار کمی تجربه کنند و حدس بزنند که در بسیاری از موقعیت‌های مختلف نتیجه چه خواهد بود. با این حال، نویسندگان استوارت راسل و پیترو نوریگ در کتاب هوش مصنوعی که در بیش از ۱۴۰۰ دانشگاه به‌عنوان مرجع محسوب می‌شود می‌گویند هوش مصنوعی «مطالعه عامل‌هایی است که ادراکاتی را از محیط دریافت می‌کنند و متناسب آن اعمالی را انجام می‌دهند». در واقع تعریف هوش مصنوعی چهار نوع رویکرد زیر را شامل می‌شود [14].

۱. انسان اندیشی: تقلید از اندیشه مبتنی بر ذهن انسان است.

۲. عقلانی فکر کردن: تقلید فکر بر اساس استدلال منطقی است.

۳. رفتار انسانی: رفتار به‌گونه‌ای که رفتار انسان را تقلید کند.

۴. منطقی عمل کردن: عمل کردن به‌نحوی که برای رسیدن به هدفی خاص باشد.

تحقیقات اخیر نشان داده است که نوآوری هوش مصنوعی از قانون مور بهتر عمل کرده است و هر شش ماه یا بیشتر در مقایسه با دو سال دو برابر می‌شود. با این منطق، پیشرفت‌هایی که هوش مصنوعی در صنایع مختلف ایجاد کرده است، در چند سال گذشته بسیار مهم بوده است و پتانسیل تاثیر حتی بیشتر در چند دهه آینده کاملاً اجتناب‌ناپذیر به نظر می‌رسد [15].

۲-۴- یادگیری ماشین

یادگیری ماشینی، اصلی‌ترین شاخه هوش مصنوعی است که هوش مصنوعی را به یکی از جذاب‌ترین زمینه‌های تحقیقات علوم کامپیوتری تبدیل کرده [16] و شامل الگوریتم‌ها و تکنیک‌هایی است که به کامپیوترها امکان یادگیری را می‌دهد؛ به عبارت دیگر یادگیری ماشین حوزه‌هایی مانند داده‌کاوی، برنامه‌ریزی‌های دشوار و برنامه‌های نرم‌افزاری را پوشش می‌دهد و می‌تواند برای رگرسیون^۱ یا طبقه‌بندی^۲ چند متغیره، غیرخطی، غیرپارامتری مورد استفاده قرار گیرد. قابلیت‌های قابل توجه شبیه‌سازی روش‌های مبتنی بر یادگیری ماشین منجر به کاربردهای گسترده آن‌ها در علم و مهندسی شده است [17].

فرضیه اصلی یادگیری ماشین این است که یک ماشین (به‌عنوان مثال الگوریتم یا مدل) قادر به پیش‌بینی جدید بر اساس داده باشد. برخی از الگوریتم‌ها تحت نظارت قرار می‌گیرند (یادگیری با نظارت)، به این معنی که به آن‌ها داده‌های پیش‌بینی نشان داده می‌شود و سپس بر اساس داده‌های قبلی در مورد داده‌های جدید پیش‌بینی می‌کنند. برخی نیز بدون نظارت (یادگیری بدون نظارت) هستند، به این معنی که می‌توانند بدون داده‌های پیش‌بینی پیش‌بینی کنند. بعضی از آن‌ها نیز ترکیبی از این دو (نیمه نظارت^۳) هستند [17]. از یادگیری با نظارت برای کارها طبقه‌بندی و رگرسیون و از الگوریتم‌های یادگیری بدون نظارت برای کارهای خوشه‌بندی^۴ استفاده می‌شود [15]. البته یادگیری دیگری به نام یادگیری تقویتی^۵ نیز وجود دارند. الگوریتم‌های زیادی در یادگیری ماشین از قبیل روش بیز^۶ نزدیکترین k همسایه^۷، ماشین بردار پشتیبان، جنگل تصادفی^۸، شبکه‌های عصبی مصنوعی^۹، یادگیری عمیق^{۱۰} و غیره وجود دارد که شکل ۲ بخشی را نشان می‌دهد [18]. الگوریتم‌های اساسی در داده‌کاوی و یادگیری ماشینی،

¹ Regression

² Classification

³ Semisupervised learning

⁴ Clustering

⁵ Reinforcement learning

⁶ Naive bayes classifier

⁷ K-nearest neighbors

⁸ Random forest

⁹ Artificial neural network

¹⁰ Deep learning

اساس علم داده را تشکیل می‌دهند و از روش‌های خودکار برای تجزیه و تحلیل الگوها و مدل‌ها برای انواع داده‌ها در برنامه‌های کاربردی اعم از اکتشاف علمی تا تجزیه و تحلیل تجاری استفاده می‌کنند [16].

از آنجایی که هم داده‌کاوی و هم یادگیری ماشین از داده‌ها استفاده می‌کنند، جز علوم داده محسوب می‌شوند. هر دو تکنولوژی برای حل مسایل پیچیده استفاده می‌شوند؛ بنابراین، در نتیجه، بسیاری از افراد (به اشتباه) از این دو اصطلاح به جای یکدیگر استفاده می‌کنند. در داده‌کاوی و هم یادگیری ماشین از الگوریتم‌های یکسانی برای کشف الگوهای داده استفاده می‌کنند [15]. برای ارزیابی مدل‌های یادگیری ماشین یا داده‌کاوی از ماتریس سردرگمی^۱ استفاده می‌شود که در جدول ۱ قابل مشاهده می‌باشد [19]. در این جدول TP میزان داده‌هایی است که به درستی در وضعیت کلاس درست تشخیص داده شده‌اند و TN میزان داده‌هایی است که به درستی در وضعیت کلاس منفی تشخیص داده شده‌اند، می‌باشد و FP و FN میزان خطاهای کلاس درست و کلاس منفی می‌باشد.

جدول ۱- ماتریس سردرگمی [19].

Table 1- Confusion matrix [19].

	پیش‌بینی مثبت	پیش‌بینی منفی
وضعیت مثبت	TP	FN
وضعیت منفی	FP	TN



شکل ۲- الگوریتم‌های یادگیری ماشین [17].

Figure 2- Machine learning algorithms [17].

با استفاده از جدول فوق یکسری پارامترهای ارزیابی تحت عنوان دقت (پیش‌بینی مثبت)^۲، دقت^۳، یادآوری^۴، حساسیت^۵، معیار F ^۱ و سطح زیر منحنی ROC ^۶ (نمودار نرخ مثبت واقعی در مقابل نرخ مثبت کاذب است) به شرح ذیل وجود دارد:

دقت: نسبت پیش‌بینی‌های صحیح به کل پیش‌بینی‌ها می‌باشد که از رابطه زیر محاسبه می‌شود [17]:

$$Accuracy = \frac{TP + TN}{TP + TN + FN + FP}$$

پیش‌بینی مثبت (دقت): نسبت موارد مثبت پیش‌بینی‌شده که درست تشخیص داده شده است که به صورت زیر تعریف می‌شود [20]:

$$Precision = \frac{TP}{TP + FP}$$

¹ Confusion matrix

² Accuracy

³ Precision (positive predictive)

⁴ Sensitivity (recall) hit rate

⁵ Specificity

⁶ F-measure

⁷ ROC curve (AUC)

میزان ضربه حساسیت (یادآوری): این مقدار به صورت زیر تعریف می شود [20]:

$$Recall = \frac{TP}{P} = \frac{TP}{TP + FN}$$

حساسیت: نسبت منفی های واقعی به کل منفی ها است. مقدار به صورت زیر تعریف می شود [20]:

$$Specificity = \frac{TN}{N} = \frac{TN}{TN + FP}$$

معیار F که از رابطه زیر به دست می آید [20].

$$F - measure = 2 \frac{Precision \times Recall}{Precision + Recall}$$

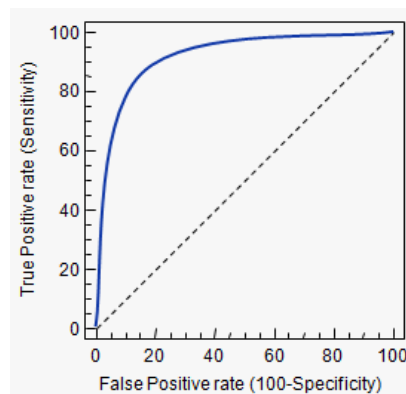
در روابط فوق:

TN : تعداد مقادیری است که دسته واقعی آن ها منفی بوده و الگوریتم دسته بندی نیز دسته آن ها را به درستی منفی تشخیص داده است.

TP : تعداد مقادیری است که دسته واقعی آن ها مثبت بوده و الگوریتم دسته بندی نیز دسته آن ها را به درستی مثبت تشخیص داده است.

FP : تعداد مقادیری است که دسته واقعی آن ها منفی بوده و الگوریتم دسته بندی دسته آن ها را به اشتباه مثبت تشخیص داده است.

FN : تعداد مقادیری است که دسته واقعی آن ها مثبت بوده و الگوریتم دسته بندی دسته آن ها را به اشتباه منفی تشخیص داده است [20].



شکل ۳- نمودار ROC.

Figure 3- ROC chart.

معیار ناحیه زیر منحنی^۱ است. AUC نشان دهنده سطح زیر نمودار ویژگی عملیاتی گیرنده^۲ می باشد و هر چه مقدار این عدد مربوط به یک دسته بند بزرگتر باشد کارایی نهایی دسته بند مطلوب تر ارزیابی می شود. نمودار ROC روشی برای بررسی کارایی دسته بندها می باشد. در واقع منحنی های ROC منحنی های دو بعدی هستند که در آن ها نرخ تشخیص صحیح دسته مثبت^۳ روی محور Y و نرخ تشخیص غلط دسته^۴ روی محور X رسم می شوند. به بیان دیگر یک منحنی ROC مصالحه نسبی میان سودها و هزینه ها را نشان می دهد. برای درک بهتر نمودار ROC شکل ۳ را مشاهده کنید.

¹ Area Under the Curve (AUC)

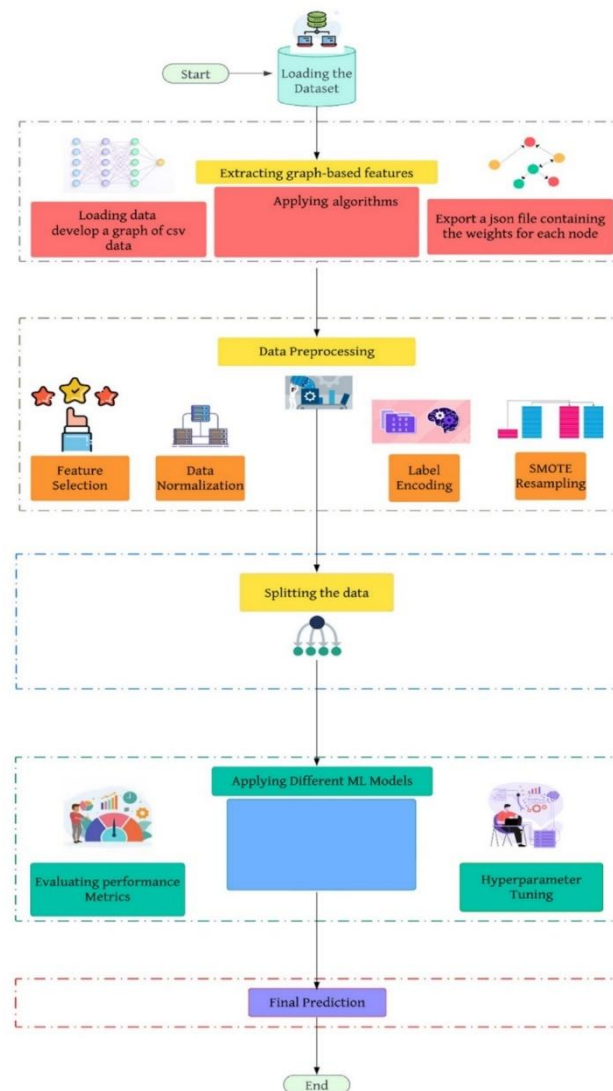
² Receiver Operating Characteristic (ROC)

³ True Positive Rate (TPR)

⁴ False Positive Rate (FPR)

۳- نتیجه‌گیری

با توجه به مطالب ارایه‌شده، با تحلیل داده‌های بانکی امکان تشخیص ناهنجاری‌های بانکی از جمله تراکنش‌های مشکوک وجود دارد، لذا تحلیل برخط داده‌های این سیستم می‌توان روش کارآمدی برای تشخیص تراکنش‌های بانکی ارایه نمود. شکل ۴ این الگوریتم را نشان می‌دهد.



شکل ۴- الگوریتم پیشنهادی.

Figure 4- Proposed algorithm.

مراجع

- [1] Karpoff, J. M. (2021). The future of financial fraud. *Journal of corporate finance*, 66, 101694. <https://doi.org/10.1016/j.jcorpfin.2020.101694>
- [2] Mandal, A., & Shanmugam, A. (2023). Fathoming fraud: unveiling theories, investigating pathways and combating fraud. *Journal of financial crime*, 31(5), 1106-1125. <http://dx.doi.org/10.1108/JFC-06-2023-0153>
- [3] Olufemi, B., Bello, O., Olufemi, K., & Author, C. (2024). Artificial intelligence in fraud prevention: Exploring techniques and applications challenges and opportunities. *Computer science & it research journa*, 5(6), 1505-1520. <http://dx.doi.org/10.51594/csitrj.v5i6.1252>
- [4] Hassan, M., Aziz, L. A. R., & Andriansyah, Y. (2023). The role artificial intelligence in modern banking: an exploration of AI-driven approaches for enhanced fraud prevention, risk management, and regulatory compliance. *Reviews of contemporary business analytics*, 6(1), 110-132. <https://researchberg.com/index.php/rcba/article/view/153>
- [5] Shoetan, P. O., & FAMILONI, B. T. (2024). Transforming fintech fraud detection with advanced artificial intelligence algorithms. *Finance & accounting research journal*, 6(4), 602-62. <https://doi.org/10.51594/farj.v6i4.1036>
- [6] Devan, M., Prakash, S., & Jangoan, S. (2023). Predictive maintenance in banking: Leveraging AI for real-time data analytics. *Journal of knowledge learning and science technology*, 2(2), 483-490. <https://doi.org/10.60087/jklst.vol2.n2.p490>

- [7] Shabbir, A., Shabir, M., Javed, A. R., Chakraborty, C., & Rizwan, M. (2022). Suspicious transaction detection in banking cyber-physical systems. *Computers & electrical engineering*, 97, 107596. <https://doi.org/10.1016/j.compeleceng.2021.107596>
- [8] Hawkins, D. M. (1980). *Identification of outliers*. Springer Netherlands. <https://books.google.com/books?id=fb0OAAAAQAAJ>
- [9] Adhikari, R., & Agrawal, R. K. (2013). *An introductory study on time series modeling and forecasting*. <https://doi.org/10.48550/arXiv.1302.6613>
- [10] Chandola, V., Banerjee, A., & Kumar, V. (2009). Anomaly detection: A survey. *ACM computing surveys (CSUR)*, 41(3), 1–58. <https://doi.org/10.1145/1541880.1541882>
- [11] Xiong, H., Pandey, G., Steinbach, M., & Kumar, V. (2006). Enhancing data analysis with noise removal. *IEEE transactions on knowledge and data engineering*, 18(3), 304–319. <https://doi.org/10.1109/TKDE.2006.46>
- [12] Hilal, W., Gadsden, S., & Yawney, J. (2021). Financial fraud: A review of anomaly detection techniques and recent advances. *Expert systems with applications*, 193(8), 116429. <http://dx.doi.org/10.1016/j.eswa.2021.116429>
- [13] Hodge, V., & Austin, J. (2004). A survey of outlier detection methodologies. *Artificial intelligence review*, 22, 85–126. <http://dx.doi.org/10.1023/B:AIRE.0000045502.10941.a9>
- [14] Russell, S. J., & Norvig, P. (2016). *Artificial intelligence: A modern approach*. Pearson. <https://thuvienso.hoasen.edu.vn/handle/123456789/8967>
- [15] Mirhoseini, S. R., Jahedmotlagh, M.-R., Ardakani, B., & Sabeti, M. (2022). *Survey the use of artificial intelligence and machine learning to predict and prevent crimes and violations of goods and currency smuggling from the perspective of data governance*. Hidden economy scientific and research. Kordestan, Iran. **(In Persian)**. <https://b2n.ir/ju5609>
- [16] Jin, Z., Yang, H., & Yin, G. (2021). A hybrid deep learning method for optimal insurance strategies: Algorithms and convergence analysis. *Insurance: Mathematics and economics*, 96, 262–275. <https://doi.org/10.1016/j.insmatheco.2020.11.012>
- [17] El Bouchefry, K., & de Souza, R. S. (2020). Learning in BIG Data: Introduction to machine learning. In *Knowledge discovery in big data from astronomy and earth observation* (pp. 225–249). Elsevier. <https://doi.org/10.1016/B978-0-12-819154-5.00023-0>
- [18] Mirhoseini, S. R., & Minaei, B. (2020). Network security: Artificial intelligence method for attack detection (Survey study). *3rd national conference on computer, information technology and applications of artificial intelligence*, Ahvaz, Iran. **(In Persian)**. <https://b2n.ir/nb4803>
- [19] Mirhoseini, S. R., JahedMotlagh, M. R., & Pooyan, M. (2016). *Improve accuracy of early detection sudden cardiac deaths (scd) using decision forest and SVM*. Proceedings of the international conference on robotics and artificial intelligence (pp. 20–22). Nagoya, Japan. <https://b2n.ir/xg8686>
- [20] Mirhoseini, S. R., Vahedi, F., & Nasiri, J. A. (2020). *E-mail phishing detection using natural language processing and machine learning techniques*. Iran, Tech. Rep. **(In Persian)**. <https://b2n.ir/xm4617>