



Anomaly Detection in Accounting Entries Using Deep Learning with Autoencoder Neural Networks

Mozaffar Jamalianpour

Assistant Professor of Accounting, Allameh Tabataba'i University, Tehran, Iran,
m.jamalianpour@atu.ac.ir

ARTICLE INFO	ABSTRACT
Received: 2025-04-12 Revised: 2025-09-19 Accepted: 2025-09-23	The detection of anomalies and fraudulent activities in accounting records has become increasingly critical in modern auditing practices, particularly in the era of big data where traditional sampling methods are insufficient. This study proposes a novel approach utilizing deep autoencoder neural networks for anomaly detection at the transaction level within accounting information systems. Two large-scale datasets were used: 36,538 journal entries from the Rahkaran system and 30,000 from the Sepidar system. Artificial anomalies were injected to evaluate performance. The autoencoder was trained in an unsupervised manner using PyTorch, with reconstruction error as the anomaly indicator. The empirical results indicate that the proposed model significantly outperforms conventional detection techniques, demonstrating a strong ability to identify both global anomalies (e.g., unusual amounts or transaction timings) and contextual anomalies (e.g., rare attribute combinations). Key features included subsidiary account, general ledger code, cost center, and last modification date. The findings provide strong evidence that deep learning-based anomaly detection can substantially improve fraud risk assessment and enhance the reliability of financial reporting, thereby offering a powerful tool for auditors, regulators, and financial system designers.
* Corresponding author: Dr. Mozaffar Jamalianpour Assistant Professor of Accounting, Allameh Tabataba'i University, Tehran, Iran Email: m.jamalianpour@atu.ac.ir	

1. Introduction

The primary purpose of this research is to address the critical challenge of detecting anomalies and fraudulent activities in accounting records

through advanced machine learning techniques. Financial reporting and auditing environments are increasingly characterized by high data volumes, complexity, and the necessity for timely detection of irregularities. Traditional audit methods—primarily based on sampling and red-flag indicators—are limited in scope and often incapable of identifying rare, sophisticated, or concealed fraudulent activities. These limitations create substantial audit risks, reduce the reliability of financial reporting, and undermine stakeholder confidence.

From a theoretical perspective, the study seeks to fill a gap in the accounting and auditing literature by applying the application of deep learning methods, specifically autoencoder neural networks, to journal-level transaction data. While anomaly detection has been widely studied in fields such as cybersecurity and network intrusion detection, its application in accounting information systems remains underexplored. This research responds to this gap by adapting deep learning methodologies to the unique characteristics of financial transactions, where anomalies may not be obvious single-point outliers but instead subtle deviations in complex attribute relationships.

From a practical perspective, the study aims to contribute to the advancement of auditing practice by offering auditors, regulators, and organizations a robust AI-driven tool for full-population testing. Unlike sample-based approaches, the proposed model enables continuous monitoring and detection of anomalies across all journal entries, thereby enhancing the timeliness and effectiveness of fraud detection and internal control systems. Moreover, the integration of anomaly detection into accounting information systems can support proactive risk management, reduce the likelihood of undetected fraudulent activity, and ultimately strengthen investor confidence and corporate governance.

Accordingly, the overarching purpose of this research is twofold: (i) to extend the academic understanding of anomaly detection in the accounting domain by applying deep autoencoder neural networks at the transaction level, and (ii) to provide actionable insights and technological solutions that can transform audit methodologies in line with the data-intensive demands of modern financial systems.

2. Methodology

This applied and quantitative study was conducted using two real datasets extracted from widely implemented Iranian accounting systems: (1) Rahkaran containing 36,538 journal entries, and (2) Sepidar containing 30,000 journal entries. Controlled artificial anomalies were injected into both datasets to systematically test the performance of the proposed model. A deep autoencoder network was designed and trained using PyTorch. The model minimized reconstruction error between input and output vectors, with anomalies identified when reconstruction errors exceeded a statistically determined threshold. The performance of the model was benchmarked against alternative approaches, including statistical outlier detection and clustering-based techniques, to ensure result robustness.

3. Results

The study yielded several significant findings:

The autoencoder demonstrated superior detection performance, outperforming traditional methods in detecting rare and complex patterns.

Increasing network depth enhanced detection capabilities, highlighting the need to model nonlinear relationships in accounting data.

Reconstruction error distributions varied across subsystems, emphasizing context-sensitive detection (e.g., sales, inventory, payroll).

Key features included subsidiary account, general ledger code, cost center, and last modification date.

4. Conclusion

This study provides theoretical and practical contributions to the accounting and auditing literature. Theoretically, it extends anomaly detection research by applying deep autoencoder networks directly to journal entries, rather than to aggregated financial statements. Practically, it offers auditors and regulators a data-driven, AI-enabled tool for moving from traditional sample-based procedures toward full-population testing. By integrating such models into audit workflows, stakeholders can significantly reduce the risk of undetected fraud, strengthen internal control systems, and enhance investor confidence in financial reporting. The findings confirm the potential of deep learning approaches to transform the auditing profession in an era of increasingly complex and

data-intensive financial systems.

Keywords: Anomaly detection, accounting information systems, autoencoder neural networks, deep learning, fraud detection.



تشخیص ناهنجاری‌ها در ثبت‌های حسابداری با استفاده از یادگیری عمیق

در شبکه‌های خود رمزگذار

دکتر مظفر جمالیان‌پور^۱

چکیده: کشف ناهنجاری و تقلب همواره یکی از چالش‌های مد نظر حساب‌رسان و ممیزین بوده است. در پژوهش حاضر سعی شده است پس از بیان مبانی نظری انواع ناهنجاری در سیستم دفترداری و ثبت اسناد مالی، با بهره‌گیری از الگوریتم شبکه‌های عصبی رمزگذار خودکار که یکی از روش‌های یادگیری عمیق بشمار می‌رود، ساز و کار برای استفاده از فنون نوظهور ارائه شود. بدین منظور از دو مجموعه داده سیستم اطلاعاتی حسابداری نرم افزار راهکاران (شامل ۳۶۵۳۸ آرتیکل) و سپیدار سیستم (شامل ۳۰۰۰۰ آرتیکل) استفاده شده است. در این پژوهش چند رویداد ناهنجار به مجموعه داده‌های موجود اضافه شد تا بدین وسیله معیار مطلوب برای قدرت تشخیص الگوریتم مشخص شود. نتایج نشان داد که در وهله اول یادگیری عمیق به وسیله الگوریتم شبکه عصبی خود رمزگذار توانایی بالایی در تشخیص آرتیکل‌های ناهنجار دارد و به علاوه به نسبت سایر فنون مورد بررسی (شامل درخت تصمیم، Extra Trees Regressor، جنگل تصادفی، Ada Boost Classifier و Quadratic Discriminant Analysis) عملکرد مطلوب‌تری را نشان می‌دهد. افزون بر این با افزایش عمق شبکه عملکرد تشخیصی بهبود یافت. به علاوه نتایج برآمده از پژوهش مبین آن بود که اسناد و آرتیکل‌های صادر شده از زیر سیستم‌های مختلف دارای خطای بازسازی متفاوتی بودند و لذا در تشخیص ناهنجاری باید به زیر سیستم پایه صادر کننده سند توجه کرد. به علاوه با اتکا به نتایج پژوهش حساب تفصیلی (طرف مقابل)، حساب معین، مرکز هزینه و تاریخ آخرین ویرایش سند، مهم‌ترین ویژگی‌ها برای تشخیص ناهنجاری در سیستم صدور سند هستند.

واژه‌های کلیدی: ناهنجاری داده، سیستم‌های اطلاعاتی حسابداری، شبکه عصبی خود رمزگذار، کشف ناهنجاری داده محور.

^۱ m.jamalianpour@atu.ac.ir

پذیرش: ۱۴۰۴/۰۷/۰۱

۱. استادیار حسابداری، دانشگاه علامه طباطبائی، تهران، ایران (نویسنده مسئول)

اصلاحات نهایی: ۱۴۰۴/۰۶/۲۸

دریافت: ۱۴۰۴/۰۱/۲۳

۱. مقدمه

علیرغم اینکه سیستم دفترداری یکی از اجزای اصلی حسابداری مدرن بشمار می‌آید و عملاً دروازه اصلی ورود داده‌های مالی و تجاری به سیستم‌های اطلاعاتی سازمانی است، لیکن به دلایل مختلفی از جمله عدم در دسترس بودن و محدودیت‌های محرمانگی به موضوع‌های مرتبط با کیفیت و کاربرد پذیری آنها بسیار کم پرداخته شده است. گیو و همکاران^۱ (۲۰۲۲)، با بررسی مجلات علمی حسابداری به این مهم اشاره داشته‌اند. دبریسنی و گری (۲۰۱۰)^۲ بیان می‌دارند که روش‌های سنتی حسابرسی عمدتاً بر بررسی نمونه‌ای از داده‌ها و استفاده از روش‌های تحلیلی سنتی مانند آزمون‌های آماری و کنترل داخلی متکی هستند. این روش‌ها معمولاً به تمامی داده‌های حسابداری دسترسی ندارند و ممکن است موارد تقلب پنهان و الگوهای غیرعادی را تشخیص ندهند.

انجمن بازرسان خبره تقلب^۳ در پژوهش تقلب جهانی سال ۲۰۲۴ خود تخمین می‌زند که سازمان‌ها و نهادها معمولاً حدود ۵ درصد از درآمد سالانه خود را به دلیل تقلب از دست می‌دهند. این پژوهش نشان داده است که سیستم اطلاعاتی حسابداری یکی از اصلی‌ترین ابزارهای است که مرتکبین تقلب برای کتمان و مخفی سازی تراکنش‌ها و رویدادهای تحریفی استفاده شده است. همچنین سیستم‌های نظارتی خودکار و تطبیق حساب‌ها سریع‌ترین ابزار برای کشف تحریف و تقلب بوده است. تقلب در گزارش‌ها و صورت‌های مالی به نسبت تقلب استفاده نادرست از دارایی‌ها و فساد بیشترین مبلغ نوع تقلب را به خود اختصاص داده است (انجمن بازرسان خبره تقلب، ۲۰۲۴). در واقع به دلیل آنکه ثبت‌های حسابداری در دفاتر مالی و حسابداری اولین مرحله از شناخت رویدادهای مالی و تجاری هستند، هر گونه خطا و تحریف در این بخش موجب خواهد شد تا خروجی سیستم اطلاعاتی حسابداری نیز نامطلوب و بعضاً گمراه کننده شود. از این‌رو موضوع تقلب و تحریف همواره مورد توجه سازمان‌های نظارتی و همچنین استفاده کنندگان از اطلاعات حسابداری بوده است.

سازمان‌ها همواره در حال استفاده از فناوری‌های اطلاعاتی در راستای دیجیتالی کردن و پیکربندی مجدد فرآیندهای کسب و کاری خود هستند که این موضوع بر سیستم‌های اطلاعات حسابداری به صورت خاص و به طور کلی‌تر بر سیستم‌های برنامه‌ریزی منابع سازمانی^۴ تأثیر می‌گذارد. این سیستم‌ها به طور پیوسته مقادیر زیادی از شواهد الکترونیکی و یا فیزیکی را در سطح جزئی و خرد (با بیشترین جزئیات) جمع‌آوری می‌کنند. این امر عمدتاً به وسیله ثبت داده‌های ورودی در دفتر روزنامه و کل صورت می‌پذیرد و در نهایت منجر به گردش حساب‌های معین خواهد شد. برای شناسایی فعالیت‌های بالقوه متقلبانه، استانداردهای حسابرسی، حساب‌رسان را به ارزیابی

مستقیم داده‌های ورودی به دفاتر حسابداری و همچنین سیستم‌های اطلاعاتی که خروجی آنها بر گزارش‌ها و صورت‌های مالی اثر می‌گذارد توصیه کرده است. استاندارد حسابرسی ۳۱۵ (تشخیص و ارزیابی خطرهای تحریف بااهمیت از طریق شناخت واحد تجاری و محیط آن)، ۳۳۰ (برخوردهای حسابرسان با خطرهای ارزیابی شده)، ۵۰۰ (شواهد حسابرسی)، ۵۲۰ (شواهد حسابرسی - روش‌های تحلیلی)، ۵۳۰ (نمونه‌گیری) و ۲۴۰ (مسئولیت حسابرسان در ارتباط با تقلب، در حسابرسی صورت‌های مالی) هر کدام فراخور موضوع به بیان ابعادی از کشف ناهنجاری (مسئولیت، پیامد، تعریف و چگونگی کشف) در سیستم حسابداری و مالی کارفرما اشاره دارند.

امروزه، اغلب فنون کاربردی برای بررسی ثبت‌های دفاتر به قوانینی اشاره دارد که توسط حسابداران خبره یا متخصصین کشف تقلب تعریف شده است و عمدتاً به صورت ذهنی و قضاوتی تدوین می‌شوند و اغلب به صورت سنتی و غیرسیستماتیک اجرا می‌شوند. این فنون، معمولاً بر اساس سناریوهای تقلب شناخته‌شده، اغلب تحت عنوان آزمون‌های پرچم قرمز^۵ (برای نمونه اسناد ثبت شده در زمان‌های خاص، تغییرات و اصلاح چندباره حساب بانکی، تعدیل حساب هزینه‌های معوق) یا تحلیل‌های آماری پیشرفته (مانند قانون بنفورد) نامیده می‌شوند. این شیوه از کشف تقلب از موارد تقلب شناخته شده قبلی فراتر نمی‌روند و بنابراین در کشف طرح‌های جدید تحریف و تقلب ناتوان هستند (سان و وازرالی^۶، ۲۰۱۷؛ باو و همکاران^۷، ۲۰۲۰؛ هوآنگ و همکاران^۸، ۲۰۲۲؛ شارما و لوکانان^۹، ۲۰۲۵). به علاوه مبتنی بر رسیدگی نمونه‌ای بودن باعث بروز ضعفی عمده در تعیین رویدادهای نادر اما پر ریسک است (بیک و سالامون^{۱۰}، ۱۹۸۵ و تتلبوام و رابینسون^{۱۱}، ۱۹۷۵). به همین صورت، چنین قوانینی به سرعت منسوخ می‌شوند چرا که کلاهبرداران و متقلبین به طور مداوم راه‌هایی برای دور زدن این روش‌ها پیدا می‌کنند. پیشرفت‌های اخیر مبتنی بر ابزارهای تحلیلگری داده محور مانند یادگیری عمیق^{۱۲}، محققین را قادر می‌سازد تا ویژگی‌ها و روندهای پیچیده غیرخطی را از داده‌های خام استخراج کنند و باعث پیشرفت‌های بسیاری در حوزه‌های مختلف شوند (لکان و همکاران، ۲۰۱۵^{۱۳}).

تحریف و تقلب عمدتاً بر پایه الگویی متفاوت با آنچه در حالت طبیعی رخ می‌دهد، اتفاق می‌افتد. لذا ثبت‌های حسابداری چنین موارد غیرعادی و ناهنجار تعداد بسیار محدودی را شامل می‌شوند. از این رو به منظور کشف این موارد خاص که عملاً از پیچیدگی بالایی برخوردار هستند باید از ابزارهای نوین و توانمند در تشخیص موارد نادر استفاده کرد. شبکه‌های عصبی خودرمزگذار^{۱۴} یکی از روش‌های یادگیری عمیق است که دارای چنین توانایی است. از این رو در پژوهش حاضر سعی خواهد شد بر

پایه داده‌ها و ثبت‌های واقعی در نرم افزارهای حسابداری، یک روش جدید برای شناسایی ثبت‌های غیرعادی حسابداری رایانه‌ای پیشنهاد شود. لذا برای این موضوع در وهله اول، میزان خطای بازسازی یک ثبت حسابداری بر پایه یک شبکه رمزگذار خودکار عمیق آموزش دیده محاسبه می‌شود و سپس با خصیصه‌ها و ویژگی‌های ثبت‌های حسابداری عادی مقایسه می‌شود تا بر اساس میزان اختلاف میان آنچه از اسناد واقعی یادگرفته شده است، موارد مشکوک به ناهنجاری مشخص شود. علاوه بر این، در صورتی که این مقادیر از آستانه امتیازدهی از پیش تعریف شده فراتر رود، امکان مظنون بودن ثبت به عنوان یک ثبت غیرعادی وجود خواهد داشت.

بدین منظور از مجموعه داده‌های واقعی مربوط به ثبت‌های دو سیستم اطلاعاتی حسابداری سپیدار و راهکاران سیستم استفاده شده است. به علاوه اثربخشی روش پیشنهادی با ارزیابی مقایسه‌ای در برابر الگوریتم‌های تشخیص ناهنجاری پیشرفته مورد بررسی قرار گرفته است. شایان ذکر است که عمده پژوهش‌های پیشین در حوزه کشف ناهنجاری در سطح صورت‌ها و گزارش‌های مالی بوده است و یکی از نوآوری‌های پژوهش حاضر ارائه ساز و کاری برای تشخیص ناهنجاری در سطح ریز اسناد سیستم حسابداری (سیستم دفتر کل) است.

در پژوهش حاضر ناهنجاری در اسناد حسابداری به عنوان موضوع محوری مورد نظر است. همان‌طور که در ادامه پژوهش آورده شده است، ناهنجاری یکی از مواردی است که می‌تواند باعث بروز تحریف (تقلب و یا اشتباه) در سیستم‌های اطلاعاتی حسابداری شود. لذا در این پژوهش ناهنجاری به معنای تحریف نبوده و تنها به عنوان یک عامل تقویت کننده ریسک تحریف مورد توجه بوده است. در ادامه ابتدا پیشینه و مبانی نظری مربوط به پژوهش ارائه و سپس تعریف دقیقی از ناهنجاری‌های امکان‌پذیر در سیستم اطلاعاتی حسابداری تشریح شده است. پس از آن مطالب مربوط به شیوه پیاده سازی شبکه عصبی خودرزم‌گذار تشریح شد. در پایان پس از بیان نتایج پژوهش، برونادهای مربوط مورد بحث و نتیجه گیری قرار گرفت.

۲. پیشینه و مبانی نظری

موضوع کشف تقلب و ناهنجاری‌های حسابداری هم توسط نهادهای متولی کشف تقلب و هم توسط پژوهشگران مورد مطالعه قرار گرفته است (ویی و همکاران^{۱۵}، ۲۰۲۴؛ امانی و فدالالا،^{۱۶} ۲۰۱۷؛ ولز،^{۱۷} ۲۰۱۷ و سینگلتون و سینگلتون،^{۱۸} ۲۰۱۰). کشف ناهنجاری در سیستم‌های حسابداری و مالی

عمدتاً بر نتیجه و خروجی نهایی، یعنی گزارش‌های مالی و صورت‌های مالی معطوف بوده است.^{۱۹} لیکن به دلیل عدم توجه به ساختار و ریز اسناد مالی و همچنین شرایط رویدادهای مالی، در این رویکرد تنها بخشی از ناهنجاری‌های بالقوه در سطح کلی و گزارش‌های مالی قابل شناسایی است. در واقع بخش عمده تحقیقات به حوزه بررسی ناهنجاری در مدیریت سود و اعداد گزارش شده در اعلامیه‌های سودآوری شرکت‌ها متمرکز بوده است (تئو و همکاران، ۱۹۹۸^{۲۰}؛ پینکوس و همکاران، ۲۰۰۷^{۲۱}؛ دراک و همکاران، ۲۰۱۴^{۲۲}؛ یانگ و ژنگ، ۲۰۱۷^{۲۳}؛ لاکانان و همکاران، ۲۰۱۹^{۲۴}؛ و چن و ویلیکو، ۲۰۲۳^{۲۵}) در بخش‌های بعدی، پژوهش‌های پیشینی که بر روی (۱) شناسایی فعالیت‌های متقلبانه در داده‌های سیستم‌های اطلاعاتی حسابداری و برنامه ریزی منابع سازمانی و (۲) تشخیص ناهنجاری‌ها با استفاده از شبکه‌های رمزگذار خودکار بوده است، آورده شده است.

۲.۱ تشخیص تقلب در داده‌های سیستم‌های اطلاعاتی حسابداری و برنامه ریزی منابع سازمانی

تجزیه و تحلیل رایانه‌ای حسابداری قضایی^{۲۶} (کشف تقلب و تحریف داده محور) پیرامون ثبت‌های حسابداری در سیستم‌های اطلاعاتی و برنامه ریزی منابع سازمانی با افزایش حجم داده‌های ثبت شده در سیستم‌های اطلاعاتی امکان‌پذیر شده است. بای و همکاران (۲۰۰۶)^{۲۷} از روش‌های ساده بیزین برای شناسایی حساب‌های دفتر کل مشکوک، با ارزیابی ویژگی‌های مشتق‌شده از ثبت ورودی‌های دفاتر برای اندازه‌گیری هرگونه فعالیت غیرمعمول حساب دفتر کل استفاده کردند. رویکرد آنها توسط مک گلوهون و همکاران (۲۰۰۶)^{۲۸} پیرامون استفاده از تجزیه و تحلیل پیوند (شبکه‌ای) برای شناسایی گروه حساب‌های پرخطر دفتر کل توسعه یافت.

خان و کورنی (۲۰۰۹)^{۲۹} و خان و همکاران (۲۰۱۰)^{۳۰} پروفایل تراکنش کاربران را بر اساس الگوی فعالیت کاربری مبتنی بر شیوه ثبت رویدادها در دفاتر سیستم مدیریت منابع سازمانی سپ^{۳۱} را ایجاد کردند تا بدین منظور رفتار مشکوک و نقض تفکیک وظایف کاربران را شناسایی کنند. به همین ترتیب اسلام و همکاران (۲۰۱۰)^{۳۲} از گزارش‌های حسابرسی سیستم سپ (که مبتنی بر فایل‌ها واقعه نگر و لاگ‌ها است) برای تشخیص سناریوهای تقلب شناخته شده، تبانی، کلاهبرداری و تطبیق آنها با سناریوهای تقلب بر پایه علائم خطر استفاده کردند.

دبریسنی و گری (۲۰۱۰)^{۳۳} مقادیر پولی ثبت رویدادها در دفاتر ۲۹ شرکت ایالات متحده را تجزیه و تحلیل کردند تا بر اساس قانون بنفورد، ترکیبات غیرعادی ارقام و همچنین الگوی زمانی

غیرمعمول مانند ثبت‌های غیر معمولی پایان سال مالی را شناسایی کنند.

جانس و همکاران (۲۰۱۱)^{۳۴} با استفاده از خوشه بندی تراکنش‌ها به طبقه بندی تک متغیره و چند متغیره تراکنش‌های سفارش خرید پرداختند. از این‌رو، معاملاتی که به طور قابل توجه از مرکز خوشه‌ها فاصله داشتند، غیرعادی قلمداد و برای بررسی دقیق توسط حسابرسان پیشنهاد شدند. این رویکرد با استفاده از روش فرآیند کاوی برای تشخیص بدهی‌های نادرست (تقلب در پرداخت‌ها)، توسعه و تقویت شد.

نانماچر و گومز (۲۰۲۱)^{۳۵}، به بررسی پژوهش‌هایی پرداختند که از روش‌های بدون ناظر برای تشخیص ناهنجاری در حسابرسی داخلی استفاده کرده بودند. آنها خلاصه‌ای از روش‌ها و فنون قابل استفاده در این حوزه را جمع‌بندی و ارائه کردند. آنها نشان دادند که در جریان تحول دیجیتال و افزایش حجم داده‌ها، حسابرسی ناگزیر به بهره‌گیری از روش‌های نوین است. به‌کارگیری قواعد بر روی داده‌ها اگرچه امکان آزمون کل جمعیت داده‌های حسابرسی را فراهم می‌کند، اما محدود به کشف خطاها و انحرافات پیش‌بینی‌شده توسط حسابرس است. در مقابل، رویکرد شناسایی ناهنجاری بدون نظارت توانایی شناسایی انحرافات فرآیندی جدید و الگوهای تقلب نوظهور را دارد. مرور نظام‌مند مطالعات پیشین آنها نشان داد که اغلب پژوهش‌ها تنها بر یک مجموعه داده خاص متمرکز بوده و به مسئله تعمیم‌پذیری، ادغام در فرایند حسابرسی و نحوه ارائه مناسب نتایج به حسابرسان نپرداخته‌اند. بنابراین، شکاف پژوهشی مهمی در زمینه کاربست فراگیر و عملی این رویکرد در حسابرسی وجود دارد.

گیو و همکاران (۲۰۲۲)، در پژوهشی یک روش نوآورانه مبتنی بر ساختارشناسی گرافی حسابداری^{۳۶} را برای تحلیل داده‌های ثبت‌های روزنامه و کشف تقلب در حسابرسی معرفی کردند. با تکیه بر نظریه تناسب شناختی^{۳۷} و نظریه گراف، این روش امکان تجسم روابط درونی و میانی ثبت‌های روزنامه را فراهم می‌کند. یافته‌های آنها نشان داد که این روش می‌تواند در شناسایی ضعف‌های کنترل داخلی و تراکنش‌های مشکوک موثر باشد. به علاوه آنها نشان دادند که این رویکرد باعث افزایش کارایی و دقت حسابرسی شده و یک ابزار عملی برای حسابرسان جهت تجزیه و تحلیل سیستماتیک داده‌های مالی و حسابداری است.

ویی و همکاران (۲۰۲۴)، در پژوهشی به بررسی کاربرد یادگیری بدون نظارت در کشف نقاط پرت در داده‌های دفتر کل^{۳۸} با هدف بهبود فرآیندهای حسابرسی پرداختند. آنها یک چارچوب چندسطحی^{۳۹} برای شناسایی تراکنش‌های غیرعادی در سه سطح تراکنش، حساب و ترکیب متغیرها

ارائه کردند. نتایج آزمایش بر روی مجموعه داده‌های واقعی و مصنوعی نشان داد که این روش می‌تواند ریسک‌های نامشخص را شناسایی کند. همچنین، ترکیب این فنون با روش‌های سنتی حسابرسی می‌تواند دقت تحلیل‌ها را افزایش دهد. به علاوه یافته‌ها نشان داد که یادگیری بدون نظارت ابزار مؤثری در شناسایی ریسک‌های پنهان در حسابرسی است.

با توجه به آنچه بیان شد، اغلب پژوهش‌ها یا (۱) از تاریخچه و سوابق حسابداری و دانش حسابداری قضایی در راستای تعیین انواع علائم خطر (پرچم‌های قرمز) و طرح‌های تقلب استفاده کرده‌اند و یا (۲) بر روی فنون سنتی یادگیری غیرعمیق متمرکز بودند. در این پژوهش سعی خواهد شد تا با استفاده از یادگیری عمیق به تشخیص سناریوهای ناشناخته برای کشف معاملات مشکوک اقدام شود.

۲.۲ تشخیص ناهنجاری با استفاده از شبکه‌های عصبی رمزگذار خودکار

امروزه، شبکه‌های رمزگذار خودکار به طور گسترده در طبقه‌بندی تصویر، ترجمه ماشینی و پردازش گفتار برای قابلیت‌های فشرده‌سازی داده‌های بدون نظارت خودکار استفاده می‌شوند. هاوکینز و همکاران^{۴۰} (۲۰۰۲) و ویلیامز و همکاران^{۴۱} (۲۰۰۲) اولین کسانی بودند که شبکه‌های عصبی بازگشت پذیر^{۴۲} را برای تشخیص ناهنجاری پیشنهاد کردند. شبکه‌های رمزگذار خودکار یکی از انواع پیشرفته شبکه‌های عصبی برگشت پذیر هستند. این نوع از شبکه‌ها برای شناسایی رکوردهای غیرعادی در حوزه‌های مختلف مورد قرار گرفته است.

شبکه‌های رمزگذار خودکار در حوزه تحلیل داده‌ها برای کشف تقلب نیز استفاده می‌شوند. در واقع با گسترش حجم داده‌های حسابداری در سیستم‌های برنامه ریزی منابع سازمانی، شناسایی خودکار ورودی‌های غیرعادی و ناهنجاری‌ها اهمیت یافته است. شبکه‌های عصبی خودرمزنگار به دلیل توانایی یادگیری نمایش فشرده داده‌ها و تشخیص نمونه‌های ناهمخوان، به عنوان یک روش کلیدی در این حوزه مطرح شده‌اند. این روش در یادگیری، بدون نظارت عمل می‌کنند و نیازی به برچسب‌گذاری دقیق داده‌های ناهنجار ندارند. به همین علت در مطالعات اخیر استفاده از آن‌ها برای شناسایی خطاها و تقلب‌های مالی مورد توجه قرار گرفته است (هرناندز و همکاران^{۴۳}، ۲۰۲۴؛ باکومنکو و ال‌راگال^{۴۴}، ۲۰۲۲). در واقع به دلیل توانایی بالای شبکه‌های خودرمزنگار در خلاصه سازی و کاهش ابعاد و سپس امکان بازیابی مجدد مشاهدات می‌توانند در موضوع‌های تکرار پذیر مناسب باشند. در بخش قبلی به برخی از پژوهش‌هایی که از این روش استفاده کرده بودند، اشاره شد. پژوهش حاضر

یکی از اولین پژوهش‌های است که سعی دارد با الهام گرفته از یادگیری عمیق، ثبت‌های غیرعادی در داده‌های حسابداری واقعی در پایگاه داده‌ای دو نرم افزار سپیدار و راهکاران در ایران را مورد ارزیابی قرار دهد.

با توجه به آنچه گفته شد، هم با نگاه اکتشافی (کشف رویدادهای مشکوک) و هم با نگاه اجتنابی (پیشگیری از وقوع و ثبت اسناد غیر معمول) سیستم کشف ناهنجاری هنگامی مطلوبیت بالایی دارد که در زمان صدور سند و یا پردازش اسناد، موارد مشکوک کشف و اقدام مقتضی صورت پذیرد. این در حالی است که رویدادهای متقلبانه و غیرعادی عمدتاً دارای رفتاری جدید و نوظهور هستند. لذا استفاده از روش‌هایی نظیر الگوریتم‌های یادگیری بدون ناظر^{۴۵} می‌تواند در این راستا مناسب باشد.

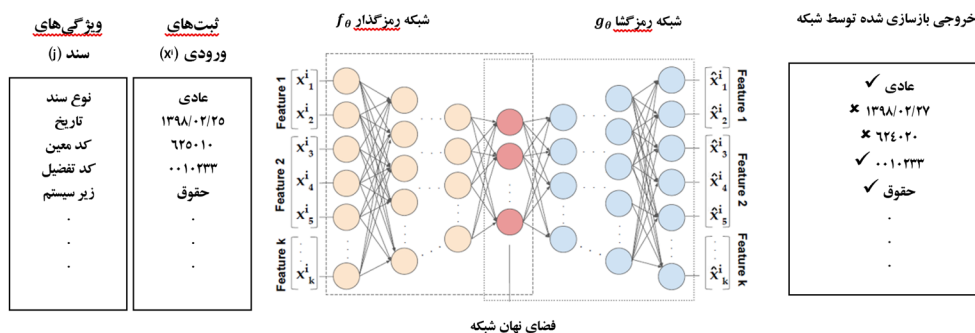
۳. تشخیص ناهنجاری‌های حسابداری

در این بخش ابتدا به معرفی عناصر اصلی شبکه‌های عصبی رمزگذار خودکار پرداخته شده است و سپس مفهوم ناهنجاری در سیستم صدور سند حسابداری تبیین می‌شود. همچنین چگونگی استفاده از خطای بازسازی چنین شبکه‌هایی برای شناسایی ثبت‌های غیرعادی در داده‌های حسابداری توضیح داده خواهد شد.

۳.۱ شبکه‌های عصبی خودکار رمزگذار عمیق

آموزش شبکه عصبی خود رمزگذار با استفاده از مجموعه ای از N ثبت دفاتر صورت می‌پذیرد ($X = \{x^1, x^2, x^3, \dots, x^N\}$) که در آن هر ثبت ورودی x^i از یک سری ویژگی^{۴۶} تشکیل شده است ($x^i = \{x_1^i, x_2^i, \dots, x_k^i\}$) لذا x_k^i نشان دهنده K امین ویژگی ثبت نام است. ویژگی‌های هر ثبت شامل جزئیات حسابداری آن ثبت خاص است، به عنوان مثال نوع سند، تاریخ، مبلغ، زیرسیستم صادر کننده، زمان صدور سند، ماهیت حساب، نوع حساب درگیر و سایر موارد مشابه نمونه ویژگی‌های یک ثبت حسابداری است. یک شبکه عصبی خود رمزگذار یا همانندساز^{۴۷}، نوع خاصی از شبکه عصبی چندلایه پیشرو^{۴۸} است که می‌تواند برای بازسازی ورودی بر اساس ویژگی‌های از قبل تعیین شده، آموزش ببیند. تفاوت بین ورودی اصلی (آنچه توسط کاربر ثبت شده است) و بازسازی شده (آنچه

شبکه عصبی حدس زده است) به عنوان خطای بازسازی^{۴۹} نامیده می‌شود. نگاره ۱ نمای کلی از یک شبکه عصبی خود رمزگذار را نشان می‌دهد.



نگاره ۱: نمای کلی شبکه عصبی خود رمزگذار

همان‌طور که در نگاره فوق مشهود است، ابتدا بر پایه ویژگی‌های سند و مقادیر مربوط به ثبت‌های ورودی داده‌های مربوط به مشاهدات وارد شبکه رمزگذار f و بر پایه فضای نهان شبکه عصبی نگاشت کلی از موضوع ایجاد می‌کند و سپس در مرحله بعد شبکه رمزگشای g سعی بر بازسازی مجدد ورودی‌ها بر پایه مدل ایجاد شده دارد. در نهایت خروجی بازسازی شده توسط شبکه مقادیر را بازسازی کرده تا بر پایه آن استنتاج لازم صورت پذیرد.

به طور کلی، شبکه‌های رمزگذار خودکار از دو نگاشت غیرخطی تشکیل شده‌اند که به آن‌ها شبکه رمزگذار f_θ شبکه رمزگشا g_θ گفته می‌شود (رامپهارت و همکاران^{۵۰}، ۱۹۸۶). معمولاً رمزگذار و رمزگشا دارای معماری متقارن هستند که از چندین لایه نورون تشکیل شده است و هر کدام یک تابع غیرخطی و پارامترهای مشترک را در درون خود قرار داده‌اند. نگاشت رمزگذار $f_\theta(\cdot)$ یک بردار ورودی x_i (ثبت نام) را به یک نمایش فشرده z_i در فضای نهان Z نگاشت می‌کند (تصویری از ثبت را بازسازی می‌کند). این نمایش نهان z_i سپس توسط رمزگشا $g_\theta(\cdot)$ به بردار بازسازی شده $\hat{x}_i$ از فضای ورودی اصلی نگاشت می‌شود. نگاشت رمزگذار و رمزگشای غیرخطی در یک شبکه رمزگذار خودکار که دارای چندین لایه نورون است را می‌تواند به صورت زیر تعریف کرد:

$$F_\theta^l(\cdot) = \sigma^l \left(W^l \left(F_\theta^{l-1}(\cdot) \right) + b^l \right), \text{ and } g_\theta^l(\cdot) = \sigma'^l \left(W'^l \left(g_\theta^{l-1}(\cdot) \right) + d^l \right)$$

معادله (۱)

جایی که σ و σ' نشان دهنده فعال‌سازهای غیرخطی (مانند تابع سیگموئید^{۵۱}) هستند، θ نشان

دهنده پارامترهای مدل است و $W \in \mathbb{R}^{dx*dz}$, $W' \in \mathbb{R}^{dz*dy}$, $\{W, b, W', d\}$ ماتریس‌های وزن هستند و بر همین منوال $d \in \mathbb{R}^{dy}$, $b \in \mathbb{R}^{dz}$ بردارهای بایاس انتقال^{۵۲} هستند و ۱ تعداد لایه‌های پنهان را نشان می‌دهد.

برای دستیابی به $x^i \approx \hat{x}^i$ و ایجاد بالاترین درجه انطباق، شبکه رمزگذار خودکار آموزش می‌بیند تا مجموعه‌ای از پارامترهای مدل رمزگذار-رمزگشای بهینه θ^* را به نحوی محاسبه کند که عدم تشابه یک ثبت حسابداری x^i و سند بازسازی شده آن $\hat{x}^i = g_\theta(f_\theta(x^i))$ به حداقل برسد. از این‌رو، هدف آموزش شبکه رمزگذار خودکار، یادگیری الگویی است که معادله زیر را بهینه‌سازی کند:

$$\arg \min_{\theta} \|X - g_\theta(f_\theta(X))\| \quad \text{معادله (۲)}$$

در فرآیند آموزش شبکه سعی خواهد شد برای تمام ثبت‌های X که به عنوان داده‌های آموزش برگزیده شده اند مقدار تابع هزینه^{۵۳} $\mathcal{L}_\theta$ به حداقل برسد. به بیان ریاضی سعی خواهد شد تا مقدار تابع زیر حداقل شود:

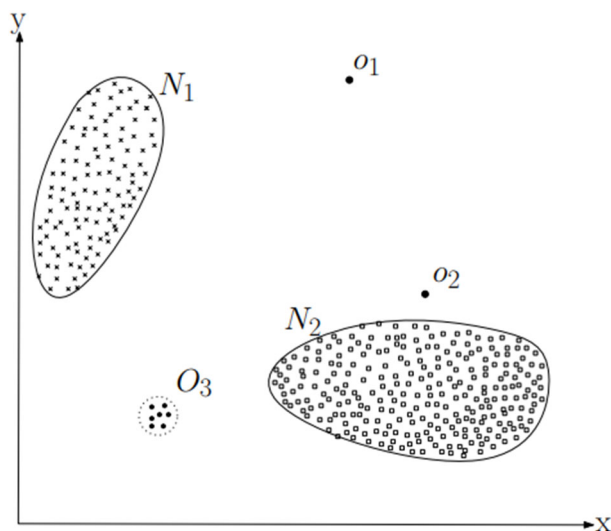
$$\mathcal{L}_\theta(x^i : \hat{x}^i) = \frac{1}{n} \sum_{i=1}^n \sum_{j=1}^k x_j^i \ln(\hat{x}_j^i) + (1 + x_j^i) \ln(1 - \hat{x}_j^i) \quad \text{معادله (۳)}$$

پس از اتمام یادگیری برای مجموعه‌ای از ورودی‌های دفاتر (داده‌های آزمون) و بازسازی مربوط که بر اساس ویژگی‌های ثبت‌های ورودی به دفاتر تخمین زده شده است، بررسی و آزمون‌های نهایی صورت می‌پذیرد تا با استفاده از مقادیر ویژگی کدگذاری شده، انحراف توزیع مستقل برنولی^{۵۴} چند متغیره محاسبه شود (بینگیو و همکاران^{۵۵}، ۲۰۱۳).

برای جلوگیری از بیش انطباقی^{۵۶}، تعداد نوروهای لایه‌های پنهان شبکه باید به حداقل ممکن کاهش یابد. به بیان ریاضی باید $R^{dx} > \in \mathbb{R}^{dz}$ باشد^{۵۷}. تحمیل چنین محدودیتی بر لایه پنهان شبکه، الگوریتم رمزگذار خودکار را مجبور می‌کند تا مجموعه بهینه‌ای از پارامترهای θ^* را بیاموزد و در نهایت منجر به یک مدل فشرده از رایج‌ترین توزیع‌های مقادیر ویژگی ثبت‌های دفاتر شود. به منظور پردازش و ارزیابی شبکه‌های عصبی مورد استفاده در این پژوهش از کتابخانه منبع باز PyTorch در بستر زبان برنامه‌نویسی پایتون استفاده شده است. با توجه به اینکه یادگیری ماشین به یک موضوع پر بسامد مبدل شده است، در زبان‌های مختلف برنامه‌نویسی از جمله پایتون بدین منظور افزونه‌های متعددی تعبیه شده است. کتابخانه PyTorch یکی از مرسوم‌ترین ابزارهایی است که با تخصیص بهینه منابع رایانه‌ای امکان ایجاد شبکه‌های عصبی و یادگیری عمیق را فراهم می‌سازد (فانگ و وانگ^{۵۸}، ۲۰۲۴).

۳.۲ طبقه بندی ناهنجاری‌های حسابداری

به طور کلی به منظور تعیین هنجار و تبیین ناهنجاری روش‌ها و الگوهای مختلفی ارائه شده است. کاندولا و همکاران (۲۰۰۹)^{۵۹} بدین منظور سه دسته ناهنجاری را تعریف می‌کنند. اولین دسته، ناهنجاری نقطه‌ای^{۶۰} است که در آن برخی از داده‌ها بنابر رفتار کلی داده‌ها در وضعیت غیر عادی قرار می‌گیرند. در نگاره ۲ نقاط o_1 و o_2 از این دسته ناهنجاری بشمار می‌آیند. این دسته از داده‌ها ساده‌ترین ناهنجاری‌ها را نگاره می‌دهند که عمدتاً به وسیله روش‌های سنتی مانند فنون آماری و پراکندگی قابل شناسایی هستند.



نگاره ۱: نمونه‌ای از انواع ناهنجاری در مجموعه داده‌ها (اقتباس از کاندولا و همکاران، ۲۰۰۹)

ناهنجاری‌های زمینه‌ای^{۶۱} یا ناهنجاری شرطی^{۶۲} دسته دیگری از ناهنجاری‌ها بشمار می‌آیند. این دسته از ناهنجاری‌ها به استناد ساختار کلی داده‌ها قابل تشخیص و تعریف هستند. لذا به منظور تشخیص ناهنجاری زمینه‌ای می‌توان به دو موضوع ناهنجاری خصیصه‌های زمینه‌ای^{۶۳} و ناهنجاری‌های خصیصه‌های رفتاری^{۶۴} داده‌ها توجه کرد. صدور سند حسابداری حقوق و دستمزد در اواسط ماه در شرکتی که عمدتاً حقوق دستمزد در اوایل هر ماه ثبت می‌شود، می‌تواند یک ناهنجاری رفتاری بشمار آید و یا اینکه صدور و تایید یک سند توسط شخصی که عمدتاً اسناد را صادر نمی‌کند، می‌تواند یک ناهنجاری خصیصه‌ای بشمار آید. اصلی‌ترین موضوع برای کشف این نوع از ناهنجاری‌ها

داشتن سبد مناسبی از ویژگی‌ها و خصیصه‌ها است. این ویژگی‌ها ممکن است از ذاتی و مبتنی بر واقعیت باشند و یا ساختگی و بر اساس الگوبندی مبتنی بر داده‌ها ایجاد شوند.

ناهنجاری‌های جمعی^{۶۵} نوع دیگری از ناهنجاری است که می‌توان در حوزه کشف ناهنجاری داده محور دنبال کرد. این نوع از ناهنجاری‌ها در بستر رفتار جمعی و عمومی داده‌های یک حوزه مشخص قابل تعریف و تبیین هستند. در واقع در این نوع از ناهنجاری رفتار داده‌ها بر اساس گروه‌بندی مشخص شده تعیین کننده عادی و غیرعادی بودن داده‌ها است.

برای شناسایی ثبت‌های غیرعادی ابتدا باید عادی بودن را با توجه به داده‌های حسابداری تعریف کرد. لذا فرض می‌شود که اغلب ثبت‌ها در سیستم‌های اطلاعاتی حسابداری یک شرکت مربوط به فعالیت‌های تجاری عادی و روزمره است. لذا برای انجام کلاهبرداری و تحریف، مرتکبین باید از روال عادی و معمول منحرف شوند. چنین رفتار انحرافی و متقلبانه بسیار محدود و نادر هستند. لذا ثبت‌های ورودی انحرافی و غیر معمول خواهند بود که هنجارهای حسابداری (متداول‌ترین رفتارهای موجود) و ویژگی‌های عمومی یک ثبت را رعایت نکرده باشد.

هنگام بررسی دقیق ثبت‌های حسابداری که عمدتاً در سیستم‌های برنامه ریزی منابع سازمانی و سیستم‌های اطلاعاتی حسابداری ثبت می‌شوند، دو ویژگی رایج را می‌توان مشاهده کرد: اول، ویژگی‌ها و خصیصه‌های ثبت‌های ورودی دفاتر تنوع بالایی دارند و دوم، وابستگی شدیدی بین حساب‌های کل و معین وجود دارد. به علاوه مقادیر و ویژگی‌های مشخصی در اسناد مربوط به یک حساب دفتر کل خاص وجود دارد. لذا با توجه به این موضوع و مشابه دسته بندی انحرافات در طبقه بندی برونینگ و همکاران (۲۰۰۰)^{۶۶}، می‌توان در ثبت‌های حسابداری دو نوع ناهنجاری را در نظر گرفت: ناهنجاری عمومی^{۶۷} (در سطح کلیت اسناد) و ناهنجاری محلی^{۶۸} (در سطح درونی و اختصاصی اسناد).

ناهنجاری‌های عمومی حسابداری، به ثبت‌هایی اشاره دارد که مقادیر غیرمعمول یا نادر ویژگی‌های یک سند را نشان می‌دهند. چنین ناهنجاری‌هایی معمولاً مربوط به ویژگی‌هایی است که در یک سند به نسبت سایر اسناد غیر معمول باشند. برای مثال گردش یافتن حساب‌های نادر در دفتر کل و یا ثبت‌هایی که در زمان‌های غیر معمول زده شده اند. آزمون‌های مرسوم و متداول که توسط حسابرسان در طول حسابرسی سالانه انجام می‌شود عموماً برای کشف این نوع از ناهنجاری‌ها طراحی شده‌اند. با این حال، چنین آزمون‌هایی اغلب به دلیل تعدد ثبت‌های اصلاحی و معکوس، ثبت‌های مبتنی بر مقررات و تعدیل‌های پایان سال که معمولاً دارای خطر تقلب پایین هستند، منجر

به حجم بالایی از هشدارهای مثبت کاذب^{۶۹} می‌شود. لذا ناهنجاری‌های عمومی اغلب منجر به کشف خطا می‌شوند و کمتر تقلبی را نشان می‌دهند.

ناهنجاری‌های محلی حسابداری، ثبت‌هایی هستند که ترکیبی غیرعادی یا نادری از مقادیر ویژگی‌ها در یک سند نسبت آنچه قبلاً در موارد مشابه اتفاق می‌افتد. برای مثال ردیف‌ها و آرتیکل‌های غیر معمول در یک سند، ترکیب نامنظم و غیر مرسوم در حساب‌های کل از این دسته ناهنجاری‌ها به شمار می‌آیند. تشخیص این نوع ناهنجاری به نسبت ناهنجاری عمومی دشوارتر است، زیرا متخلفین قصد دارند با تقلید از یک الگوی فعالیت منظم، رفتار متقلبانانه خود را پنهان کنند. در نتیجه، وجود چنین ناهنجاری‌هایی، مبین خطر تقلب بالایی هستند.

در حسابرسی‌های مرسوم، حسابرسان مایلند موارد ناهنجاری از هر دو دسته ناهنجاری بیان شده را کشف و برای بررسی دقیق‌تر آنها اقدام کنند. لذا در این پژوهش سعی خواهد شد تا الگویی ارائه شود که بر اساس آن ابتدا مقادیر و ویژگی‌های غیرمعمول اسناد (ناهنجاری‌های عمومی) و سپس ترکیب غیرعادی از مقادیر و ویژگی‌های مشاهده شده (ناهنجاری‌های محلی) بررسی می‌شود.

۳.۳ امتیاز دهی به ناهنجاری‌های حسابداری

یکی از خروجی‌های پژوهش حاضر، ارائه الگویی برای امتیازدهی و تشخیص ناهنجاری‌های عمومی و محلی در مجموعه داده‌های حسابداری دنیای واقعی است. امتیاز برای هر دو نوع ناهنجاری مشاهده شده، یعنی (۱) هر گونه وقوع مقدار غیر معمول در ویژگی‌ها (ناهنجاری عمومی) و (۲) هر گونه مقدار غیر معمول در ترکیب ویژگی‌ها (ناهنجاری محلی) به صورت زیر محاسبه و برآورد خواهد شد:

بررسی احتمال وقوع ویژگی‌ها: به منظور بررسی غیرمعمول یا نادر بودن ویژگی‌ها (X_j) ، احتمال وقوع آن در میان تمامی ثبت‌های ورودی تعیین می‌شود. بدین منظور می‌توان از معادله $\frac{n_j^i}{N}$ استفاده نمود که در آن n_j^i به معنای تعداد ویژگی‌های j ام در سند i ام است و N مجموع تعداد کل اسناد موجود در مجموعه داده‌ها است. سپس مجموع مقدار احتمال لگاریتمی مشخصه و ویژگی‌های (j) یک سند (n^i) به صورت زیر محاسبه می‌شود:

$$P(x^i) = \sum_{j=1}^k \ln \left(1 + \frac{n_j^i}{N} \right) \quad \text{معادله (۴)}$$

در نهایت، به منظور نرمال‌سازی مقادیر از روش بیشینه - کمینه به صورت زیر برای هر یک

از اسناد استفاده می‌شود:

$$AP(x^i) = \frac{P(x^i) - p_{min}}{p_{max} - p_{min}} \quad \text{معادله (۵)}$$

که در این معادله مقدار بیشینه و کمینه مربوط به مقادیر در کل اسناد موجود است.

بررسی ترکیب (همزمانی) مقدار ویژگی‌ها: سطح همزمانی مقادیر ویژگی‌ها و تعیین میزان ناهنجاری‌های محلی آنها، بر اساس خطای بازسازی یک سند که با آموزش یک شبکه عصبی خود رمزگذار به دست آمده است، برآورد می‌شود. به عنوان مثال، احتمال مشاهده یک حساب معین در ترکیب با یک حساب معین دیگر و یا سایر ویژگی‌ها برآورد و تخمین زده می‌شود. اگر این ترکیب غیرعادی باشد توسط شبکه در سطح پایینی بازسازی خواهد شد و لذا چنین ثبتي دارای خطای بازسازی بالایی خواهد بود. به منظور تعیین سطح خطای بازسازی شبکه رمزگذار خودکار (E) به صورت زیر محاسبه و برآورد خواهد شد.

$$E_{\theta^*}(x^i; \hat{x}^i) = \frac{1}{k} \sum_{j=1}^k (x_j^i - \hat{x}_j^i)^2 \quad \text{معادله (۶)}$$

که در آن $\hat{x}_j^i$ مقدار بازسازی شده x_j^i بر اساس مقادیر بهینه شده پارامترها (θ^*) است. سپس به منظور نرمال‌سازی همانند احتمال وقوع ویژگی‌ها از روش بیشینه - کمینه به صورت زیر استفاده خواهد شد:

$$RE_{\theta^*}(x^i; \hat{x}^i) = \frac{E_{\theta^*}(x^i; \hat{x}^i) - E_{\theta^*, min}}{E_{\theta^*, max} - E_{\theta^*, min}} \quad \text{معادله (۷)}$$

امتیازدهی ناهنجاری حسابداری: با محاسبه دو مقدار قبل، می‌توان غیرعادی بودن یک سند را از منظر عمومی و محلی محاسبه کرد. لذا به منظور غیر عادی خواندن یک سند حسابداری باید مقدار و امتیاز ناهنجاری سند به صورت زیر محاسبه و برآورد شود:

$$AS(x^i; \hat{x}^i) = \alpha \times RE_{\theta^*}(x^i; \hat{x}^i) + (1 - \alpha) \times AP(x^i) \quad \text{معادله (۸)}$$

در معادله فوق α ضریب ایجاد تعادل میان دو نوع ناهنجاری ممکن در اسناد است که بر اساس شرایط کلی محاسبه و در نظر گرفته می‌شود. پس از محاسبه مقدار و امتیاز نهایی ناهنجاری برای تمامی اسناد لازم است تا مقدار آستانه‌ای تشخیص ناهنجار بودن اسناد حسابداری محاسبه شود. این مقدار که با β نمایش داده می‌شود مبین سطحی است که مازاد بر آن، سند مورد نظر ناهنجار قلمداد خواهد شد. هر چه قدر مقدار پایین‌تری برای β در نظر گرفته شود (بازه تغییرات این مقدار ۰ تا یک است) ریسک عدم کشف پایین‌تری برنامه‌ریزی می‌شود و به تبع سطح دقت بیشتری باید اعمال شود (اسچریور و همکاران، ۲۰۱۷، ۲۰۱۹).

۴. پیاده سازی و استفاده از شبکه

در این بخش شیوه آموزش تجربی مدل شرح داده می‌شود. بدین منظور عملکرد تشخیص ناهنجاری در ۹ معماری مختلف شبکه رمزگذار خودکار مجزا و بر پایه دو مجموعه داده واقعی از ثبت‌های ورودی‌های زیر سیستم دفتر کل حسابداری در دو سیستم مجزای حسابداری (نرم افزار راهکاران و نرم افزار سپیدار) مورد ارزیابی قرار گرفت.

۴.۱ مجموعه داده‌ها و آماده سازی داده‌ها

با توجه به اینکه داده‌های سیستم‌های اطلاعاتی به عنوان یکی از منابع حیاتی و مهم سازمان‌ها بشمار می‌آید، دسترسی به آنها به منظور پژوهش امری دشوار است. در این پژوهش از دو مجموعه داده در دسترس محقق استفاده شده است.^{۷۱} برای بررسی مجموعه داده A از زیر سیستم دفتر کل نرم افزار راهکاران و جهت مجموعه داده‌ی B از زیر سیستم حسابداری سپیدار استفاده شده است. این اسناد مربوط به رویدادهای حین سال مالی است و لذا تمامی اسناد به جز سند افتتاحیه و اختتامیه (شامل اسناد تعدیلی مربوط) مد نظر قرار گرفته شده است. به منظور حفظ حریم خصوصی داده‌ها، تمام ویژگی‌های هویتی و متنی ثبت‌های ورودی با استفاده از یک تابع هش یک طرفه برگشت ناپذیر در طول فرآیند استخراج داده، ناشناس و گمنام شد. لازم به توضیح است که ساختار کلی اسناد حسابداری در سیستم‌های پردازش تراکنشی سازمان‌ها، فارغ از نرم افزار مورد استفاده، به هم مشابه است و عملاً تفاوت معناداری در ساختار و اجزای اسناد حسابداری در سیستم‌های مختلف وجود ندارد.

اغلب ویژگی‌های مربوط به اسناد ثبت شده در سیستم‌های حسابداری شامل متغیرهای طبقه‌بندی شده (گسسته) است، به عنوان مثال تاریخ ثبت، عنوان حساب، نوع سند و امثالهم همگی حالتی گسسته دارند. به منظور آموزش شبکه‌های عصبی رمزگذار خودکار، این ویژگی‌های باید به شکل باینری (کدگذاری شده بر پایه دو مقدار ۰ و یک) پیش پردازش شوند. در مجموع این پیش پردازش منجر به ۳۶۱۲ بعد برای مجموعه داده A و ۲۱۸۵ بعد رمزگذاری شده برای مجموعه داده B شد.

به منظور تجزیه و تحلیل دقیق و ارزیابی کمی آزمایش‌ها، تعدادی از آرتیکل‌های موجود به ناهنجاری‌های عمومی (۱۱ مورد برای مجموعه داده A و ۱۰ مورد برای مجموعه B) و تعدادی به

ناهنجاری محلی (۱۰ مورد برای مجموعه داده A و ۱۰ مورد برای مجموعه داده B) تبدیل شد^{۲۲}. این موضوع باعث شد تا مشابه یک حسابرسی حقیقی توزیع بسیار نامتعادل از ثبت‌های غیرعادی در مقابل ثبت‌های معمولی طراحی شود. ناهنجاری‌های عمومی از مقادیر ویژگی‌های تشکیل شده‌اند که در داده‌های اصلی معمول و عادی نیستند، در حالی که ناهنجاری‌های محلی ترکیبی از زیر مجموعه‌های از مقادیر ویژگی‌ها را نشان می‌دهند که در داده‌های اصلی رخ نمی‌دهند.

به بیان دیگر مجموعه داده A در مجموع شامل ۳۶.۵۳۸ خط آرتیکل ثبت حسابداری است که از ۱۳ ویژگی تشکیل شده است (شماره سند، شماره روزانه سند، تاریخ سند، صادر کننده، تاریخ ثبت، تاریخ تایید، کد حساب، مقدار ریالی آرتیکل، ماهیت آرتیکل {بدهکار یا بستانکار}، نوع سند {عمومی، حقوق و دستمزد، خزانه داری، انبار و اموال}، طرف حساب و مرکز هزینه). در مجموع برای این مجموعه داده ۲۱ (۰.۰۵۷٪) آرتیکل دستکاری شده مصنوعی به مجموعه داده‌ها اضافه شده است. این ثبت‌ها شامل ۱۰ (۰.۰۲۷٪) ناهنجاری عمومی و ۱۱ (۰.۰۳۰٪) ناهنجاری محلی است. نگاره شماره ۲ تعداد هر یک از ویژگی‌های موجود در مجموعه داده اول را نشان می‌دهد. این داده‌ها مربوط به نرم افزار راهکاران است.

نگاره ۱: خلاصه‌ای از خصیصه‌ها و ویژگی‌های مجموعه داده اول

تعداد آرتیکل	تعداد سند	شماره روزانه سند	تاریخ اسناد	صادر کننده	روزهای ثبت سند	روزهای تایید اسناد	تعداد حساب‌های در گردش	ماهیت	زیر سیستم‌ها	طرف حساب	مرکز هزینه
۳۶۵۳۸	۲۶۱۵	۶۳	۳۰۶	۵	۲۴۴	۱۲	۱۷۵	۲	۵	۵۶۸	۱۲۳

برای مجموعه داده دوم مجموعه ۳۰.۰۰۰ خط آرتیکل ثبت حسابداری وجود دارد که نگاره ۳، خلاصه‌ای از ویژگی‌های مربوط به آنها را نشان می‌دهد. شایان ذکر است در این بخش از پایگاه داده‌ای نرم افزار سپیدار سیستم استفاده شده است. در این مجموعه داده نیز در مجموع ۱۰ ناهنجاری محلی (۰.۰۳۳٪) و ۱۰ ناهنجاری عمومی (۰.۰۳۳٪) وارد شد.

نگاره ۲: خلاصه‌ای از خصیصه‌ها و ویژگی‌های مجموعه داده دوم

تعداد آر تیکل	تعداد سند	شماره روزانه سند	تاریخ اسناد	صادر کننده	روزهای ثبت سند	روزهای تایید اسناد	تعداد حساب‌های در گردش	ماهیت	زیر سیستم‌ها	طرف حساب
۳۰.۰۰۰	۱۵۸۶	۵۴	۳۱۳	۴	۱۹۰	۱۹۵	۱۴۴	۲	۴	۳۹۱

۴.۲ آموزش شبکه عصبی خود رمزگذار

در حسابرسی سالانه هدف حساب‌رسان نمونه‌گیری و محدود کردن تعداد اسناد و ثبت‌های مشمول آزمون‌های محتوا است تا هیچ خطا یا ورودی مرتبط با تحریف را از دست ندهند. برگرفته از این موضوع، سه هدف در روند آموزش مورد توجه قرار گرفت: (۱) به حداقل رساندن خطای کلی بازسازی شبکه رمزگذار خودکار، (۲) تمرکز بر مدل‌هایی که ۱۰۰٪ آر تیکل‌های مصنوعی را تشخیص دهد و (۳) به حداکثر رساندن تشخیص شبکه رمزگذار خودکار به نحوی که تعداد هشدارهای مثبت کاذب به حداقل ممکن برسد. در این پژوهش از نه معماری متمایز شبکه‌های رمزگذار خودکار کم عمق (AE 1) تا عمیق (AE 9) برای آموزش شبکه عصبی استفاده شده است. به دلیل محدود بودن منابع رایانه‌ای و کاهش پیچیدگی بالای مدل و احتمال بیش‌برازشی با گسترش لایه‌های پنهان، معماری تا عمق ۹ لایه مورد بررسی قرار گرفت (برهمند و همکاران^{۷۳}، ۲۰۲۴). نگاره ۳ نمای کلی از معماری‌های استفاده شده را نشان می‌دهد.

نگاره ۳: معماری و ساختار شبکه استفاده شده

نماد	ساختار کلی شبکه
AE 1	[input features] – 3 – [output features]
AE 2	[input features] – 4 – 3 – 4 – [output features]
AE 3	[input features] – 8 – 4 – 3 – 4 – 8 – [output features]
AE 4	[input features] – 16 – 8 – 4 – 3 – 4 – 8 – 16 – [output features]
AE 5	[input features] – 32 – 16 – 8 – 4 – 3 – 4 – 8 – 16 – 32 – [output features]
AE 6	[input features] – 64 – 32 – 16 – 8 – 4 – 3 – 4 – 8 – 16 – 32 – 64 – [output features]
AE 7	[input features] – 128 – 64 – 32 – 16 – 8 – 4 – 3 – 4 – 8 – 16 – 32 – 64 – 128 – [output features]

نماد	ساختار کلی شبکه
AE 8	[input features] – 256 – 128 – 64 – 32 – 16 – 8 – 4 – 3 – 4 – 8 – 16 – 32 – 64 – 128 – 256 – [output features]
AE 9	[input features] – 512 – 256 – 128 – 64 – 32 – 16 – 8 – 4 – 3 – 4 – 8 – 16 – 32 – 64 – 128 – 256 – 512 – [output features]
Input features و output features مبین ابعاد ویژگی‌های مربوط به مجموعه داده‌های مورد استفاده است که برای مجموعه داده اول (دوم) ۳۶۱۲ (۲۱۸۵) بعد است. در این نگاره ساختار کلی شبکه عصبی برای یادگیری عمیق ارائه شده است.	

به منظور آموزش شبکه عصبی در این پژوهش از تابع فعال‌ساز تابع یکسو کننده خطی تعدیل شده^{۷۴} استفاده شده است و ضریب مقیاس مقادیر منفی برای آن $a = 0.4$ در نظر گرفته شد. هر کدام از شبکه معماری رمزگذار خودکار فوق با اعمال نرخ یادگیری 10^{-2} (در گردیان کاهشی) در سطح تمام لایه‌ها و با استفاده از دسته‌های ۱۲۸ تایی^{۷۵} از آرتیکل‌های حسابداری آموزش داده شد. علاوه بر این، به منظور تعیین روش بهینه یابی از تابع تخمین گشتاور تطبیقی^{۷۶} (ADAMS) استفاده شد و وزن هر یک از لایه‌های شبکه مقداردهی اولیه گردید. آموزش از طریق پس انتشار استاندارد تا زمان همگرایی^{۷۷} (مقداری که تابع هزینه الگوریتم بر اساس دوره آموزشی^{۷۸} بیشترین کاهش را داشته باشد) انجام شد. برای هر معماری، آزمایش‌ها پنج بار با استفاده از مقادیر پارامتر اولیه^{۷۹} متفاوت اجرا شد تا صحت نتایج مورد بررسی قرار گیرد.

پس از همگرایی آموزش شبکه، از مدل‌های آموزش دیده برای به دست آوردن خطاهای بازسازی^{۸۰} هر یک از آرتیکل‌های ثبت‌های ورودی استفاده شد. در فرآیند آموزش شبکه رمزگذار خودکار مورد استفاده (AE 9) از داده‌های مربوط به آرتیکل‌های اسناد حسابداری استفاده می‌شود. به علاوه نوع آموزش تقویتی و عمیق است. به منظور آموزش شبکه از تابع هزینه BCEWithLogitsLoss استفاده شده است. این تابع هزینه شامل یک لایه با تابع فعال‌ساز سیگموئیدی^{۸۱} و تابع هزینه آنتروپی مقطعی باینری^{۸۲} است. شبکه با استفاده از این تابع به دنبال حداقل سازی مقدار $\ell(x,y)$ در تابع زیر است:

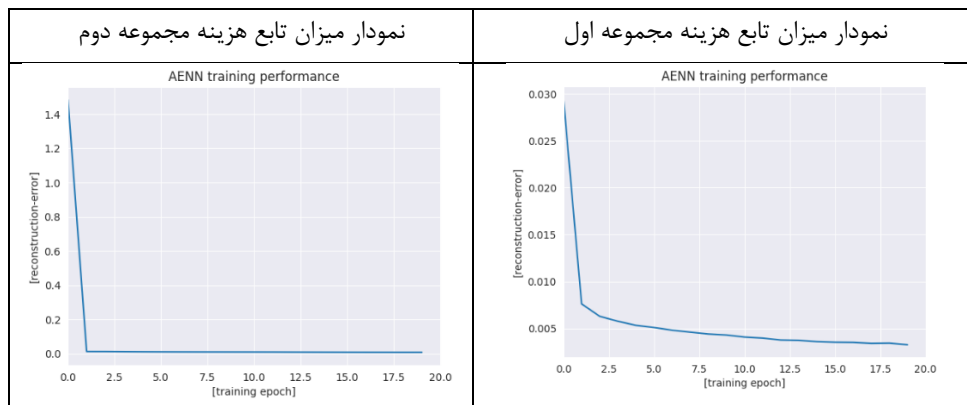
$$\ell(x,y)=L=\{l_1,\dots,l_N\} T, l_n=-wn[yn \cdot \log \sigma(xn) + (1-yn) \cdot \log(1-\sigma(xn))] \quad (9) \text{ معادله}$$

در این تابع N مبین حجم دسته‌های^{۸۳} ورودی برای آموزش است، y نشان‌دهنده خروجی است که بر اساس ویژگی‌ها و خصیصه‌های مربوط به اسناد بازسازی شده (X) است. در واقع تابع فوق سعی دارد بر اساس آنتروپی تعیین شده و تعیین بهترین وزن‌های ممکن (w) مقدار خطا و تابع زیان فوق

را به حداقل ممکن تقلیل دهد.

به منظور تعیین مناسب بودن دوره آموزش از نمودار عملکرد تابع آموزش برای شبکه مورد استفاده برای هر دو مجموعه داده‌های اول و دوم استفاده و نتایج در نگاره ۴ آمده است.

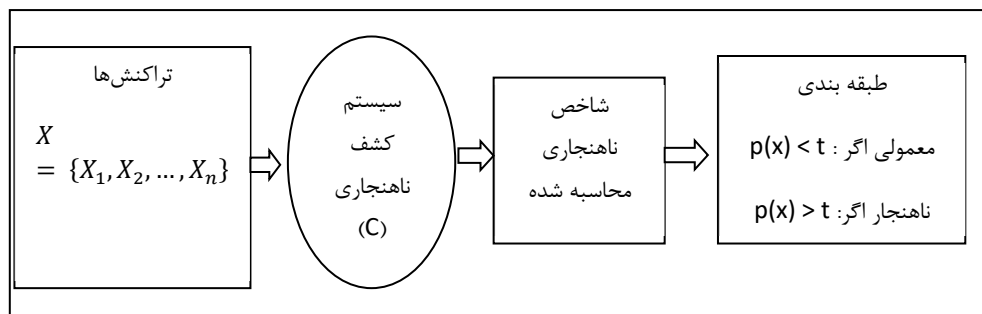
نگاره ۴: عملکرد تابع آموزش برای دو مجموعه اول و دوم



۴.۳ ارزیابی نتایج

به طور کلی به منظور پیاده سازی و استفاده از سیستم شناسایی ناهنجاری از شاخص ارزیابی آستانه‌ای^{۸۴} استفاده می‌شود. که شمای کلی این روش در نگاره ۵ آمده است.

نگاره ۵-ارزیابی آستانه ای نتایج



در این الگو سیستم کشف کننده ناهنجاری، تابع تمایزی (C) ایجاد می‌شود که بر اساس ورودی‌ها (X) احتمال تطبیق (P) را مشخص می‌کند، سپس با تعیین مقدار بهینه t (آستانه) بهترین امکان تمایز میان حالت‌های مختلف مشخص می‌شود.

با توجه به آنچه گفته شد می‌توان میزان دقت، صحت و خطای نوع اول و دوم سیستم تمایزکننده را در سطوح مختلف را مورد ارزیابی قرار داد. در این پژوهش سیستم کشف کننده ناهنجاری به وسیله شبکه عصبی خود رمزگذار تبیین شده است. به منظور بررسی مناسب بودن توانایی و خروجی نهایی از روش‌های یادگیری ماشین پر کاربرد در حوزه کشف ناهنجاری نیز استفاده شد تا مناسب بودن نتایج مورد ارزیابی قرار گیرد.

۵. تفسیر نتایج

پس از آموزش مناسب شبکه، امتیاز ارائه شده (فاصله میان آرتیکل‌های اسناد واقعی تا آنچه توسط شبکه بازسازی شده است) بر اساس دو معیار ارزیابی شده است:

(۱) آیا معماری‌های شبکه عصبی خود رمزگذار آموزش دیده، قادر به یادگیری مدلی از آرتیکل‌های اسناد صحیح و متعاقب آن تشخیص ناهنجاری‌ها است؟

(۲) آیا آرتیکل‌های صحیحی که به عنوان آرتیکل‌های ناهنجار دسته بندی شده‌اند، آنقدر مشکوک هستند که توسط حسابرسان مورد رسیدگی و آزمون قرار بگیرند؟

برای ارزیابی سوال اول و بررسی اثربخشی رویکرد پیشنهادی، طیفی از معیارهای ارزیابی شامل حساسیت، امتیاز $f1$ استاندارد و سطح زیر نمودار AUC برای بهترین روش‌های یادگیری ماشین در مقابل روش شبکه رمزگذار خودکار عمیق در نگاره زیر گزارش شده‌اند. شایان ذکر است که به دلیل عدم تقارن تعداد ناهنجاری‌ها (مشابه با آنچه در واقعیت رخ می‌دهد) در مقابل کل داده‌ها (کمتر از ۱ ده هزارم) برای ارزیابی از شاخص‌های مذکور استفاده شده است. نتایج این بخش در نگاره شماره ۶ آورده شده است.

نگاره ۶: شاخص‌های عملکرد مقایسه‌ای شبکه خود رمزگذار در مقابل سایر روش‌های مرسوم

مدل	حساسیت (Recall)	امتیاز fl	AUC	توضیحات تکمیلی
مجموعه داده اول				
Extra Trees	۰.۶۰۰۰	۰.۶۶۶۷	۰.۹۸۴۵	با توجه به آموزش و یادگیری صورت پذیرفته، مهم‌ترین ویژگی‌های مرتبط برای کشف ناهنجاری در اسناد حسابداری به ترتیب حساب تفضیلی (طرف مقابل)، حساب معین، مرکز هزینه و تاریخ آخرین ویرایش سند بودند.
Decision Tree Classifier	۰.۵۰۰۰	۰.۵۳۳۳	۰.۷۵۰۰	
Quadratic Discriminant Analysis	۰.۵۰۰۰	۰.۵۳۳۳	۰.۷۹۸۶	
AE_NN	۰.۸۰۹۵	۰.۶۶۶۷	۰.۹۹۵۵	
مجموعه داده دوم				
Ada Boost Classifier	۰.۲۵۰۰	۰.۲۶۶۷	۰.۸۹۱۰	با توجه به آموزش و یادگیری صورت پذیرفته، مهم‌ترین ویژگی‌های مرتبط برای کشف ناهنجاری در اسناد حسابداری به ترتیب حساب تفضیلی (طرف مقابل)، حساب معین، تاریخ آخرین ویرایش سند و تاریخ ثبت اولیه سند بودند.
Decision Tree Classifier	۰.۱۵۰۰	۰.۱۶۶۷	۰.۵۷۵۰	
Random Forest Classifier	۰.۱۵۰۰	۰.۱۶۶۷	۰.۸۴۹۹	
AE_NN	۰.۲۵۰۰	۰.۳۰۵۶	۰.۹۴۸۸	
Extra Trees یک الگوریتم یادگیری ماشین در دسته‌ی روش‌های ترکیبی^{۸۵} است که به‌طور خاص برای مسائل رگرسیونی و همچنین طبقه‌بندی به کار می‌رود. Decision Tree Classifier یک الگوریتم یادگیری نظارت‌شده است که داده‌ها را با استفاده از قواعد if-then به‌صورت درختی تقسیم‌بندی کرده و برای طبقه‌بندی نمونه‌ها به کار می‌رود. Quadratic Discriminant Analysis (QDA) یک الگوریتم طبقه‌بندی مبتنی بر مدل بیزین است که با فرض ماتریس کوواریانس متفاوت برای هر کلاس، مرزهای تصمیم غیرخطی (درجه دو) ایجاد می‌کند. AdaBoost Classifier یک الگوریتم ترکیبی از نوع تقویتی^{۸۶} است که چندین مدل ضعیف (مانند درخت‌های تصمیم کم‌عمق) را به‌صورت ترتیبی ترکیب می‌کند تا یک مدل قوی و دقیق برای طبقه‌بندی بسازد. Random Forest Classifier یک الگوریتم ترکیبی از نوع تقویتی است که با ساخت مجموعه‌ای از درخت‌های تصمیم و ترکیب آن‌ها، دقت طبقه‌بندی را افزایش داده و خطر بیش‌برازشی را کاهش می‌دهد. AE_NN روش شبکه عصبی خود رمزنگار است که در این پژوهش بر آن تمرکز شده است.				

بر اساس نتایج مربوط به ارزیابی عملکرد (حساسیت، امتیاز fl و سطح زیر منحنی AUC) و

امتیازهای بالاتر اخذ شده شبکه عصبی خود رمزگذار به نسبت سایر روش‌های یادگیری ماشین، این روش از مطلوب بالاتری برخوردار است و عملاً بهترین عملکرد را از لحاظ تشخیص ناهنجاری‌های درج شده دارد.

افزون بر مورد فوق، یکی از کارکردهای سیستم تشخیص ناهنجاری افزایش کارایی و اثربخشی در رسیدگی و کشف موارد ناهنجار است. به منظور بررسی این موضوع در دو مجموعه داده مورد بررسی و بر اساس خروجی نهایی شبکه رمزگذار خودکار در سه سطح ۱۰، ۵ و ۱ درصد رسیدگی به اسناد بر اساس نمره عدم انطباق و خطای بازیابی شبکه، تعداد ناهنجاری‌های کشف شده در نگاره ۷ ارائه شده است.

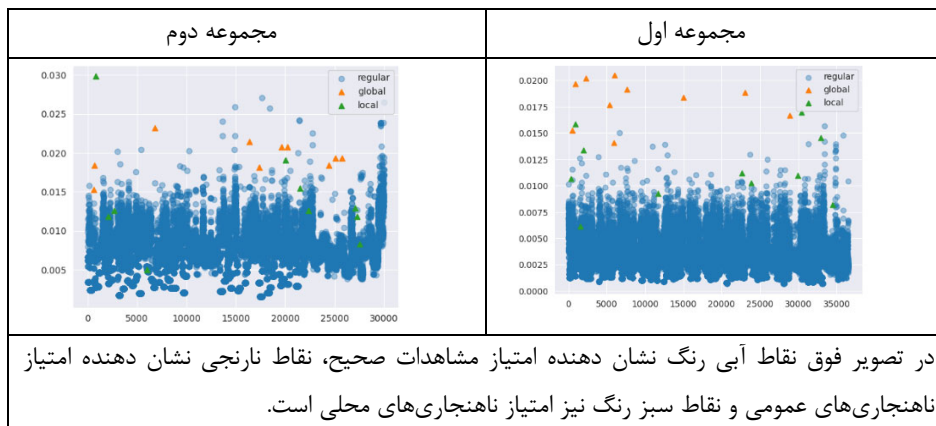
نگاره ۷ وضعیت عملکرد شبکه آموزش دیده برای کشف ناهنجاری‌ها

آستانه بررسی	تعداد کل مشاهده	تعداد ناهنجاری محلی کشف شده (درصد)	تعداد ناهنجاری عمومی کشف شده (درصد)
مجموعه داده اول			
۱۰ درصد	۳۶۵۴	۱۱ (٪۱۰۰)	۱۰ (٪۱۰۰)
۵ درصد	۱۸۲۷	۱۰ (٪۹۱)	۱۰ (٪۱۰۰)
۱ درصد	۳۶۶	۱۰ (٪۹۱)	۱۰ (٪۱۰۰)
مجموعه داده دوم			
۱۰ درصد	۳۰۰۰	۸ (٪۸۰)	۱۰ (٪۱۰۰)
۵ درصد	۱۵۰۰	۸ (٪۸۰)	۱۰ (٪۱۰۰)
۱ درصد	۳۰۰	۳ (٪۳۰)	۱۰ (٪۱۰۰)

نگاره فوق وضعیت کلی پس از اجرای شبکه و اخذ میزان خطای انطباق و بازیابی، عملکرد شبکه را نشان می‌دهد. به نحوی که در سطح رسیدگی به یک درصد بالایی خطای بازیابی برای مجموعه داده اول تنها یک مورد از ناهنجاری‌های محلی کشف نشده بود که این موضوع در سطح ۱۰ درصد رسیدگی کشف شده است. در سطح مجموعه داده دوم نیز تمامی ناهنجاری‌های عمومی از همان سطح یک درصد بالایی قابل کشف بوده است. لیکن برای ناهنجاری‌های محلی (که مطابق آنچه در بخش‌های ابتدایی مقاله گفته شد سطح پیچیدگی و دشواری بالاتری دارند) تا سطح ۵ درصد تنها سه مورد از ده مورد کشف شده بود. لیکن با افزایش آستانه رسیدگی به ۵ درصد، ۸۰ درصد از این

موارد نیز با استفاده از نمره خطای بازیابی قابل کشف بودند. نگاره ۸ شمای کلی از تشخیص خروجی نهایی شبکه رمزگذار خودکار را با توجه به سطح خطای بازیابی اسناد برای دو مجموعه اول و دوم را نشان می‌دهد.

نگاره ۸: نمودار میزان پراکندگی خطای بازیابی



در ادامه به منظور اطمینان از عدم بیش انطباقی باید وضعیت میانگین وزن خروجی تابع فعال‌ساز را در مرون‌های لایه گلوگاهی^{۸۷} مورد بررسی قرار گیرد. نگاره زیر این موضوع را برای دو مجموعه داده استفاده شده در عمیق‌ترین شبکه آموزشی (AE9) نشان می‌دهد. نتایج نشان می‌دهد که میانگین خروجی تابع فعال‌ساز برای سه دسته داده عادی (معمولی)، ناهنجاری محلی و ناهنجاری عمومی متفاوت است و لذا شبکه به درستی میان سه نوع داده تمایز قائل شده است. میانگین وزن‌های تابع فعال‌ساز نرون‌های شبکه در معماری شبکه AE9 و ۲۰ دوره آموزش به قرار نگاره ۹ است.

نگاره ۹: میانگین وزن‌های فعال‌ساز در معماری ۹ لایه‌ای

نوع داده‌ها	میانگین فعال‌ساز نرون‌ها
مجموعه داده اول	
همه داده‌ها	[۳.۷۱۶۱۹- , ۳.۱۹۹۹۶۱۴ , ۵.۳۹۵۴۰۸-]
معمولی	[۳.۷۱۶۷۵۶۶- , ۳.۲۰۱۱۱۵۸ , ۵.۳۹۶۶۴۰۳-]
ناهنجار محلی	[۴.۱۶۹۰۹۷۴- , ۱.۲۱۲۴۵۷۸ , ۴.۷۵۷۵۳۹۷-]
ناهنجار عمومی	[۱.۱۵۰۸۵۰۲- , ۱.۱۷۰۵۵۸ , ۱.۵۹۸۶۴۲۶-]

نوع داده‌ها	میانگین فعال‌ساز نرونها
مجموعه داده دوم	
همه داده‌ها	[۵.۲۱۲۴۱۵ , ۲.۲۴۴۶۹۲۶ , ۶.۸۰۵۴۱۶۶]
معمولی	[۵.۲۱۱۱۰۸ , ۲.۲۴۳۶۱۲۷ , ۶.۸۰۵۸۳۳۳]
ناهنجار محلی	[۵.۴۸۷۳۱۶ , ۲.۹۸۹۵۱۰۸ , ۵.۹۲۹۶۳۳]
ناهنجار عمومی	[۸.۸۵۷۶۴۶ , ۴.۷۳۸۵۵۰۷ , ۶.۴۳۲۸۸۳۳]

یکی دیگر از نتایج مورد بررسی این است که آیا در زیر سیستم‌های مختلف میزان خطای انطباق و بازسازی شبکه مشابه است یا خیر. هرچه میزان خطای بازسازی در یک زیر سیستم کمتر باشد (میانگین و انحراف معیار) تشخیص موارد غیر معمول و ناهنجار آسان‌تر است. منظور از خطای بازیابی متفاوت بودن نتایج نهایی انطباق اجزا با مقادیر واقعی است. نگاره ۱۰ میزان خطای بازسازی شبکه عصبی خود رمزگذار برازش شده را برای دو مجموعه داده اول و دوم نشان می‌دهد.

نگاره ۱۰: میزان خطای بازسازی شبکه عصبی خود رمزگذار به ازای هر یک از زیر سیستم‌های استفاده شده

زیر سیستم	میانگین	انحراف معیار
مجموعه داده راهکاران		
دفتر کل عمومی	۰.۰۰۴۴۷۶	۰.۰۰۱۹۶۲
دریافت و پرداخت	۰.۰۰۲۷۸۷	۰.۰۰۱۶۱۱
حقوق و دستمزد	۰.۰۰۳۱۱۶	۰.۰۰۰۷۵۹
انبار و تدارکات	۰.۰۰۵۷۸۶	۰.۰۰۱۲۵۷
دارایی ثابت و اموال	۰.۰۰۴۴۶۶	۰.۰۰۱۳۰۲
مجموعه داده سپیدار		
دفتر کل عمومی	۰.۰۰۶۵۱۰	۰.۰۰۳۰۶۷
دریافت و پرداخت	۰.۰۱۱۶۰۹	۰.۰۰۲۲۵۸
حقوق و دستمزد	۰.۰۰۶۹۹۵	۰.۰۰۱۵۹۴

نتایج این بخش از پژوهش نشان داد که زیر سیستم‌های مختلف دارای خطای بازیابی و انطباق^{۸۸} متفاوتی هستند. لذا به منظور بررسی دقیق‌تر می‌توان اسناد و آرتیکل‌های هر دسته از زیر سیستم‌های

نرم افزار حسابداری را به صورت جداگانه مورد بررسی قرار داد. این نتیجه‌گیری با توجه به تفاوت ماهوی هر یک از چرخه‌های عملیاتی در عملکرد و ساختار اسناد آنها منطقی است.

۶. بحث و نتیجه‌گیری

بر پایه نتایج پژوهش حاضر، روش یادگیری عمیق بر پایه شبکه عصبی خود رمزگذار توانایی بالایی در تشخیص اسناد و آرتیکل‌های نامتعارف و غیر معمول را دارد. این نوع از یادگیری ماشین سعی می‌کند در مرحله اول به اسناد ویژگی‌ها و ورودی‌های اولیه، درک مناسبی از وضعیت آنها بدست آورد و سپس اقدام به بازسازی اسناد بر اساس الگوهای یادگرفته شده با استفاده از ورودی و ویژگی‌های واقعی اسناد می‌کند. پس از این مرحله تصویر و موضوع ایجاد شده با واقعیت مورد بررسی قرار می‌گیرد و به اسناد میزان تفاوت، سطح عدم انطباق و تناظر با واقعیت به عنوان معیار اختلاف و خطا مورد بررسی قرار داده می‌شود. کشف تقلب و تحریف در اسناد مالی و حسابداری که عمدتاً برای مخفی‌سازی این رویدادها در دل سایر اسناد حسابداری قرار می‌گیرند، همواره مورد توجه شرکت‌ها و ارکان نظارتی بوده است. در پژوهش حاضر سعی شد تا بر اساس سناریوسازی و درج تعداد محدودی تقلب و تحریف از دو نوع محلی و عمومی در پایگاه داده‌ای دو نرم افزار پر استفاده حسابداری و مالی در ایران، کارایی و اثربخشی این روش یادگیری ماشین مورد بررسی قرار گیرد. نتایج مبین این موضوع بود که روش شبکه عصبی خود رمزگذار می‌تواند به شیوه مناسبی میان داده‌های درست و داده‌های متقلبانه تمایز قائل شود و به علاوه این نوع از شبکه‌ها به نسبت سایر روش‌های خوشه بندی و طبقه بندی متداول توانمندتر و مناسب‌تر بوده است.

۷. پیشنهادها و محدودیت‌های پژوهش

یکی از نوآوری‌های پژوهش حاضر استفاده از داده‌های تراکنشی^{۸۹} سیستم اطلاعاتی حسابداری است. بهره‌گیری از این داده‌ها با وجود ایجاد فرصت‌های پژوهشی جدید، موجب کارآمدی هر چه بیشتر پژوهش‌های حسابداری و حسابرسی خواهد شد. پژوهش حاضر نشان داد که با استفاده از داده تراکنش‌های حسابداری و تکنیک شبکه عصبی خود رمزنگار می‌توان به شیوه‌ای مناسب ناهنجاری موجود در داده‌ها را کشف کرد. لذا در راستای بهبود روش‌های کشف ناهنجاری (تحریف یا اشتباه)، استفاده از فنون نوین و مبتنی بر یادگیری عمیق توصیه می‌شود.

به دلیل مسائل محرمانگی و حریم خصوصی داده‌های واقعی تراکنش‌های حسابداری با

محدودیت زیادی در دسترس پژوهشگران قرار می‌گیرد. لذا ایجاد پایگاه داده‌ای در نهادهایی نظیر سازمان امور مالیاتی، سازمان حسابرسی، دیوان محاسبات و سایر دستگاه‌های مشابه و با حفظ محرمانگی و حریم خصوصی می‌تواند منجر به ترویج و گسترش پژوهش‌های کاربردی در مطالعات مشابه خواهد شد.

یکی از محدودیت‌های دیگر حاکم بر پژوهش‌های یادگیری عمیق و شبکه‌های عصبی عدم امکان تشریح دقیق و مناسب عملکرد مدل نهایی است. لذا استفاده از فنون جدید نظیر هوش مصنوعی توضیح دهنده^{۹۰} در پژوهش‌های آتی می‌تواند منجر به شفافیت و افزایش درک نسبت به عملکرد شبکه‌های عصبی شود. به علاوه استفاده از روش‌های نوظهور دیگر (نظیر شبکه‌های عصبی شناختی) می‌تواند منجر به بهبود و گسترش روش‌های کشف ناهنجاری در سیستم‌های اطلاعاتی حسابداری شود.

پژوهش حاضر در سطح اسناد حسابداری بوده است. لذا تکرار این پژوهش در سایر حوزه‌ها از جمله داده‌های موجود در سیستم‌های تراکنشی خزانه‌داری و سیستم‌های مدیریت منابع انسانی می‌تواند منجر به بسط استفاده از روش‌های کشف ناهنجاری‌های نوین شود. افزون بر این استفاده از معماری‌های دیگر شبکه‌های رمزگذار خودکار می‌تواند به عنوان بخشی از پژوهش‌های آتی در این حوزه باشد.

یادداشت‌ها

- | | |
|---|-----------------------------|
| 1. Guo et al | 2. Debreceeny and Gray |
| 3. Association of Certified Fraud Examiners (ACFE) | |
| 4. ERP | 5. Reg Flag |
| 6. Sun and Vasarhelyi | 7. Bao et al |
| 8. Huang et al | 9. Sharma and Lokanan |
| 10. Beck and Solomon | 11. Teitlebaum and Robinson |
| 12. Deep Learning | 13. LeCun et al |
| 14. Deep autoencoder neural networks | 15. Wei et al |
| 16. Amani and Fadlalla | 17. Wells |
| 18. Singleton and Singleton | |
| ۱۹. شاید دلیل اصلی این موضوع، عدم امکان دسترسی مناسب به زیر سیستم‌ها و ریز اسناد توسط پژوهشگران مستقل باشد. | |
| 20. Teoh et al | 21. Pincus et al |
| 22. Drake et al | 23. Yan and Zheng |
| 24. Lokanan et al | 25. Chen and Velikov |
| 26. Forensic Accounting | 27. Bay et al |
| 28. McGlohon et al | 29. Khan and Corney |
| 30. Khan et al | 31. ERP SAP R/3 |
| 32. Islam et al | 33. Debreceeny and Gray |
| 34. Jans et al | 35. Nonnenmacher and Gómez |

36. Accounting Graph Topology (AGT) 37. Cognitive Fit Theory
 38. General Ledger (GL)
 39. Multilevel Outlier Detection Framework (MODF)
 40. Hawkins et al 41. Williams et al
 42. Replicator Neural Networks (RNN) 43. Hernandez et al
 44. Bakumenko and Elragal 45. Unsupervised
 46. attributes 47. replicator neural network
 48. Feed forward multilayer 49. reconstruction error
 50. Rumelhart et al 51. sigmoid function
 52. offset bias vectors 53. Loss Function
 54. Bernoulli 55. Bengio et al
 56. Overfitting
۵۷. این موضوع معمولاً تحت عنوان معماری گلوگاه (bottleneck) شناخته می‌شود.
58. Fang and Wang 59. Chandola et al
 60. Point Anomalies 61. Contextual
 62. Conditional 63. Contextual attributes
 64. Behavioral attributes 65. Collective
 66. Breunig et al 67. Global accounting anomalies
 68. Local accounting anomalies 69. False positive
 70. Schreyer et al
۷۱. داده‌ها بر پایه ثبت‌های حسابداری دو شرکت ارائه خدمات رایانه‌ای بوده است.
۷۲. یکی از روش‌های یادگیری عمیق با استفاده از شبکه‌های عصبی خود رمزگذار روش خود رمزگذار پارازیت (نویز) زد (Denoising Autoencoder) است. که در آن به منظور بررسی عملکرد نهایی یادگیری داده‌های اولیه پارازیت (نویز) دار را ایجاد می‌کنند تا توانایی مدل در تشخیص این موضوع را مشخص کنند (وینسنت و همکاران، ۲۰۰۸). این تعداد ناهنجاری همانند پژوهش‌های مشابه بر پایه نظر محقق تعیین شده است.
73. Berahmand et al 74. Leaky ReLU
 75. Batch Size 76. Adaptive moment estimation
 77. Convergence 78. Epochs
 79. Seed 80. Reconstruction errors
 81. Sigmoid 82. Binary Cross Entropy
 83. Batch 84. Threshold-based metrics
 85. Ensemble Methods
۸۶. روش تقویتی یا Boosting یک روش ترکیبی است که چندین مدل ضعیف (Weak Learners) را به‌صورت توالی (Sequential) آموزش می‌دهد. هر مدل جدید روی خطاهای مدل قبلی تمرکز می‌کند تا در نهایت یک مدل قوی و دقیق ساخته شود.
۸۷. لایه‌ای است که ویژگی‌ها به کمترین ابعاد خلاصه و بعد آن فرآیند بازسازی آغاز می‌شود. در پژوهش حاضر و بنابر معماری بیان شده، این لایه متشکل از سه نرون است.
۸۸. به منظور محاسبه خطا از تفاضل میان خروجی نهایی مدل (آنچه از بازسازی داده‌های اولیه و بر پایه نتایج یادگیری حاصل شده است) با واقعیت استفاده شده است (معادله ۶).
89. Transactional Data 90. Explainable AI (XAI)

منابع

ACFE (2024). Report to the Nations on Occupational Fraud and Abuse, The 2016 Global Fraud Study. Association of Certified Fraud Examiners,

- <https://www.acfe.com/-/media/files/acfe/pdfs/rtn/2024/2024-report-to-the-nations.pdf>
- AICPA (2022). SASs Consideration of Fraud in a Financial Statement Audit. American Institute of Certified Public Accountants, <https://www.aicpa-cima.com/resources/download/aicpa-statements-on-auditing-standards-currently-effective/>
- Amani, F.A., & Fadlalla, A.M. (2017). Data mining applications in accounting: A review of the literature and organizing framework. *International Journal of Accounting Information Systems* 24, 32.
- Bakumenko, A., & Elragal, A. (2022). Detecting anomalies in financial data using machine learning algorithms systems, *10(5)*, 130. <https://doi.org/10.3390/systems10050130>
- Bao, Y., Ke, B., Li, B., Yu, Y.J. and Zhang, J. (2020), Detecting accounting fraud in publicly traded U.S. firms using a machine learning approach. *Journal of Accounting Research*, 58: 199-235. <https://doi.org/10.1111/1475-679X.12292>
- Bay, S., Kumaraswamy, K., Anderle, M.G., Kumar, R., Steier, D.M., Blvd, A., & Jose, S. (2006). Large scale detection of irregularities in accounting data. In: Data Mining. ICDM'06. *Sixth International Conference on Data Mining (ICDM'06)*, Hong Kong, China, 2006, pp. 75-86, doi: 10.1109/ICDM.2006.93.
- Beck, P. J., & Solomon. I. (1985). Sampling risks and audit consequences under alternative testing approaches. *The Accounting Review* 60 (4): 714–723. <https://www.jstor.org/stable/247467>
- Bengio, Y., Yao, L., Alain, G., & Vincent, P. (2013). Generalized denoising autoencoders as generative models. In Proceedings of the 27th International Conference on Neural Information Processing Systems - Volume 1 (NIPS'13), Vol. 1. Curran Associates Inc., Red Hook, NY, USA, 899–907.
- Berahmand, K., Daneshfar, F., Salehi, E. S., Li, Y., & Xu, Y. (2024). Autoencoders and their applications in machine learning: A survey. *Artificial Intelligence Review*, 57(2), 28. <https://doi.org/10.1007/s10462-023-10662-6>
- Breunig, M.M., Kriegel, H.P., Ng, R.T., Sander, J. (2000). LOF: Identifying density-based local outliers. In: Proceedings of the Acm Sigmod International Conference on Management of Data. pp. 1–12.
- Chandola, V., Banerjee, A., & Kumar, V. (2009). Anomaly detection: A survey. *ACM computing surveys (CSUR)*, 41(3), 1-58
- Chen, A. Y., & Velikov, M. (2023). Zeroing In on the Expected Returns of

- Anomalies. *Journal of Financial & Quantitative Analysis*, 58(3), 968–1004. doi:10.1017/S0022109022000874
- Debreceeny, R.S., & Gray, G.L. (2010). Data mining journal entries for fraud detection: An exploratory study. *International Journal of Accounting Information Systems* 11(3), 157 – 181
- Deng, Q., & Mei, G. (2009). Combining self-organizing map and K-means clustering for detecting fraudulent financial statements. *IEEE International Conference on Granular Computing*, 126–131.
- Drake Michael S. & Guest Nicholas M. & Twedt Brady J. (2014). The Media and mispricing: The role of the business press in the pricing of accounting information, *The Accounting Review*, 89 (5): 1673–1701.
- Domingos, S. L., Carvalho, R. N., Carvalho, R. S., & Ramos, G. N. (2016). Identifying IT purchases anomalies in the Brazilian government procurement system using deep learning. 2016 15th IEEE International Conference on Machine Learning and Applications (ICMLA), 722–727. <https://doi.org/10.1109/ICMLA.2016.0129>
- Fang, Z., & Wang, S. (2024). Boosting financial market prediction accuracy with deep learning and big data. *Journal of Organizational & End User Computing*, 36(1). <https://doi.org/10.4018/JOEUC.358454>
- Gomes, T. A., Carvalho, R. N., & Carvalho, R. S. (2017). Identifying anomalies in parliamentary expenditures of Brazilian Chamber of Deputies with deep autoencoders, *16th IEEE International Conference on Machine Learning and Applications (ICMLA)*, Cancun, Mexico, 2017, pp. 940-943, doi: 10.1109/ICMLA.2017.00-33.
- Guo, K. H., Yu, X., & Wilkin, C. (2022). A picture is worth a thousand journal entries: Accounting graph topology for auditing and fraud detection. *Journal of Information Systems*, 36(2), 53–81. <https://doi.org/10.2308/ISYS-2021-003>
- Hawkins, S., He, H., Williams, G., & Baxter, R. (2002). Outlier Detection Using Replicator Neural Networks. In: *International Conference on Data Warehousing and Knowledge Discovery Lecture Notes in Computer Science*, vol 2454. Springer, Berlin, Heidelberg. https://doi.org/10.1007/3-540-46145-0_17
- Hernandez Aros, L., Bustamante Molano, L. X., Gutierrez-Portela, F., Moreno Hernandez, J. J., & Rodríguez Barrero, M. S. (2024). Financial fraud detection through the application of machine learning techniques: A literature review. *Humanities & Social Sciences Communications*, 11(1), Article 1130. <https://doi.org/10.1057/s41599-024-03606-0>

- Huang, F., No, W. G., Vasarhelyi, M. A., & Yan, Z. (2022). Audit data analytics, machine learning, and full population testing. *The Journal of Finance & Data Science*, 8, 138–144. <https://doi.org/10.1016/j.jfds.2022.05.002>
- Islam, A.K., Corney, M., Mohay, G., Clark, A., Bracher, S., Raub, T., & Flegel, U. (2010). Fraud detection in ERP systems using Scenario matching. In: Rannenber, K., Varadharajan, V., Weber, C. (eds) Security and Privacy – Silver Linings in the Cloud. SEC 2010. IFIP Advances in Information and Communication Technology, vol 330. Springer, Berlin, Heidelberg. https://doi.org/10.1007/978-3-642-15257-3_11
- Jans, M., Lybaert, N., & Vanhoof, K. (2007). Data mining for fraud detection: Toward an improvement on internal control systems? *Proceedings of the 30th Annual Congress European Accounting Association (EAA2007)*.
- Jans, M., Lybaert, N., & Vanhoof, K. (2010). Internal fraud risk reduction: Results of a data mining case study. *International Journal of Accounting Information Systems*, 11(1), 17–41.
- Jans, M., Van DerWerf, J.M., Lybaert, N., & Vanhoof, K. (2011). A business process mining application for internal transaction fraud mitigation. *Expert Systems with Applications*, 38(10), 13351-13359
- Khan, R., Corney, M., Clark, A., Mohay, G. (2010). Transaction Mining for Fraud Detection in ERP Systems. *Industrial Engineering & Management Systems* 9(2), pp. 141 – 156
- Khan, R., & Corney, M. (2009). A role mining inspired approach to representing user behavior in ERP systems. In: Proceedings of the 10th Asia Pacific Industrial Engineering and Management Systems Conference. pp. 2541 - 2552
- Kogan, A., Alles, M. G., Vasarhelyi, M. A., & Wu, J. (2014). Design and Evaluation of a Continuous Data Level Auditing System. *AUDITING: A Journal of Practice & Theory*, 33(4), 221–245. <https://doi.org/10.2308/ajpt-50844>
- Kuna, H. D., García-Martínez, R., & Villatoro, F. R. (2014). Outlier detection in audit logs for application systems. *Information Systems*, 44, 22–33. <https://doi.org/10.1016/j.is.2014.03.001>
- Lokanan, M., Tran, V. & Vuong, N.H. (2019), Detecting anomalies in financial statements using machine learning algorithm: The case of Vietnamese listed firms, *Asian Journal of Accounting Research*, 4(2), 181-201.
- Lu, F., Boritz, J. E., & Covey, D. (2006). Adaptive fraud detection using Benford's law. *Canadian AI*, 347–358. https://doi.org/10.1007/11766247_30.

- McGlohon, M., Bay, S., Anderle, M.G.M., Steier, D.M., & Faloutsos, C. (2009). SNARE: A Link Analytic System for Graph Labeling and Risk Detection. *Kdd-09: 15th Acm Sigkdd Conference on Knowledge Discovery and Data Mining*.
- Nonnenmacher, Jakob & Gómez Jorge Marx (2021). Unsupervised anomaly detection for internal auditing: Literature review and research agenda. *The International Journal of Digital Accounting Research*. Vol. 21, pp. 1-22.
- Paula, E. L., Ladeira, M., Carvalho, R. N., & Marzagão, T. (2016). Deep Learning Anomaly Detection as Support Fraud Investigation in Brazilian Exports and Anti-Money Laundering. 2016 15th IEEE International Conference on Machine Learning and Applications (ICMLA), 954–960. <https://doi.org/10.1109/ICMLA.2016.0172>
- Pincus, Morton & Rajgopal, Shivaram & Venkatachalam, Mohan (2007). The Accrual Anomaly: International Evidence. *The Accounting Review* 82 (1): 169–203.
- Rumelhart, D.E., Hinton, G.E., & Williams, R.J. (1986). Learning internal representation by error propagation. In: *Parallel distributed processing: explorations in the microstructure of cognition*, MIT Press, 1987, pp.318-362.
- Schreyer, M., Sattarov, T., Borth, D., Dengel, A.R., & Reimer, B. (2017). Detection of anomalies in large scale accounting data using Deep Autoencoder Networks. *ArXiv*, abs/1709.05254.
- Schreyer, M., Sattarov, T., Schulze, C., Reimer, B., & Borth, D. (2019). Detection of accounting anomalies in the latent space using Adversarial Autoencoder Neural Networks. <https://arxiv.org/abs/1908.00734>
- Sharma, S., & Lokanan, M. (2025). The use of machine learning algorithms to predict financial statement fraud. *The British Accounting Review*, 57(1), 101560. <https://doi.org/10.1016/j.bar.2025.101560>
- Singleton, T., Singleton, A.J. (2010). *Fraud auditing and forensic accounting* (4th ed). John Wiley & Sons.
- Sun, T., & Vasarhelyi, M. A. (2017, June). Deep learning and the future of auditing: How an evolving technology could transform analysis and improve judgment. *The CPA Journal*, 87(6), 24–29. <https://www.cpajournal.com/2017/06/19/deep-learning-future-auditing/>
- Teitlebaum, A. D., & C. F. Robinson. (1975). The real risks in audit sampling. *Journal of Accounting Research* 13: 70–91. <https://doi.org/10.2307/2490480>
- Teoh S.H. & Welch I. & Wong T.J. (1998). Earnings management and the

- underperformance of seasoned equity offerings *Journal of Financial Economics*, 50, pp. 63-99
- Thiprungsri, S., & Vasarhelyi, M. A. (2011). Cluster analysis for anomaly detection in accounting data: An audit approach. *International Journal of Digital Accounting Research*, 11, 69-84. https://doi.org/10.4192/1577-8517-v11_4
- Vincent, P., Larochelle, H., Bengio, Y., & Manzagol, P. (2008). Extracting and composing robust features with denoising autoencoders. *International Conference on Machine Learning*.
- Wei, D., Cho, S., Vasarhelyi, M. A., & Te-Wierik, L. (2024). Outlier detection in auditing: Integrating unsupervised learning within a multilevel framework for general ledger analysis. *Journal of Information Systems*, 38(2), 123–142. <https://doi.org/10.2308/ISYS-2022-026>
- Wells, J.T. (2017). *Corporate Fraud Handbook: Prevention and Detection*. John Wiley & Sons.
- Williams, G., Baxter, R., He, H., Hawkins, S., & Gu, L. (2002) A comparative study of RNN for outlier detection in data mining. *IEEE International Conference on Data Mining*, 1–16.
- Yan, X. (S.) & Zheng, L. (2017). Fundamental Analysis and the Cross-Section of Stock Returns: A Data-Mining Approach, *The Review of Financial Studies*, 30(4), 1382–1423.