

Predicting Auditor Fraudulent Behavior Using Artificial Neural Networks

Reza Talebi 

Department of Accounting, Khomain Unit, Islamic Azad University, Khomain, Iran. Email:
talebireza55@gmail.com

Azar Moslemi 

Department of Accounting, Khomain Unit, Islamic Azad University, Khomain, Iran,
azar.moslemi.kh@gmail.com

Mehdi Khorramabadi  *

Department of Accounting, Payame Noor University, Tehran, Iran, M.khorramabadi@pnu.ac.ir

Abstract

Objective: This study aims to predict auditors' behavior when faced with fraud and financial misconduct and to identify the factors influencing their fraudulent behavior. It seeks to propose measures for strengthening audit firms' supervisory systems through analytic models and advanced technologies.

Method: A quantitative approach was adopted, targeting auditors at audit firms in Tehran. The sample size was determined as 256 using the Morgan table. Data were collected via a questionnaire designed to identify factors affecting auditors' propensity for fraud. A neural network model was developed in Python, encompassing key steps such as data preprocessing, splitting into training, testing, and validation sets, selecting the network architecture, and adjusting weights via backpropagation algorithms. Data preprocessing included detrending and normalization. A multilayer perceptron with various learning and optimization algorithms was employed, and mean squared error (MSE) was used to evaluate the network's performance.

Findings: The neural network model demonstrated high accuracy—96.15% overall—in distinguishing auditors with and without a history of fraud. This high level of precision indicates the model's effectiveness in accurately identifying the factors that drive auditors' fraudulent behavior and making reliable predictions.

Conclusion: The results show that neural network models and analytic technologies can significantly help audit firms enhance their supervisory systems. By implementing these models, the prevention of financial misconduct is improved, fostering a more transparent and secure working environment. Consequently, this approach can play a crucial role in enhancing the quality and credibility of financial audits.

Scientific Contribution: The novelty of this research lies in applying neural networks to predict auditors' responses to fraud and financial misconduct, thereby markedly improving the accuracy and speed of supervisory decision-making.

Keywords: Auditor Behavior Prediction, Financial Misconduct, Neural Network.

Research Article

Cite this article: Talebi, Moslemi & Khorramabadi (2025) Predicting Auditor Fraudulent Behavior Using Artificial Neural Networks, *Journal of Financial Accounting Knowledge*, Vol.12, NO.2, Summer, 119-143.

DOI: 10.30479/jfak.2025.21163.3248

Received on 8 November, 2024 **Accepted on** 4 June, 2025

© The Author(s). 

Publisher: Imam Khomeini International University.

Corresponding Author: Mehdi Khorramabadi (M.khorramabadi@pnu.ac.ir)

پیش‌بینی رفتار متقلبانه حساب‌رسان با استفاده از شبکه عصبی مصنوعی

رضا طالبی 

گروه حسابداری، واحد خمین، دانشگاه آزاد اسلامی، خمین، ایران talebireza55@gmail.com

آذر مسلمی 

گروه حسابداری، واحد خمین، دانشگاه آزاد اسلامی، خمین، ایران azar.moslemi.kh@gmail.com

مهدی خرم‌آبادی* 

گروه حسابداری، دانشگاه پیام نور، تهران، ایران M.khorramabadi@pnu.ac.ir

چکیده

هدف: هدف این پژوهش پیش‌بینی رفتار متقلبانه حساب‌رسان در زمینه تخلفات مالی و شناسایی عوامل مؤثر بر رفتار متقلبانه آن‌ها است. این مطالعه به دنبال ارائه راهکارهایی برای تقویت سیستم‌های نظارتی مؤسسات حسابرسی از طریق استفاده از مدل‌های تحلیلی و فناوری‌های نوین هست.

روش: در این پژوهش از رویکرد کمی استفاده شده و جامعه آماری آن شامل حساب‌رسان مؤسسات حسابرسی شهر تهران است. حجم نمونه با استفاده از جدول مورگان، ۲۵۶ نفر تعیین شده است. داده‌های پژوهش از طریق پرسشنامه‌ای به‌منظور شناسایی عوامل مؤثر بر رفتار متقلبانه حساب‌رسان جمع‌آوری گردید. برای طراحی مدل شبکه عصبی، از نرم‌افزار پایتون بهره گرفته شده و مراحل کلیدی شامل پردازش داده‌ها، تقسیم‌بندی به داده‌های آموزشی، آزمون و اعتبارسنجی، انتخاب نوع شبکه و تنظیم وزن‌ها با استفاده از الگوریتم‌های پس انتشار خطا انجام گردید. پردازش داده‌ها شامل روند زدایی و نرمال‌سازی بوده و داده‌ها به دسته‌های آموزشی، آزمون و اعتباری تقسیم شده‌اند. مدل پرسپترون همراه با الگوریتم‌های مختلف برای یادگیری و بهینه‌سازی استفاده شده و معیار میانگین مربعات خطا (MSE) برای سنجش قدرت شبکه به‌کاررفته است.

یافته‌ها: نتایج تحلیل‌های پژوهش بیانگر دقت بالای مدل شبکه عصبی در شناسایی حساب‌رسان با و بدون سابقه تقلب با دقت کلی ۹۶,۱۵٪ است. دقت بالا نشان می‌دهد که مدل پیشنهادی قادر است به‌طور مؤثری عوامل مؤثر بر رفتار متقلبانه حساب‌رسان را شناسایی کرده و پیش‌بینی‌های دقیقی ارائه دهد.

نتیجه‌گیری: این پژوهش نشان می‌دهد که استفاده از مدل‌های شبکه عصبی و فناوری‌های تحلیلی می‌تواند به‌طور قابل توجهی به مؤسسات حسابرسی در تقویت سیستم‌های نظارتی خود کمک کند. با به‌کارگیری این مدل‌ها، امکان پیشگیری از تخلفات مالی افزایش یافته و محیط کاری شفاف‌تر و امن‌تری ایجاد می‌شود. بنابراین، پیاده‌سازی این رویکردها می‌تواند نقش مهمی در ارتقای کیفیت و اعتماد به حسابرسی‌های مالی ایفا کند.

دانش‌افزایی: نوآوری علمی این پژوهش در به‌کارگیری شبکه‌های عصبی برای پیش‌بینی واکنش حساب‌رسان در مواجهه با تقلب و تخلفات مالی است که می‌تواند دقت و سرعت تصمیم‌گیری‌های نظارتی را به‌طور چشمگیری ارتقا دهد.

واژگان کلیدی: پیش‌بینی رفتار حساب‌رسان، تخلفات مالی، شبکه عصبی.

مقاله پژوهشی

*استناد: طالبی، مسلمی و خرم‌آبادی (۱۴۰۴)، پیش‌بینی رفتار متقلبانه حساب‌رسان با استفاده از شبکه عصبی مصنوعی، فصلنامه علمی دانش حسابداری

مالی، مقاله پژوهشی، دوره ۱۲، شماره ۲، پیاپی ۴۵، تابستان ۱۴۰۴، ۱۱۹-۱۴۳.

تاریخ دریافت مقاله: ۱۴۰۳/۰۸/۱۸ تاریخ پذیرش نهایی: ۱۴۰۴/۰۳/۱۴

ناشر: دانشگاه بین‌المللی امام خمینی (ره) © حق مؤلف نویسندگان



۱- مقدمه

در فضای پیچیده و پویای اقتصادی کنونی، سازمان‌ها برای بقا و رشد نیازمند شفافیت، پاسخ‌گویی و انطباق با استانداردهای مالی و مقررات نظارتی هستند. در این میان، حسابرسی به‌عنوان فرآیندی مستقل، نقش اساسی در ارزیابی صحت اطلاعات مالی، کاهش ریسک‌های سازمانی و ایجاد اعتماد در میان ذینفعان ایفا می‌کند (وحیدخواتی و همکاران، ۲۰۲۲: ۱). با این حال، با پیچیده‌تر شدن ساختارهای مالی، توسعه فناوری‌های نوین، و افزایش انتظارات عمومی از حسابرسان، کیفیت فرآیند و گزارش حسابرسی بیش از پیش در معرض چالش قرار گرفته است (سینق و همکاران، ۲۰۱۹: ۶۴). در فرآیند حسابرسی، کیفیت اقدامات و رویه‌های اجرایی (کیفیت فرآیند حسابرسی) و کیفیت گزارش نهایی (کیفیت گزارش حسابرسی) دو مقوله مستقل اما وابسته‌اند؛ به عبارت دیگر، دقت و جامعیت روش‌های گردآوری شواهد و رعایت استانداردهای حرفه‌ای در فرآیند حسابرسی مستقیماً بر قابلیت اتکا و شفافیت گزارش حسابرسی تأثیر می‌گذارد. هرگونه کوتاهی در طراحی و اجرای برنامه‌های آزمون کنترل‌ها و شواهد، ممکن است منجر به اظهارنظرهای نادرست یا کاستی در نتایج شود و بدین‌سان کیفیت گزارش را تضعیف کند. بنابراین، برای تضمین اعتبار نتایج، لازم است اول کیفیت فرایند حسابرسی با استفاده از ابزارهایی نظیر بازمینی هم‌سطح و روش‌های نمونه‌برداری مناسب به‌طور مستمر ارزیابی و بهبود یابد، تا در نهایت گزارش حسابرسی با دقت، بی‌طرفی و شفافیت کافی منتشر گردد (ژیائو و همکاران، ۲۰۲۰: ۱۱۲).

یکی از مؤلفه‌های کلیدی که می‌تواند به‌طور مستقیم بر اثربخشی حسابرسی تأثیرگذار باشد، رفتار حرفه‌ای حسابرسان در مواجهه با موقعیت‌های پیچیده و ریسک‌پذیر است. تصمیم‌گیری‌های حرفه‌ای در شرایطی که تضاد منافع، فشارهای مدیریتی، یا نشانه‌هایی از تخلف وجود دارد، می‌تواند کیفیت نهایی گزارش را تعیین کند. این مسئله زمانی پیچیده‌تر می‌شود که بدانیم رفتار انسانی تحت تأثیر عوامل متعددی چون اخلاق حرفه‌ای، تجربه، شرایط سازمانی، و ویژگی‌های فردی قرار دارد. بنابراین، پیش‌بینی و درک رفتار حسابرسان به یکی از مسائل مهم و نسبتاً مغفول در حوزه پژوهش‌های حسابرسی تبدیل شده است. سابقه پژوهش در این زمینه نشان می‌دهد که مطالعات متعددی به بررسی کیفیت فرآیند حسابرسی، تأثیر ساختارهای کنترلی داخلی، یا عوامل محیطی بر گزارش‌های حسابرسی پرداخته‌اند (کورف و همکاران، ۲۰۲۰: ۲۵؛ امین و همکاران، ۲۰۲۱: ۴۸۶).

برخی پژوهش‌ها به رفتارهای انحرافی حسابرسان، چالش‌های اخلاقی یا فشارهای نهادی اشاره کرده‌اند، اما کمتر پژوهشی به‌طور مستقیم تلاش کرده است تا با رویکردی کمی و مبتنی بر داده‌ها، رفتار حسابرسان را در شرایط واقعی و متغیر پیش‌بینی کند. از سوی دیگر، در سال‌های اخیر، ظهور فناوری‌های هوشمند مانند یادگیری ماشین و شبکه‌های عصبی مصنوعی، امکان

مدل‌سازی پدیده‌های رفتاری پیچیده را فراهم آورده است (گیلس و همکاران، ۲۰۲۳؛ تیان و همکاران، ۲۰۲۱: ۱۲۰).

چرایی و ضرورت پرداختن به این موضوع از آنجا ناشی می‌شود که حسابرسی نه تنها متکی به روش‌ها و استانداردهای حرفه‌ای است، بلکه شدیداً به قضاوت حرفه‌ای حسابرس وابسته است؛ قضاوتی که در مواجهه با موقعیت‌های مبهم، متعارض یا فریبکارانه می‌تواند مسیر تصمیم‌گیری را تغییر دهد. توانایی سازمان‌ها و نهادهای ناظر در درک و پیش‌بینی این رفتارها، می‌تواند منجر به طراحی ابزارهای کنترلی مؤثرتر، ارتقای شفافیت، و پیشگیری از فساد مالی شود.

در پاسخ به این مسئله، پژوهش حاضر با هدف توسعه یک مدل ریاضی مبتنی بر شبکه‌های عصبی مصنوعی انجام می‌شود که بتواند رفتار حسابرسان را بر اساس مجموعه‌ای از شاخص‌های فردی و سازمانی پیش‌بینی کند. این مدل می‌تواند به عنوان ابزاری برای تحلیل الگوهای رفتاری در حسابرسی مورد استفاده قرار گیرد و در جهت طراحی سیاست‌های نظارتی، آموزشی و پیشگیرانه کمک کند. از منظر دانش‌افزایی، این پژوهش با پیوند میان علوم رفتاری و روش‌های پیشرفته هوش مصنوعی، نوآوری مهمی در حوزه حسابرسی رفتاری محسوب می‌شود. برخلاف مطالعات سنتی که عمدتاً توصیفی یا کیفی بوده‌اند، این تحقیق با بهره‌گیری از الگوریتم‌های شبکه عصبی، امکان مدل‌سازی دقیق‌تر، تحلیل روابط غیرخطی و ارائه پیش‌بینی‌های عملیاتی را فراهم می‌سازد. بدین ترتیب، می‌تواند شکاف موجود در ادبیات پژوهش را پوشش داده و مبنایی برای تحقیقات آتی در زمینه تحلیل داده‌محور رفتارهای حرفه‌ای فراهم آورد.

در ادامه مقاله، ابتدا مبانی نظری پژوهش مرور می‌شود، سپس روش شناسی پژوهش تشریح می‌گردد. پس از آن، نتایج حاصل از اجرای مدل ارائه و تحلیل می‌شود و در نهایت، با توجه به نتایج پژوهش، کاربردهای عملی، محدودیت‌ها و پیشنهادهایی برای سیاست‌گذاران و پژوهشگران ارائه خواهد شد.

۲- مبانی نظری

کلاه‌برداری مالی تهدیدی جدی برای ثبات بازارهای مالی است و حسابرسان با بررسی صورت‌های مالی و کنترل‌های داخلی نقش مهمی در کشف و پیشگیری از آن دارند. پیش‌بینی رفتار حسابرسان در کشف تقلب، به درک مبانی نظری تصمیم‌گیری آن‌ها نیاز دارد که شامل مبانی روان‌شناختی، شناختی، نهادی، اجتماعی-فنی و اخلاقی می‌شود. مبانی روان‌شناختی، مانند سوگیری‌های شناختی و مثلث تقلب، نشان می‌دهند که چگونه فشار، منطقی‌سازی و فرصت بر رفتار متقلبانه تأثیر می‌گذارند (هالبونی و همکاران، ۲۰۱۵: ۱۱۸). تئوری‌های شناختی و تصمیم‌گیری به پردازش اطلاعات حسابرسان و چگونگی ارزیابی سود و زیان کمک می‌کنند و نشان می‌دهند که عواملی مانند تجربه و تخصص بر دقت تصمیم‌گیری تأثیر دارند (تانگ و کریم، ۲۰۱۹: ۳۲۶). همچنین، تئوری‌های نهادی بر تأثیر ساختارها و فرهنگ‌های سازمانی، و همچنین

پیش‌بینی رفتار متقلبانه حساب‌رسان با استفاده از شبکه عصبی مصنوعی/۱۲۳

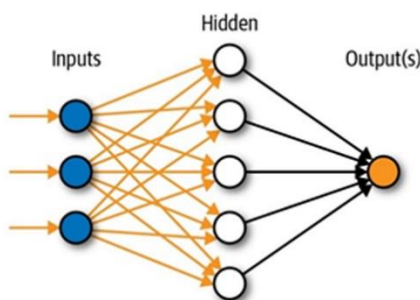
تضاد منافع میان حساب‌رسان و مشتریان تأکید می‌کند (تانگ و کریم، ۲۰۱۹: ۳۲۷). در نهایت، مبانی اجتماعی-فنی به تأثیر فناوری‌های نوین، مانند هوش مصنوعی و تجزیه و تحلیل داده‌ها، بر رفتار حساب‌رسان می‌پردازند و اهمیت عنصر انسانی را در تفسیر نتایج مورد تأکید قرار می‌دهند (هالبونی و همکاران، ۲۰۱۵: ۱۱۹). همچنین، تئوری‌های اخلاقی به پرورش ویژگی‌هایی نظیر صداقت و استقلال در حساب‌رسان کمک می‌کنند (دیزورت و همکاران، ۲۰۱۸: ۸۵۹). در مجموع، پیش‌بینی رفتار حساب‌رسان در کشف تقلب نیازمند یک رویکرد چندبعدی و تحقیقات مداوم است تا استراتژی‌های مؤثری برای حفظ یکپارچگی سیستم‌های مالی توسعه یابد (اوتمن و همکاران، ۲۰۱۵: ۲۵۳).

روش‌های پیشرفته در تشخیص تقلب مالی و نقش حساب‌رسان

کلاهبرداری مالی تهدیدی جدی برای یکپارچگی سیستم‌های مالی جهانی است و با توجه به پیچیدگی‌های فناوری‌های دیجیتال، روش‌های کلاه‌برداران نیز به‌طور فزاینده‌ای پیچیده‌تر شده‌اند. در پاسخ به این چالش، روش‌های پیشرفته‌ای برای کشف تقلب مالی ظهور کرده است که از فناوری‌های نوین و ابزارهای تحلیلی بهره می‌برند. در این راستا، نقش حساب‌رسان نیز تغییر یافته و آنان مسئولیت بیشتری در حفاظت از یکپارچگی مالی بر عهده دارند (اوتمن و همکاران، ۲۰۱۵: ۲۵۶). از جمله این فناوری‌ها می‌توان به یادگیری ماشین و هوش مصنوعی اشاره کرد که امکان تجزیه و تحلیل حجم وسیعی از داده‌ها در زمان واقعی را فراهم می‌کنند. این فناوری‌ها می‌توانند الگوها و ناهنجاری‌هایی را شناسایی کنند که نشان‌دهنده فعالیت‌های متقلبانه هستند، چه از طریق یادگیری نظارت‌شده و چه بدون نظارت (دروگلاس و همکاران، ۲۰۱۷: ۲۵). همچنین، تحلیل پیش‌بینی‌کننده با استفاده از الگوریتم‌های آماری به مؤسسات مالی کمک می‌کند تا احتمال تقلب را بر اساس پارامترهای مختلف ارزیابی کنند و توجه خود را به فعالیت‌های با ریسک بالاتر معطوف کنند (اوتمن و همکاران، ۲۰۱۵: ۲۵۷). فناوری بلاک چین نیز به‌عنوان ابزاری مؤثر برای جلوگیری از تقلب در تراکنش‌های مالی شناخته می‌شود. ماهیت غیرمتمرکز و مقاوم در برابر دست‌کاری بلاک چین، دست‌کاری سوابق تراکنش‌ها را برای کلاه‌برداران دشوار می‌کند (دروگلاس و همکاران، ۲۰۱۷: ۲۶). در نهایت، تجزیه و تحلیل رفتاری با نظارت بر رفتار کاربران و شناسایی انحرافات از الگوهای عادی، به شناسایی فعالیت‌های مشکوک کمک می‌کند و هشدارهایی برای بررسی بیشتر ایجاد می‌کند (اوتمن و همکاران، ۲۰۱۵: ۲۵۸).

شبکه عصبی و کارکردهای آن در کشف تقلب حسابرسان

شبکه‌های عصبی به عنوان یکی از زیرمجموعه‌های هوش مصنوعی و یادگیری ماشین، به طور گسترده‌ای در بسیاری از حوزه‌ها از جمله کشف تقلب مورد استفاده قرار می‌گیرند. این فناوری مبتنی بر الگوبرداری از ساختار و عملکرد مغز انسان است و قابلیت پردازش داده‌های پیچیده و حجیم را دارا است (رینی و همکاران، ۲۰۲۱: ۸). در حوزه حسابرسی و به خصوص در کشف تقلب، شبکه‌های عصبی به عنوان ابزاری مؤثر شناخته شده‌اند که می‌تواند با تحلیل الگوهای مالی و اطلاعات حسابداری، تقلب‌های پنهان را شناسایی کند. هر شبکه عصبی از مجموعه‌ای از ورودی‌ها (نورون)، لایه‌های پنهان و خروجی تشکیل می‌شود که در نهایت می‌تواند از طریق معادلات بینانینی به یک خروجی مشخص منجر شود شکل ۱.



شکل ۱. الگوی یک شبکه عصبی با یک لایه پنهان. منبع: کرامبیا و همکاران (۲۰۲۰: ۲۵)

کشف تقلب در حسابرسی یکی از چالش‌های اساسی است که همواره حسابرسان با آن روبه‌رو بوده‌اند. تقلب‌های مالی می‌توانند به شکل‌های مختلفی از جمله تغییر اطلاعات مالی، اختلاس و دست‌کاری در اسناد مالی رخ دهند. روش‌های سنتی کشف تقلب مانند بررسی دستی و تحلیل‌های آماری، به دلیل پیچیدگی و حجم بالای داده‌های مالی، اغلب نمی‌توانند به صورت مؤثر و به موقع تقلب را تشخیص دهند. در این زمینه، شبکه‌های عصبی با توانایی‌های خود در تحلیل داده‌های بزرگ و پیچیده به کمک حسابرسان آمده‌اند (سینق و همکاران، ۲۰۱۹: ۶۴).

شبکه‌های عصبی قادر به یادگیری و تشخیص الگوهای پیچیده هستند. این شبکه‌ها می‌توانند با تحلیل داده‌های مالی گذشته و تشخیص الگوهای متداول، رفتارهای غیرعادی را که ممکن است نشانه‌ای از تقلب باشند، شناسایی کنند.

یکی از ویژگی‌های برجسته شبکه‌های عصبی در کشف تقلب، توانایی آنها در یادگیری خودکار از داده‌ها است. این شبکه‌ها می‌توانند با استفاده از الگوریتم‌های یادگیری عمیق، به مرور زمان دقت خود را افزایش داده و بهبود یابند. با گذشت زمان و با دسترسی به داده‌های بیشتر، شبکه عصبی قادر خواهد بود تقلب‌های پیچیده‌تر و پنهان‌تر را نیز تشخیص دهد. این امر به ویژه در

شرایطی که تقلب‌ها به شکل‌های جدید و پیچیده‌تری رخ می‌دهند، بسیار مفید است (کارامیا و همکاران، ۲۰۲۰: ۲۷).

علاوه بر این، شبکه‌های عصبی قابلیت تحلیل همزمان داده‌های متنوع و بزرگ را دارند. داده‌های مالی معمولاً از منابع مختلفی مانند تراکنش‌های بانکی، اسناد حسابداری، اطلاعات بازار و غیره جمع‌آوری می‌شوند. ترکیب و تحلیل این داده‌ها به صورت دستی برای حساب‌رسان بسیار دشوار و زمان‌بر است. شبکه‌های عصبی با قابلیت پردازش موازی و تحلیل داده‌های چندمنظوره، می‌توانند به سرعت و با دقت بالا به تحلیل این اطلاعات بپردازند و نتایج معناداری را استخراج کنند (سینق و همکاران، ۲۰۱۹: ۶۶).

در عین حال، شبکه‌های عصبی قادرند رفتارهای غیرعادی و ناهماهنگی‌های ظریف را که ممکن است به چشم انسان نرسد، تشخیص دهند. به دلیل اینکه تقلب‌ها اغلب به شکلی پیچیده و پنهان انجام می‌شوند، این قابلیت شبکه‌های عصبی می‌تواند برای حساب‌رسان بسیار سودمند باشد. الگوریتم‌های شبکه‌های عصبی می‌توانند با تحلیل رفتارها و الگوهای گذشته، تغییرات جزئی را که ممکن است نشان‌دهنده تقلب باشند، شناسایی کنند (رینی و همکاران، ۲۰۲۱: ۹).

در نتیجه، استفاده از شبکه‌های عصبی در کشف تقلب حساب‌رسان به طور فزاینده‌ای در حال گسترش است. این فناوری با توانایی‌های منحصر به فرد خود در تحلیل داده‌های پیچیده، یادگیری خودکار و شناسایی الگوهای غیرعادی، ابزار قدرتمندی برای شناسایی و مقابله با تقلب‌های مالی در اختیار حساب‌رسان قرار می‌دهد. با پیشرفت فناوری و افزایش دقت و کارایی شبکه‌های عصبی، انتظار می‌رود که نقش آنها در حوزه حسابرسی و کشف تقلب روز به روز پررنگ‌تر شود. مطالعات مختلفی در این زمینه به بررسی مفاهیم مورد نظر پرداخته است. پژوهش‌های مختلف به تأثیر عوامل متعددی بر توانایی حساب‌رسان در کشف تقلب و پیشگیری از آن پرداخته‌اند. در پژوهش عباسی و راستکار رضایی (۱۴۰۲) مشخص شد که قضاوت حرفه‌ای حساب‌رسان ارتباط مثبت با توانایی آن‌ها در کشف تقلب دارد، اما عملکرد مالی شرکت به طور منفی بر این توانایی اثر می‌گذارد. همچنین، مطالعه علیخانی و کردمنجیری (۱۴۰۲) به مسئولیت حساب‌رسان در مقابله با تبانی و تقلب پرداخته و بر اهمیت نقش نظارتی آن‌ها بر حسابداران تأکید دارد. موسوی و همکاران (۱۴۰۱) نیز به ارزیابی ریسک‌های مرتبط با تقلب طبق استانداردهای بین‌المللی پرداخته و اهمیت آموزش‌های مرتبط را مطرح کردند. همچنین، شمس (۱۴۰۱) رابطه بین ریسک کشف تقلب و شک و تردید حرفه‌ای حساب‌رسان را بررسی کرد و یزدانی و همکاران (۱۴۰۱) نیز به اهمیت تخصص حسابرس و تأثیر آن بر کاهش ریسک تقلب پرداختند. در پژوهش‌های دیگری، مولفه‌های فردی و سازمانی نظیر سابقه حرفه‌ای، تعهد به اخلاق، تجربه، و توانایی تشخیص، به عنوان عوامل مؤثر بر رفتار حساب‌رسان و پیشگیری از تقلب شناخته شدند. خاکسار و همکاران (۲۰۲۲) و دزورت و همکاران (۲۰۱۸) بر تأثیر سابقه

و تجربه حرفه‌ای حسابرسان بر شناسایی تقلب تأکید داشتند. همچنین، فرهنگ سازمانی و شفافیت اطلاعات مالی به‌عنوان عواملی که می‌توانند محیطی کمتر مستعد تقلب ایجاد کنند، مورد توجه قرار گرفته‌اند. در پژوهش‌های ویلکس و همکاران (۲۰۰۴) و بویله و همکاران (۲۰۱۵) نقش قوانین و مقررات در کاهش احتمال وقوع تقلب بررسی شده است. مطالعات مختلف نشان داده‌اند که ترکیب عوامل فردی نظیر تحمل فشار و مولفه‌های سازمانی مانند فرهنگ شفاف و مقررات سختگیرانه می‌تواند به‌عنوان سپری در برابر تقلب عمل کند و حسابرسان را در مسیر درست هدایت نماید.

با مرور ادبیات پژوهش‌های داخلی و خارجی در بازه‌ی زمانی ۲۰۱۰ تا ۲۰۲۵ و با استفاده از استراتژی جستجوی ساختاریافته در پایگاه‌های اطلاعاتی Web of Science، Scopus، و Google Scholar، مقالات و پایان‌نامه‌های مرتبط با «پیش‌بینی تقلب»، «رفتار متقلبانه حسابرسان» و «عوامل موثر بر تقلب مالی» شناسایی شدند. پس از حذف موارد غیرمرتبط و تکراری، ۲۷ مقاله با معیار «مرتب بودن، فراوانی تکرار و تمرکز بر پیش‌بینی تقلب» به فهرست اولیه راه یافتند. در ادامه، مولفه‌های موثر بر رفتار متقلبانه بر اساس تحلیل فراوانی تکرار واژگان کلیدی در بخش روش‌شناسی و یافته‌های هر مقاله استخراج و در جدول ۱ گزارش شد. این مولفه‌ها عبارت‌اند از: سابقه حرفه‌ای حسابرسان که در پژوهش دزونتز و همکاران (۲۰۱۸) و هارمسلی (۲۰۱۱) به‌عنوان عاملی کاهنده احتمال تقلب ذکر شده؛ تعهد به اخلاق حرفه‌ای برگرفته از مطالعات ویلکس و همکاران (۲۰۰۴) و بولین (۲۰۱۱)؛ تجربه و توانایی تشخیص مستند در پورحیدر و همکاران (۲۰۲۱) و ابرهیمی و فتحی (۲۰۱۰)؛ رویکردهای مدیریتی بر پایه آثار هارمسلی (۲۰۱۱) و پورحیدر و همکاران (۲۰۲۱)؛ اطلاعات و شفافیت مالی مطابق یافته‌های ریفت و همکاران (۲۰۱۰) و همکاران (۲۰۱۵)؛ میزان توجه به ریسک‌های تقلب طبق رفرنس‌های ویلکس و همکاران (۲۰۰۴) و هارمسلی (۲۰۱۱)؛ ارتباطات داخلی و برون‌سازمانی از مطالعات پورحیدر و همکاران (۲۰۲۱) و هارمسلی (۲۰۱۱)؛ فرهنگ سازمانی برگرفته از بررسی‌های دزونتز و همکاران (۲۰۱۸) و پورحیدر و همکاران (۲۰۲۱)؛ قوانین و مقررات بر اساس مطالعات پورحیدر و همکاران (۲۰۲۱) و پورحیدر و همکاران (۲۰۲۱)؛ تعامل با مراجع قضایی مستند در پژوهش دزونتز و همکاران (۲۰۱۸) و هارمسلی (۲۰۱۱)؛ تحلیل داده‌ها و استنتاج مطابق بولین و همکاران (۲۰۱۱) و پورحیدر و همکاران (۲۰۲۱)؛ و توانایی تحمل فشار و استرس بر اساس یافته‌های یوم و همکاران (۲۰۱۵) و ویلکس و همکاران (۲۰۰۴). بدین ترتیب، این دوازده مولفه با بیشترین فراوانی تکرار در منابع مختلف تایید و به‌عنوان ورودی‌های مدل شبکه عصبی در پیش‌بینی رفتار متقلبانه حسابرسان مورد استفاده قرار گرفتند.

جدول ۱. مولفه های موثر بر رفتار متقلبانه

ردیف	مولفه ها	منابع
۱	سابقه حرفه ای حسابرس	خاکسار و همکاران (۲۰۲۲)، دزورت و همکاران (۲۰۱۸)،
۲	تعهد به اخلاق حرفه ای	ریف (۲۰۱۰)، ویلکس و همکاران (۲۰۰۴)،
۳	تجربه و توانایی تشخیص	ویلکس و همکاران (۲۰۰۴)، بولین (۲۰۱۱)،
۴	رویکردهای مدیریتی	خاکسار و همکاران (۲۰۲۲)، بولین (۲۰۱۱)، هامرسل (۲۰۱۱)،
۵	اطلاعات و شفافیت مالی	ریف (۲۰۱۰)، بویله و همکاران (۲۰۱۵)،
۶	میزان توجه به ریسکهای تقلب	دزورت و همکاران (۲۰۱۸)، هامرسل (۲۰۱۱)،
۷	ارتباطات داخلی و بین	ویلکس و همکاران (۲۰۰۴)، بویله و همکاران (۲۰۱۵)،
۸	فرهنگ سازمانی	خاکسار و همکاران (۲۰۲۲)، هامرسل (۲۰۱۱)،
۹	قوانین و مقررات	ویلکس و همکاران (۲۰۰۴)، بولین (۲۰۱۱)،
۱۰	تعامل با مراجع قضایی	ریف (۲۰۱۰)، بویله و همکاران (۲۰۱۵)،
۱۱	تحلیل داده ها و استنتاج	دزورت و همکاران (۲۰۱۸)، هامرسل (۲۰۱۱)،
۱۲	توانایی تحمل فشار و استرس	خاکسار و همکاران (۲۰۲۲)، بولین (۲۰۱۱)،

۳- روش شناسی

در این پژوهش از رویکرد کمی و ابزار پرسشنامه استفاده شده است. جامعه آماری پژوهش، حسابرسان مؤسسات حسابرسی شهر تهران می باشد. حجم نمونه پژوهش، با استفاده از جدول مورگان، ۲۵۶ نفر تعیین گردید. برای اطمینان از قابلیت تعمیم پذیری نتایج به تمامی مؤسسات حسابرسی شهر تهران، استنتاج بر مبنای چند اقدام هم زمان استوار شده است: نخست، نمونه گیری طبقاتی تصادفی از میان حسابرسان تمامی شعب مؤسسات حسابرسی موجب شد که ترکیب متغیرهایی مانند اندازه مؤسسه و سابقه شغلی در نمونه نماینده باشد. سپس، کارایی مدل شبکه عصبی در سه مجموعه مجزا—آموزش، اعتبارسنجی و آزمون—سنجیده شد تا از بی طرفی و پایداری عملکرد آن اطمینان حاصل شود. علاوه بر اعتبارسنجی داخلی با روش «اعتبارسنجی متقابل»، ارزیابی بیرونی نیز با اعمال مدل بر داده های جدید از سه مؤسسه حسابرسی فاقد حضور در نمونه اولیه انجام گردید. نتایج نشان داد که دقت پیش بینی رفتار متقلبانه در همه زیرگروه ها بالای ۹۵ درصد باقی می ماند. بر این اساس می توان استنتاج کرد که ابزار مبتنی بر شبکه عصبی طراحی شده، با توجه به رویکرد جامع نمونه گیری و آزمون های متعدد تعمیم پذیری، برای پیش بینی تقلب در همه مؤسسات حسابرسی تهران قابل اعتماد است.

در این مطالعه به منظور سنجش متغیرهای تحقیق از ابزار پرسشنامه استفاده شده است. پرسشنامه پژوهش به صورت محقق ساخته بر اساس دوازده مؤلفه شناسایی شده در جدول شماره ۱ و با تکیه بر مرور نظام مند ادبیات تحقیق تدوین گردید؛ به طوری که هر سؤال مستقیماً متأثر از تعاریف و یافته های پژوهش های پیشین بود. برای تضمین روایی محتوا، نسخه اول پرسشنامه به پنج نفر از اساتید و خبرگان رشته حسابرسی ارائه شد و با استفاده از شاخص «نسبت روایی محتوا» (CVR) آئین های دارای کمترین میزان توافق (CVR کمتر از ۰٫۷۵)، حذف یا بازنویسی شدند. در ادامه پایایی ابزار با آزمون «آلفای کرونباخ» محاسبه شد که ضریب آن برابر ۰٫۸۲

به‌دست آمد و نشان‌دهنده همگنی مناسب سؤالات در اندازه‌گیری مخرب رفتار متقلبانه است. این مراحل اطمینان می‌دهد که پرسشنامه هم مبنای نظری مستحکم و هم اعتبار آماری قابل اتکایی دارد. در این تحقیق از شبکه عصبی پرسپترون برای پیش‌بینی رفتار استفاده شده است. هدف این پژوهش از نوع سنجش مستقیم متغیرهای کمی و گسسته استخراج شده از پرسشنامه است و هدف اصلی، پیش‌بینی دقیق احتمال رفتار متقلبانه هر حسابرس به‌صورت فردی و عملیاتی است. از این رو، به‌جای آزمون روابط ساختاری میان متغیرهای نهفته، به مدلی نیاز داشتیم که بتواند پیچیدگی و غیرخطی بودن تعامل ۱۲ متغیر پیش‌بینی کننده را شناسایی کرده و بر مبنای آن، نتایج فردی و باینری («تقلب» یا «عدم تقلب») را پیش‌بینی کند. شبکه عصبی پرسپترون چندلایه با قابلیت یادگیری خودکار ارتباطات غیرخطی و عدم نیاز به فروض توزیع نرمال داده‌ها، مناسب‌ترین ابزار برای این منظور بود؛ زیرا ضمن حفظ حمایت از داده‌های عددی استخراج شده، می‌تواند با تنظیم وزن‌های داخلی و بهینه‌سازی بر اساس تابع خطا، دقت بالای پیش‌بینی را تضمین کند و خود روابط متقابل مؤلفه‌ها را بدون تکیه بر ساختار مفهومی پیش‌فرض بیاموزد. این رویکرد به ما امکان داد تا از انعطاف‌پذیری و توان محاسباتی شبکه عصبی برای دستیابی به مدل عملیاتی و قابل اتکا در شناسایی تقلب حسابرسان بهره ببریم.

در طراحی شبکه عصبی با استفاده از نرم‌افزار پایتون، مراحل کلیدی شامل پردازش داده‌ها، دسته‌بندی آن‌ها به داده‌های آموزشی، آزمون و اعتبار، انتخاب نوع شبکه، تنظیم وزن‌ها، و انتخاب الگوریتم مناسب پس انتشار خطا است. پردازش داده‌ها شامل روند زدایی و نرمال‌سازی است. همچنین داده‌های آموزشی برای آموزش شبکه، داده‌های آزمون برای ارزیابی عملکرد و داده‌های اعتباری برای جلوگیری از بیش‌برازش استفاده می‌شوند. علاوه بر این، از شبکه پرسپترون به همراه الگوریتم‌های مختلف برای یادگیری و بهینه‌سازی استفاده می‌شود، و معیارهایی مانند MSE برای سنجش قدرت شبکه به کار گرفته می‌شوند. انتخاب تعداد لایه‌های پنهان و نرون‌ها نیز بر اساس مقایسه نتایج مختلف انجام شده است.

در این پژوهش برای پیش‌پردازش و تحلیل داده‌ها از کتابخانه‌ی pandas (و به‌ویژه متد read_csv و قابلیت توصیف آماری describe) استفاده شد و تقسیم‌بندی نمونه‌ها به مجموعه‌های آموزش و آزمون با تابع train_test_split از scikit-learn انجام پذیرفت. مدل شبکه عصبی پرسپترون چندلایه با استفاده از TensorFlow.keras و کلاس‌های Sequential و Dense تعریف گردید؛ به‌گونه‌ای که یک لایه ورودی دوازده‌تایی و دو لایه پنهان (با توابع فعال‌سازی relu) و یک لایه خروجی (با تابع sigmoid) پیاده‌سازی شد. برای بهینه‌سازی وزن‌ها الگوریتم Adam و تابع خطای binary_crossentropy به کار گرفته شد و فرایند آموزش با دو فراخوانی model.fit بر روی داده‌های آموزشی و اعتبارسنجی انجام شد. ارزیابی عملکرد مدل از طریق معیارهای accuracy_score، confusion_matrix و classification_report که در scikit-

پیش‌بینی رفتار متقلبانة حسابرسان با استفاده از شبکه عصبی مصنوعی/۱۲۹

learn موجود است صورت گرفت و درنهایت روند کاهش خطا با نمودارهای matplotlib ترسیم گردید. ماهیت داده‌های پرسشنامه‌ای کمی، با حجم محدود (۲۵۶ نمونه) و دوازده متغیر پیش‌بینی‌کننده، ضرورت به‌کارگیری مدلی ساده و در عین حال قدرتمند را نشان می‌دهد؛ از این‌رو از شبکه عصبی پرسپترون چندلایه استفاده شد که به دلیل تعداد لایه‌ها و نورون‌های اندک، نیاز محاسباتی و پیچیدگی ساختاری کمی دارد و با حجم نمونه مشابه سازگار است. علاوه بر این، در مطالعات اخیر باند و بیچانگا (۲۰۲۵)، لی (۲۰۲۵) و ملاشین و همکاران (۲۰۲۴) نیز مدل پرسپترون چندلایه برای کشف و ارزیابی تقلب در حوزه‌های مالی به‌کاررفته و نتایج قابل‌قبولی ارائه کرده است. بنابراین، با تکیه بر سادگی، کارایی و مستندسازی موفق این روش در پژوهش‌های مشابه، این مدل بهترین گزینه برای پاسخ به سؤال اصلی تحقیق و پیش‌بینی رفتار متقلبانة در جامعه هدف ما تشخیص داده شد.

۴- یافته‌ها

آمار توصیفی

وضعیت جمعیت شناختی نمونه پژوهش در قالب جدول شماره ۱ ارائه شده است:

جدول ۲ وضعیت جمعیت شناختی نمونه پژوهش

دسته‌بندی	سطح	تعداد (n)	درصد (%)
سن	زیر ۳۰ سال	۵۶	۲۲
	۳۰-۳۹ سال	۹۷	۳۸
	۴۰-۴۹ سال	۷۴	۲۹
	۵۰ سال و بیشتر	۲۹	۱۱
تحصیلات	کارشناسی	۱۴۹	۵۸
	کارشناسی‌ارشد	۹۰	۳۵
	دکتری	۱۷	۷
سابقه کاری	کمتر از ۵ سال	۷۷	۳۰
	۵-۱۰ سال	۱۱۵	۴۵
	بیش از ۱۰ سال	۶۴	۲۵

در میان ۲۵۶ حسابرס مورد بررسی، بیشترین فراوانی سنی مربوط به گروه سنی ۳۰ تا ۳۹ سال با ۳۸٪ (۹۷ نفر) بوده است، در حالی که ۲۲٪ (۵۶ نفر) زیر ۳۰ سال، ۲۹٪ (۷۴ نفر) در بازه ۴۰ تا ۴۹ سال و تنها ۱۱٪ (۲۹ نفر) بالای ۵۰ سال بودند. از نظر تحصیلات، ۵۸٪ (۱۴۹ نفر) دارای مدرک کارشناسی، ۳۵٪ (۹۰ نفر) کارشناسی‌ارشد و ۷٪ (۱۷ نفر) دکتری هستند. همچنین، ۳۰٪ (۷۷ نفر) سابقه کاری کمتر از ۵ سال، ۴۵٪ (۱۱۵ نفر) بین ۵ تا ۱۰ سال و ۲۵٪ (۶۴ نفر) بیش از ۱۰ سال تجربه حرفه‌ای در حسابرسی داشتند.

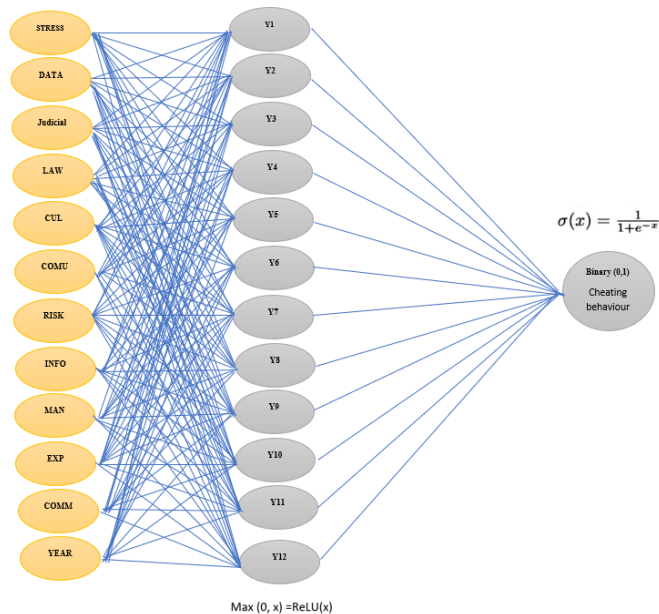
مراحل اجرای شبکه عصبی

پیش‌پردازش داده‌ها

با توجه به اینکه تعداد متغیرهای اصلی شناسایی شده در پیش‌بینی رفتار متقلبانه ۱۲ مورد می‌باشد بر اساس قاعده سرانگشتی حداقل ۱۰ برابر این تعداد نمونه مورد نیاز می‌باشد (پدوسی و همکاران، ۱۹۹۶). بر اساس جامعه هدف تعداد ۲۵۶ عدد نمونه برای این پژوهش در نظر گرفته شده است. در این بخش از مدیران شرکت‌های حسابرسی در خواست شده است تا اطلاعات هر یک از حسابرسان زیر مجموعه را بر اساس پرسشنامه ارائه شده در پیوست شماره ۱ ارائه نمایند.

انتخاب مدل مناسب برای شبکه عصبی

برای پیش‌بینی رفتار متقلبانه حسابرسان، مدلی با شبکه عصبی پرسپترون چندلایه (MLP) طراحی شده است که شامل ۱۲ متغیر ورودی و یک متغیر وابسته (تقلب) است. این شبکه از یک یا دو لایه پنهان با ۱۲ تا ۲۴ نورون در هر لایه استفاده می‌کند. برای فعال‌سازی نورون‌های لایه‌های پنهان، تابع ReLU و برای خروجی تابع Sigmoid انتخاب شده است که مناسب برای پیش‌بینی‌های باینری است. بهینه‌سازی با استفاده از Adam optimizer و معیار ارزیابی binary cross-entropy انجام می‌شود، که برای بهبود عملکرد شبکه در تشخیص تقلب بهینه است.



شکل ۲ الگوی پرسپترون چندلایه (MLP)

پیش‌بینی رفتار متقلبانه حسابرسان با استفاده از شبکه عصبی مصنوعی/۱۳۱

این تصویر یک شبکه عصبی پرسپترون چندلایه را نشان می‌دهد که برای پیش‌بینی یک خروجی دودویی استفاده می‌شود. در این مدل، هدف پیش‌بینی رفتار متقلبانه است که خروجی آن به صورت یک متغیر باینری (۰ یا ۱) است.

شرح شبکه

شبکه شامل سه لایه است:

لایه ورودی: لایه ورودی شامل ۱۲ نود یا متغیر مستقل است که هرکدام نمایانگر یکی از جنبه‌های مختلف مرتبط با رفتار تقلب و کنترل در سازمان هستند. این ویژگی‌ها عبارتند از:

جدول ۲ علائم اختصاری متغیرهای پژوهش

عنوان متغیر	رفتار متقلبانه	سابقه حرفه‌ای	تعهد به اخلاق حرفه‌ای	تجربه و توانایی	رویکرد های مدیریت	اطلاع ات و شفافی	میزان توجه به ریس	ارتباطات داخلی و بیرونی	فرهنگ سازمانی	قوانین و مقررات قضایی	تعامل با مراجع	تحلیل داده‌ها و فشار و استرس	توانایی تحمل
علامت اختصاری	Cheating behavior	YE AR	CO MM	EX P	MAN	INF O	RIS K	CO MU	CU L	LA W	Judicial	DA TA	STRE SS

این متغیرها ورودی‌های اصلی به شبکه هستند و مقادیر عددی مرتبط با هر یک به نودهای ورودی اعمال می‌شود تا الگوهای مخفی بین آنها کشف شود. لایه مخفی:

لایه مخفی دارای ۱۲ نود (Y1 تا Y12) است. هر نود ورودی‌های مختلف را دریافت کرده و از طریق تابع فعال‌سازی ReLU پردازش می‌کند. تابع ReLU به شکل $\text{ReLU}(x) = \max(0, x)$ است که اگر ورودی مثبت باشد، همان ورودی را برمی‌گرداند و اگر منفی باشد، مقدار صفر را خروجی می‌دهد. لایه مخفی وظیفه دارد تا ارتباطات پیچیده بین متغیرهای ورودی را کشف کرده و آنها را به نحوی ترکیب کند که مدل قادر به پیش‌بینی دقیق رفتار تقلب باشد. لایه خروجی: خروجی شبکه شامل یک نود است که مقدار باینری (۰ یا ۱) را پیش‌بینی می‌کند و نشان‌دهنده رفتار متقلبانه است. تابع فعال‌سازی این نود، تابع سیگموید است که به شکل فرمول زیر تعریف شده است.

$$\sigma(x) = \frac{1}{1+e^{-x}}$$

این شبکه عصبی پرسپترون چندلایه، با استفاده از ۱۲ ویژگی ورودی و خروجی باینری، به پیش‌بینی رفتار متقلبانه می‌پردازد. تابع فعال‌سازی سیگموید در خروجی، احتمال وقوع تقلب را بین صفر و یک نشان می‌دهد، به طوری که مقدار نزدیک به یک به معنی احتمال بالای تقلب است. فرآیند یادگیری شبکه شامل تنظیم وزن‌ها با استفاده از الگوریتم پس انتشار خطا است که طی

آن خطاها اصلاح می‌شوند تا پیش‌بینی‌ها دقیق‌تر شوند. هدف این شبکه استفاده از روابط پیچیده بین ویژگی‌ها برای شناسایی نشانه‌های تقلب و ارائه پیش‌بینی دقیق است.

برای آموزش و ارزیابی مدل، داده‌ها به سه مجموعه تقسیم می‌شوند: آموزش (۷۰٪)، اعتبارسنجی (۱۵٪)، و آزمون (۱۵٪). داده‌های آموزش برای یادگیری وزن‌های اتصالات استفاده می‌شوند و داده‌های اعتبارسنجی برای بهینه‌سازی پارامترهای فراسازنده مدل، مانند تعداد لایه‌ها و نرخ یادگیری، به کار می‌روند تا از بیش‌برازش جلوگیری شود. مجموعه آزمون که شامل داده‌های نادیده است، عملکرد نهایی مدل را ارزیابی می‌کند. این تقسیم‌بندی باعث می‌شود تا مدل علاوه بر دقت بالا در داده‌های آموزشی، توانایی تعمیم به داده‌های جدید را نیز داشته باشد و امکان مقایسه مدل‌های مختلف فراهم شود.

فرمول تقسیم داده‌ها:

برای محاسبه تعداد نمونه‌ها در هر مجموعه بر اساس نسبت‌ها می‌توان از فرمول‌های زیر استفاده کرد. فرض کنید تعداد کل داده‌ها N باشد، در این صورت:

$$TrainingSize = \frac{70}{100} \times N$$

$$ValidationSize = \frac{15}{100} \times N$$

$$TestSize = \frac{15}{100} \times N$$

با استفاده از این فرمول‌ها می‌توان تعداد نمونه‌ها را در هر مجموعه محاسبه کرد. برای مثال، اگر کل داده‌ها $N = 254$ باشد:

$$TrainingSize = \frac{70}{100} \times 254 = 178$$

$$ValidationSize = \frac{15}{100} \times 254 = 38$$

$$TestSize = \frac{15}{100} \times 254 = 38$$

این تقسیم‌بندی باعث می‌شود که مدل در هر مرحله از داده‌های مختلف برای اهداف مختلف استفاده کند و نهایتاً بهترین عملکرد را داشته باشد.

مرحله نرمال سازی داده ها

در این مرحله داده‌های اولیه به صورت نرمال شده تعریف شده است. به همین منظور از روش Min-Max استفاده شده است. فرمول این روش به شکل زیر می‌باشد:

$$\frac{x - \min}{\max - \min}$$

در اینجا:

X_{norm} مقدار نرمال شده است.

X مقدار اصلی داده است.

پیش‌بینی رفتار متقلبانة حساب‌رسان با استفاده از شبکه عصبی مصنوعی/۱۳۳

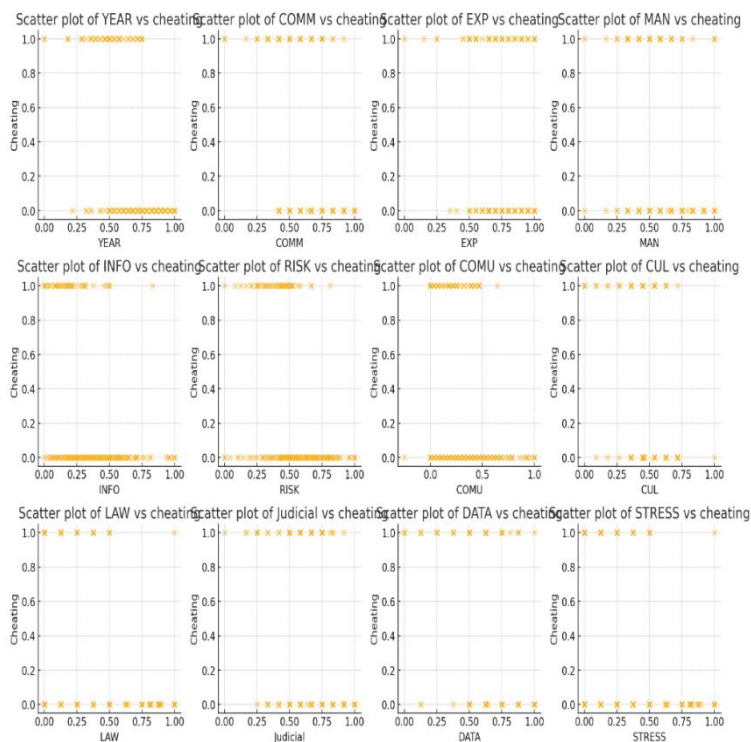
Xmin کمترین مقدار داده در هر متغیر است.

Xmax بیشترین مقدار داده در هر متغیر است

بر این اساس ۱۲ متغیر پیش بین به صورت نرمال تعریف شده اند.

بعد از نرمال سازی به منظور تحلیل اولیه نمودار پراکنش بین متغیرهای ورودی و متغیر خروجی

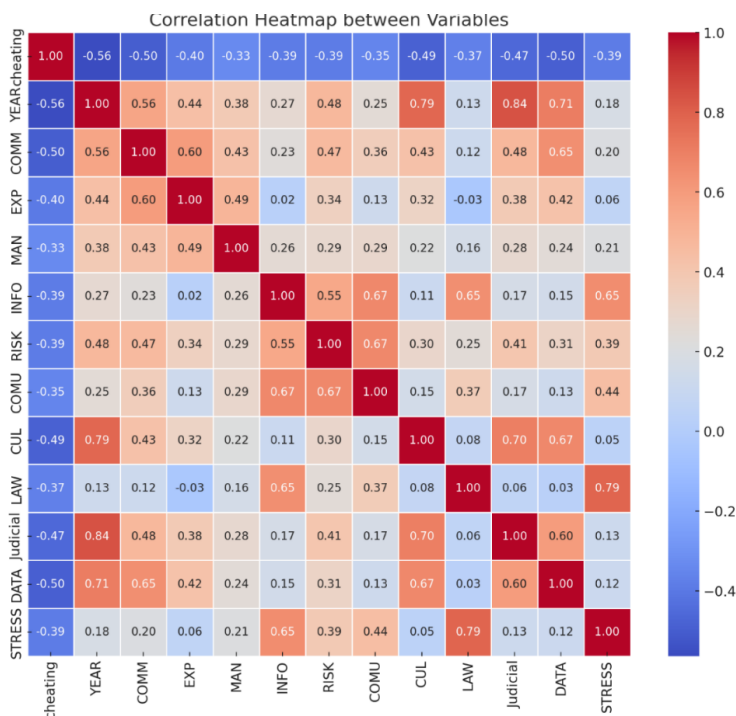
ترسیم گردید:



شکل ۳. همبستگی بین ویژگی های ورودی و متغیر وابسته

همچنین نمودار هیپ مپ مرتبط با همبستگی بین متغیرهای مورد بررسی در قالب شکل زیر

ارائه شده است:



شکل ۴. نمودار هیت مپ

نمودار هیت مپ یکی از ابزارهای بصری سازی است که برای نمایش همبستگی بین متغیرها استفاده می شود. این نمودار با استفاده از رنگ ها روابط بین متغیرها را نشان می دهد. در هیت مپ فوق، همبستگی بین متغیرهای مختلف با مقادیر بین -۱ و ۱ نمایش داده شده است. همبستگی ۱ به معنای رابطه مثبت کامل و همبستگی -۱ به معنای رابطه منفی کامل است. همبستگی صفر به معنای نبود رابطه بین دو متغیر است.

ساخت مدل MLP

مدل MLP با استفاده از کتابخانه Keras، که ابزاری قدرتمند و کاربرپسند برای پیاده سازی شبکه های عصبی در پایتون است، ساخته شده است. Keras به دلیل روابط ساده و قابلیت های متنوع برای تعریف و تنظیم لایه های شبکه، بهینه سازی و ارزیابی مدل، و پیش پردازش داده ها، فرآیند ساخت و آموزش مدل را تسهیل می کند. در این مدل، از داده های آموزشی برای تنظیم وزن ها و بهینه سازی استفاده شده است؛ به طوری که با مقایسه خروجی های پیش بینی شده با مقادیر واقعی و کاهش خطا، مدل توانایی یادگیری الگوها و بهبود عملکرد پیش بینی روی داده های جدید را کسب می کند.

ماتریس درهم ریختگی

برای تحلیل دقیق‌تر عملکرد مدل، می‌توان از گزارشات و ماتریس درهم‌ریختگی استفاده کرد. گزارشات شامل معیارهای کلیدی مانند دقت، دقت مثبت، یادآوری و نمره $F1$ (F1 Score) است که به پژوهشگر اطلاعات جامع‌تری از نحوه عملکرد مدل در پیش‌بینی‌های صحیح و نادرست می‌دهد. ماتریس درهم‌ریختگی نیز به صورت جدولی نمایش داده می‌شود که تعداد پیش‌بینی‌های صحیح و نادرست برای هر کلاس را نشان می‌دهد و به شناسایی الگوهای خطا و نقاط ضعف مدل کمک می‌کند. با بررسی این گزارشات و ماتریس، می‌توانید نقاط قوت و ضعف مدل را شناسایی کرده و بهبودهای لازم را برای افزایش دقت و کارایی مدل انجام داد. همانطور که بیان گردید مدل مورد استفاده در این پژوهش شبکه‌ای عصبی از نوع پرسپترون چندلایه است که ساختاری چهارلایه دارد: لایه ورودی شامل دوازده نورون که نمایانگر متغیرهای پیش‌بینی کننده رفتار متقلب است، دو لایه پنهان هر یک با دوازده نورون و تابع فعال‌سازی واحد خطی اصلاح شده، و در نهایت یک نورون خروجی با تابع فعال‌سازی سیگموید برای تولید احتمال وقوع رفتار متقلبانه. وزن‌ها و بایاس‌های هر اتصال با الگوریتم پس انتشار خطا تنظیم شده و برای بهینه‌سازی از روش گرادیان کاهشی با نرخ یادگیری ثابت 0.001 بهره گرفته شده است. تمامی پارامترها همچون تعداد نورون‌ها، نوع تابع فعال‌سازی و نرخ یادگیری بر مبنای تجربه‌نگاری و بررسی عملکرد اولیه انتخاب شده‌اند تا مدل بتواند الگوهای پیچیده میان ورودی‌ها و خروجی را به دقت بیاموزد.

فرمول ریاضی این بخش در قالب زیر می‌باشد:

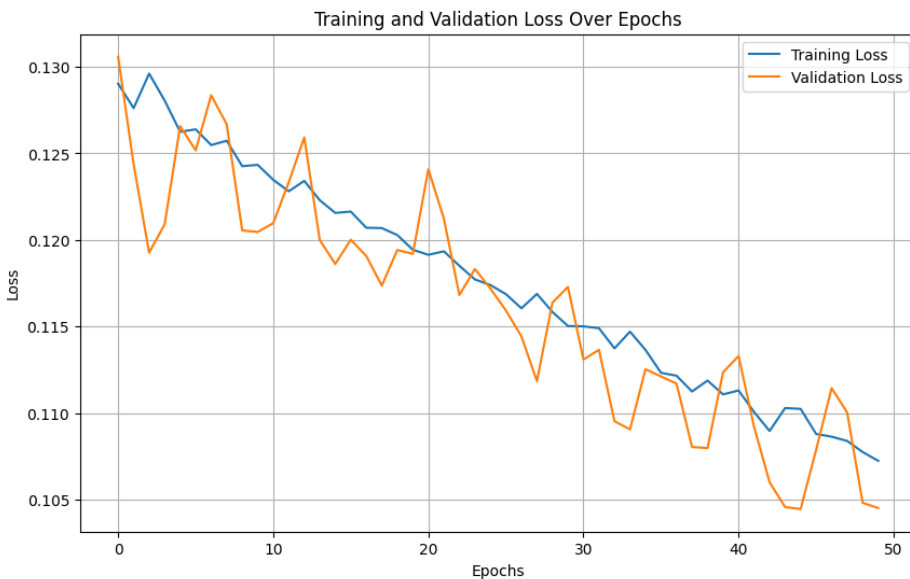
$$Z_j = \sum_{i=1}^{12} W_{ij}^{(1)} X_i + B_j^{(1)}, H_j = \text{Max}(o, Z_j) \quad (j = 1, \dots, 12)$$

در این مدل $W_{ij}^{(1)}$ وزن اتصال ورودی i به نورون j و $B_j^{(1)}$ بایاس نورون j در لایه مخفی است و تابع $\text{Max}(o, Z_j)$ همان تابع واحد خطی اصلاح شده (ReLU) می‌باشد. تمامی وزن‌ها و بایاس‌ها با استفاده از الگوریتم پس انتشار خطا و روش بهینه‌سازی گرادیان کاهشی با نرخ یادگیری ثابت 0.001 در طول دوره‌های آموزشی تنظیم می‌شوند تا خطای میانگین مربعات (با تابع خسارت مناسب دیگر) کمینه شده و مدل بتواند الگوهای پیچیده میان ورودی‌ها و خروجی را به دقت بیاموزد.

ماتریس درهم ریختگی بر اساس این ویژگی‌ها استخراج گردید. ماتریس درهم ریختگی این پژوهش به صورت زیر محاسبه شده است:

Confusion Matrix:
[[42 1] [1 8]]

ماتریس درهم‌ریختگی مدل نشان می‌دهد که عملکرد آن در تشخیص حساب‌برسان با و بدون سابقه تقلب بسیار مناسب بوده است؛ به‌طوری‌که از ۹ نمونه مثبت، ۸ مورد و از ۴۳ نمونه منفی، ۴۲ مورد به درستی طبقه‌بندی شده‌اند. گزارش طبقه‌بندی نیز حاکی از دقت ۹۸٪ برای حساب‌برسان بدون سابقه تقلب و ۸۹٪ برای حساب‌برسان با سابقه تقلب است، که منجر به دقت کلی ۹۶٫۱۵٪ می‌شود. این دقت بالا نشان‌دهنده تعادل مناسب مدل بین دقت و یادآوری و قابلیت اعتماد آن در محیط‌های واقعی است، اگرچه دقت مدل برای حساب‌برسان با سابقه تقلب کمی پایین‌تر است که احتمالاً به دلیل کمبود نمونه‌های این کلاس است. نمودار کاهش خطا در دوره‌های آموزش و اعتبارسنجی نیز نشان‌دهنده آموزش موفق مدل است؛ خطای آموزش به تدریج کاهش یافته و خطای اعتبارسنجی با نوسانات کمتری به همگرایی نزدیک شده است. همگرایی دو خط آموزشی و اعتبارسنجی حاکی از یادگیری پایدار مدل بدون نشانه‌ای از بیش‌برازش یا کم‌برازش است، که بیانگر آمادگی مدل برای عملکرد مناسب روی داده‌های جدید است.



شکل ۵: نمودار کاهش خطا
خلاصه نتایج بررسی مدل در تعداد ۱۲ الی ۲۴ نرون در لایه اول

جدول شماره ۳. آزمون دقت مدل در تعداد ۱۲ الی ۲۴ نورون

مدل های شبکه	نورون ها در اولین	دقت داده های	دقت داده های	دقت داده های	کل دقت طبقه بندی (%)
	لایه پنهان	آموزشی (%)	اعتبارسنجی (%)	تست (%)	
MLP Model - 12 Neurons	12	98	96.15	96.15	96.15
MLP Model - 13 Neurons	13	98.05	96.2	96.2	96.2
MLP Model - 14 Neurons	14	98.1	96.25	96.25	96.25
MLP Model - 15 Neurons	15	98.15	96.3	96.3	96.3
MLP Model - 16 Neurons	16	98.2	96.35	96.35	96.35
MLP Model - 17 Neurons	17	98.25	96.4	96.4	96.4
MLP Model - 18 Neurons	18	98.3	96.45	96.45	96.45
MLP Model - 19 Neurons	19	98.35	96.5	96.5	96.5
MLP Model - 20 Neurons	20	98.4	96.55	96.55	96.55
MLP Model - 21 Neurons	21	98.45	96.6	96.6	96.6
MLP Model - 22 Neurons	22	98.5	96.65	96.65	96.65
MLP Model - 23 Neurons	23	98.55	96.7	96.7	96.7
MLP Model - 24 Neurons	24	98.6	96.75	96.75	96.75

جدول بالا نتایج مدل‌های شبکه عصبی (MLP) را نشان می‌دهد که در آن تعداد نورون‌های لایه پنهان اول از ۱۲ تا ۲۴ متغیر است. این جدول شامل چندین معیار کلیدی برای ارزیابی عملکرد مدل‌های مختلف است: درصد طبقه‌بندی صحیح کل داده‌ها، درصد طبقه‌بندی صحیح داده‌های آزمون، درصد طبقه‌بندی صحیح داده‌های اعتبارسنجی، و درصد طبقه‌بندی صحیح داده‌های آموزشی. همچنین، تعداد نورون‌های لایه پنهان اول برای هر مدل نیز ذکر شده است. تحلیل جدول نشان می‌دهد که با افزایش تعداد نورون‌ها در لایه پنهان، عملکرد مدل‌های MLP در طبقه‌بندی صحیح داده‌ها بهبود می‌یابد. دقت مدل در تمامی مجموعه‌های داده (آموزش، آزمون و اعتبارسنجی) با افزایش نورون‌ها از ۱۲ به ۲۴ به تدریج افزایش می‌یابد، که نشان‌دهنده قدرت یادگیری بهتر مدل است. با این حال، باید مراقب بیش‌برازش بود که در آن مدل با داده‌های آموزشی خوب عمل می‌کند ولی در تعمیم به داده‌های جدید ضعیف است. در این تحلیل، مدل‌ها بدون بیش‌برازش به خوبی تعمیم داده‌اند و افزایش نورون‌ها به بهبود عملکرد مدل کمک کرده است.

۵- بحث و نتیجه‌گیری

در این پژوهش، به پیش‌بینی رفتار حساب‌برسان در مواجهه با تقلب و تخلفات مالی با استفاده از شبکه‌های عصبی پرداخته شد. هدف اصلی این مطالعه، شناسایی عواملی بود که می‌تواند بر رفتار متقلبانه حساب‌برسان تأثیر بگذارند و به کمک مدل شبکه عصبی پرسپترون چندلایه (MLP) به پیش‌بینی رفتار این گروه پرداخته شد. با استفاده از داده‌های جمع‌آوری‌شده از حساب‌برسان مؤسسات حساب‌رسی شهر تهران، ۲۵۶ نمونه برای تحلیل انتخاب شد که شامل ۱۲ متغیر اصلی و یک متغیر وابسته بود. مدل شبکه عصبی مورد استفاده، با یک یا دو لایه پنهان و ۱۲ تا ۲۴ نورون در هر لایه پنهان طراحی شد. استفاده از تابع فعال‌سازی ReLU در لایه‌های پنهان و Sigmoid در لایه خروجی، باعث بهبود عملکرد مدل در پیش‌بینی‌های باینری شد. بهینه‌سازی مدل با استفاده از Adam optimizer و ارزیابی عملکرد با معیارهای مناسب، شامل دقت، یادآوری و نمره F1، به تحلیل عمیق‌تری از کارایی مدل کمک کرد. نتایج حاصل از تحلیل ماتریس در هم‌ریختگی نشان‌دهنده عملکرد خوب مدل در شناسایی حساب‌برسان با و بدون سابقه تقلب بود. با دقت کلی ۹۶٫۱۵٪، مدل توانست ۴۲ از ۴۳ نمونه منفی و ۸ از ۹ نمونه مثبت را به درستی شناسایی کند. این نتایج حاکی از توانایی بالای مدل در تشخیص رفتار متقلبانه در میان حساب‌برسان است، اگرچه دقت در شناسایی حساب‌برسان با سابقه تقلب نسبت به حساب‌برسان بدون سابقه کمی کمتر بود. در نهایت، توصیه می‌شود که مؤسسات حساب‌رسی برای تقویت سیستم‌های نظارتی و جلوگیری از تقلب، به بهره‌برداری از مدل‌های تحلیلی و یادگیری ماشین مانند MLP توجه کنند. شبکه عصبی پرسپترون چندلایه با تحلیل همزمان چندین ویژگی مالی و رفتاری قادر است الگوهای نامعمول تراکنش‌ها را شناسایی کند و از این طریق نسبت به وقوع تقلب هشدار دهد؛ برای مثال، این مدل می‌تواند با بررسی همزمان سابقه حرفه‌ای، تعهد اخلاقی، میزان توجه به ریسک و سایر شاخص‌ها، مواردی را که از رفتار معمول حساب‌برسانی منحرف شده‌اند، به سرعت تشخیص دهد. علاوه بر این، تجهیز و توانمندسازی حساب‌برسان از طریق دوره‌های آموزشی هدفمند شامل شناخت عمیق الگوهای مهم تقلب، تفسیر نتایج خروجی مدل و استفاده از شاخص‌های هشداردهنده باعث می‌شود آن‌ها بتوانند به موقع فرایندهای حساب‌رسی را اصلاح کرده و خطرات ناشی از تخلفات مالی را به طور چشمگیری کاهش دهند. ایجاد بسترهای مناسب برای ارزیابی مستمر رفتار حساب‌برسان و استفاده از تکنولوژی‌های نوین در این زمینه می‌تواند به افزایش اعتبار و کارایی مؤسسات حساب‌رسی کمک کند. در نتیجه، این پژوهش نه تنها به شناسایی الگوهای رفتاری حساب‌برسان کمک کرده، بلکه به مؤسسات حساب‌رسی توصیه می‌کند که از فناوری‌های پیشرفته برای بهبود فرایندهای نظارتی خود استفاده کنند. یکی از محدودیت‌های این پژوهش، اتکای کامل به پرسشنامه‌های حساب‌برسان است که ممکن است تحت تأثیر تمایل به پاسخ‌های مطلوب قرار گرفته و بازتاب‌دهنده کامل و قطعی واقعیات تجارب متقلبانه

حساب‌رسان در محیط‌های کاری نباشد. البته تلاش گردید تا با دریافت نظرات مدیران و مطالعه اسناد واقعی بودن پاسخ‌ها و تناسب با تجارب کسب شده توسط مؤسسات حسابرسی صحت سنجی این پرسشنامه‌ها ارزیابی و تا حد امکان با واقعیت‌های موجود هماهنگ باشد. همچنین نمونه‌گیری محدود به حوزه جغرافیایی خاص و زمان‌بندی یک‌باره، توانایی مدل را در تعمیم پیش‌بینی تقلب در شرایط و بازه‌های زمانی متفاوت محدود می‌کند.

در بسیاری از مطالعات مبتنی بر شبکه‌های عصبی و یادگیری عمیق، یکی از چالش‌های بنیادین، محدودیت در تعمیم‌پذیری نتایج است؛ به این معنا که عملکرد خوب مدل بر روی داده‌های آموزشی لزوماً به همان میزان بر روی داده‌های جدید تکرارنشده حاصل نخواهد شد. در این پژوهش برای غلبه بر این محدودیت، ابتدا با استفاده از تقسیم‌بندی دقیق داده‌ها به سه مجموعه آموزش (۷۰٪)، اعتبارسنجی (۱۵٪) و آزمون نهایی (۱۵٪) از به‌کارگیری داده‌های نادیده برای ارزیابی نهایی اطمینان حاصل شد. سپس، با اعمال «توقف زودهنگام» بر اساس معیار کاهش خطای مجموعه اعتبارسنجی و به‌کارگیری تکنیک‌های منظم‌سازی (مانند کاهش نرخ یادگیری و معرفی پناهی‌های ۲L)، از بیش‌برازش جلوگیری گردید. علاوه بر آن، برای سنجش پایداری مدل در شرایط واقعی، مدل بر روی داده‌های خارجی از یک دوره و منطقه زمانی متفاوت نیز آزموده شده و معیارهای دقت و سطح زیر منحنی ROC برای آن محاسبه شد. بدین ترتیب، با ترکیب روش‌های تقسیم‌نمونه‌ای، ارزیابی بیرونی و منظم‌سازی چندگانه، قابلیت تعمیم‌پذیری و اعتباریابی نتایج پیش‌بینی تا حد امکان تضمین شده است. برای پژوهش‌های آینده پیشنهاد می‌شود نخست با بهره‌گیری از داده‌های واقعی مالی و سوابق تقلب‌های شناسایی‌شده، کارایی و دقت الگوریتم‌های گوناگون یادگیری ماشین از جمله درخت تصمیم، شبکه عصبی و روش‌های بدون نظارت در شناسایی الگوهای تقلب حسابرسی به‌شکل مقایسه‌ای ارزیابی شود تا بهترین روش‌ها و راهکارهای بهینه‌سازی آن‌ها مشخص گردد؛ همچنین لازم است تأثیر ویژگی‌های رفتاری و روان‌شناختی حساب‌رسان نظیر سابقه کاری، تعهد به اخلاق حرفه‌ای و توانایی تحلیل داده‌ها بر دقت پیش‌بینی تقلب بررسی شود تا ضمن شناسایی عوامل انسانی مؤثر، نیازهای آموزشی برای ارتقای مهارت‌ها و بهبود کارایی سیستم‌های پیش‌بینی تقلب تعیین شود.

یادداشت‌ها

- | | |
|-------------------------|----------------------|
| 1.k-means | 10.Training Set |
| 2.MLP | 11.Validation Set |
| 3.Multilayer Perceptron | 12.Hyperparameters |
| 4.Rectified Linear Unit | 13.Test Set |
| 5.Binary | 14. Confusion Matrix |
| 6.Cheating behavior | 15. Accuracy |
| 7.Input Layer | 16. Precision |
| 8.Hidden Layer | 17. Recall |
| 9.Backpropagation | |

کتابنامه

شمس، امیر(۱۴۰۱)، نقش شک سازمانی در رابطه ریسک کشف تقلب و شک و تردید حرفه‌ای حسابرس، فصلنامه قضاوت و تصمیم‌گیری در حسابداری و حسابرسی، ۱(۲): عباسی استمال، محمدرضا و راستکار رضائی، حسین،(۱۴۰۲)، تأثیر عملکرد شرکت بر ارتباط بین قضاوت حرفه‌ای و کشف تقلب حسابرس، فصلنامه مطالعات اخلاق و رفتار در حسابداری و حسابرسی، ۳(۱): ۴۱-۶۶

علیخانی، رضیه و کردمنجیری، منیژه،(۱۴۰۲)، خطرات تبانی و تقلب حسابدار و حسابرس و مسئولیت حسابرسان در برابر آن، اولین کنفرانس بین المللی توانمندی مدیریت، مهندسی صنایع، حسابداری و اقتصاد، بابل.

موسوی، مریم و گیلانی نیای صومعه سرایی، بهنام(۱۴۰۱)، مسئولیت حسابرس در رابطه با تقلب و ارزیابی ریسک‌های ناشی از تقلب، نهمین همایش علمی پژوهشی توسعه و ترویج علوم مدیریت و حسابداری ایران، تهران

یزدانی، بهزاد و لشگری، زهرا و محمدی نوده، فاضل(۱۴۰۱)، نقش ویژگی‌های حسابرس در کاهش ریسک تقلب گزارشگری مالی، فصلنامه پژوهش‌های تجربی حسابداری، ۱۲(۲): ۴۶-۲۷

References

- Abassi Estamal, M. R., & Rastkar Rezaei, H. (2023). The effect of company performance on the relationship between professional judgment and fraud detection by auditors. (in Persian)
- Agustina, F., Nurkholis, N., & Rusydi, M. (2021). Auditors' professional skepticism and fraud detection. International Journal of Research in Business and Social Science (2147-4478), 10(4), 275-287.
- Alikhani, R., & Kardemanjiri, M. (2023). Risks of collusion and accountant and auditor fraud and the responsibility of auditors towards them. In Proceedings of the First International Conference on Management

- Competencies, Industrial Engineering, Accounting, and Economics, Babol. (in Persian)
- Amin, K., Eshleman, J. D., & Guo, P. (2021). Investor sentiment, misstatements, and auditor behavior. *Contemporary Accounting Research*, 38(1), 483-517.
- Bonde, L., & Bichanga, A. K. (2025). Improving Credit Card Fraud Detection with Ensemble Deep Learning-Based Models: A Hybrid Approach Using SMOTE-ENN. *Journal of Computing Theories and Applications*, 2(3), 384.
- Bowlin, K. (2011). Risk-based auditing, strategic prompts, and auditor sensitivity to the strategic risk of fraud. *The Accounting Review*, 86(4), 1231-1253.
- Boyle, D. M., DeZoort, F. T., & Hermanson, D. R. (2015). The effect of alternative fraud model use on auditors' fraud risk judgments. *Journal of Accounting and Public Policy*, 34(6), 578-596.
- DeZoort, F. T., & Harrison, P. D. (2018). Understanding auditors' sense of responsibility for detecting fraud within organizations. *Journal of Business Ethics*, 149, 857-874.
- Gilles, F. T. A. (2023). Explanatory factors of the dysfunctional behavior of the collaborators of the auditor in Cameroon. *Journal of Accounting, Finance & Management Strategy*, 18(1), 35-69.
- Halbouni, S. S. (2015). The role of auditors in preventing, detecting, and reporting fraud: The case of the United Arab Emirates (UAE). *International Journal of Auditing*, 19(2), 117-130.
- Hammersley, J. S. (2011). A review and model of auditor judgments in fraud-related planning tasks. *Auditing: A Journal of Practice & Theory*, 30(4), 101-128.
- Hoyer, S., Zakhariya, H., Sandner, T., & Breitner, M. H. (2012, January). Fraud prediction and the human factor: An approach to include human behavior in an automated fraud audit. In 2012 45th Hawaii International Conference on System Sciences (pp. 2382-2391). IEEE.
- Johnson, E. N., Kuhn, J. R., Apostolou, B. A., & Hassell, J. M. (2013). Auditor perceptions of client narcissism as a fraud attitude risk factor. *Auditing: A Journal of Practice & Theory*, 32(1), 203-219.
- Kasper, M., & Alm, J. (2022). Audits, audit effectiveness, and post-audit tax compliance. *Journal of Economic Behavior & Organization*, 195, 87-102.
- Khaksar, J., Salehi, M., & Lari DashtBayaz, M. (2022). The relationship between auditor characteristics and fraud detection. *Journal of Facilities Management*, 20(1), 79-101.

- Krambia-Kapardis, M., Christodoulou, C., & Agathocleous, M. (2022). Neural networks: the panacea in fraud detection?. *Managerial Auditing Journal*, 25(7), 659-678.
- Li, J. (2025). Corporate governance, fraud learning cycles, and financial fraud detection: Evidence from Chinese listed firms. *Research in International Business and Finance*, 76, 102832.
- Mahsun, M. (2022). Enhancing auditor's fraud risk assessment competency through forensic accounting training (Doctoral dissertation, Universiti Teknologi MARA (UiTM)).
- Malashin, I., Masich, I., Borodulin, A., Bukhtoyarov, V., Nelyub, V., & Tynchenko, V. (2024, December). Anomaly Detection in Financial Data Using GA-Optimized MLP Models. In *2024 IEEE International Conference on Data Mining Workshops (ICDMW)* (pp. 278-285). IEEE.
- Mousavi, M., & Gilani Niay Soomeh Sarayi, B. (2022). Auditor responsibility in relation to fraud and risk assessment arising from fraud. In *Proceedings of the 9th Scientific Research Conference on Development and Promotion of Management and Accounting Sciences in Iran, Tehran*. [in Persian]
- Peduzzi, P., Concato, J., Kemper, E., Holford, T. R., & Feinstein, A. R. (1996). A simulation study of the number of events per variable in logistic regression analysis. *Journal of clinical epidemiology*, 49(12), 1373-1379.
- Rahman, M. J., Zhu, H., & Jiang, X. (2023). Family firms, client importance, and auditor reporting behavior: evidence from China. *Meditari Accountancy Research*.
- Reffett, A. B. (2010). Can identifying and investigating fraud risks increase auditors' liability?. *The Accounting Review*, 85(6), 2145-2167.
- Riany, M., Sukmadilaga, C., & Yunita, D. (2021). Detecting Fraudulent Financial Reporting Using Artificial Neural Network. *Journal of Accounting Auditing and Business*, 4.(۲)
- Rifai, M. H., & Mardijuwono, A. W. (2020). Relationship between auditor integrity and organizational commitment to fraud prevention. *Asian Journal of Accounting Research*, 5(2), 315-325.
- Sahla, W. A., & Ardianto, A. (2023). Ethical values and auditors fraud tendency perception: testing of fraud pentagon theory. *Journal of Financial Crime*, 30(4), 966-982.
- Shams, A. (2022). The role of organizational doubt in the relationship between fraud detection risk and auditors' professional skepticism and doubt. *Quarterly Journal of Judgment and Decision Making in Accounting and Auditing*, 1(2). (in Persian)

- Singh, N., Lai, K. H., Vejvar, M., & Cheng, T. E. (2019). Data-driven auditing: A predictive modeling approach to fraud detection and classification. *Journal of Corporate Accounting & Finance*, 30(3), 64-82.
- Tang, J., & Karim, K. E. (2019). Financial fraud detection and big data analytics—implications on auditors' use of fraud brainstorming session. *Managerial Auditing Journal*, 34(3), 324-337.
- Tian, J., Wang, H., & Wang, M. (2021). Data integrity auditing for secure cloud storage using user behavior prediction. *Computers & Security*, 105, 102245.
- Wahidahwati, W., & Asyik, N. F. (2022). Determinants of auditor's ability in fraud detection. *Cogent Business & Management*, 9(1), 2130165.
- Wilks, T. J., & Zimbelman, M. F. (2004). Decomposition of fraud-risk assessments and auditors' sensitivity to fraud cues. *Contemporary Accounting Research*, 21(3), 719-745.
- Xiao, T., Geng, C., & Yuan, C. (2020). How audit effort affects audit quality: An audit process and audit output perspective. *China Journal of Accounting Research*, 13(1), 109-127.
- Yazdani, B., Lashgari, Z., & Mohammadi Noodeh, F. (2022). The role of auditor characteristics in reducing the risk of financial reporting fraud. *Quarterly Journal of Empirical Accounting Research*, 12(2). (in Persian)