



Modeling the Detection of Firms Financial Fraud under the Implementation of Artificial Neural Network Evaluation Algorithms

Marzieh Poursaedi

Department of Accounting, Borujerd Branch, Islamic Azad University, Borujerd, Iran

Mahmood Hematfar

Associate Professor, Department of Accounting, Borujerd Branch, Islamic Azad University, Borujerd, Iran

Seyed Enayatallah Alavi

Assistant professor of computer Department, faculty of Engineering, Shahid Chamran university of Ahvaz. Ahvaz, Iran

Roya Nasirzadeh

Department of Statistics, Faculty of Science, Fasa University, Fasa, Iran

Abstract

The aim of this study is to model the detection of firms' financial fraud using artificial neural network evaluation algorithms. Quadratic Programming (QP) processes were applied in artificial neural network algorithms to, first, determine the basic algorithm and, second, choose the technical parameters of the artificial neural network, based on time-series data from 2013 to 2022. A diagnostic model was then developed using test and control scales to examine innovative algorithms with the highest accuracy coefficients in predicting financial fraud at the level of capital market companies. Based on systematic sampling, 95 stock exchange companies were selected, providing 950 firm-year observations. The distinction between financially healthy companies and those with the potential for financial fraud was determined through decimalization, and

*** Corresponding Author:** m.hematfar@jaub.ac.ir

How to Cite: Poursaedi, M., Hematfar, M., Alavi, S. E., Nasirzadeh, R. (2025). Modeling the Detection of Firms Financial Fraud under the Implementation of Artificial Neural Network Evaluation Algorithms, *Empirical Studies in Financial Accounting*, 22(87), 83-134. DOI: 10.22054/qjma.2025.82826.2632

companies placed in the fraud-prone deciles were examined using the parameters of the artificial neural network's effectiveness.

Keywords: Artificial Neural Network, Financial Fraud, Quadratic Programming.

1. Introduction

With the advent of technology and digital business exchanges, today's era is far more affected by the negative consequences of fraud in financial statements than in the past, making the development of fraud detection methods a technical necessity in the field of financial knowledge. By providing appropriate assessment opportunities for optimal decision-making, highly efficient financial markets facilitate a balanced flow of information among firms, thereby maintaining the attractiveness of entering this market compared to other markets, such as money markets. In particular, implementing such processes can be considered a strategic financial necessity in developing economies that face serious challenges due to resource constraints. However, reference to scientific evidence and the operational reality of financial markets in these countries indicates the existence of information rents for certain groups of financial decision-makers, which create information asymmetry without requiring expertise, technical evaluation methods, or fundamental analysis.

2 .Literature Review

Disclosure of transparent financial information is considered an important resource in stakeholders' economic decision-making and can contribute significantly to balancing investment markets. However, with theoretical and structural changes in the field of financial transparency, fraud as a behavioral and functional issue has become a challenge to the financial transparency of companies. To understand the official definition of financial fraud, the best reference is Section 240, Clause 4 of the Iranian Auditing Standards, which defines this phenomenon as the intentional actions of companies' executive managers, governing bodies, employees, or third parties that result in the acquisition of illegitimate benefits and cause widespread damage to other stakeholders. An important point noted in Clause 9 of the same standard is the distinction between fraud and error, where

intent is considered the only distinguishing feature. In practice, however, drawing a clear boundary between these two in order to protect stakeholders' rights is not an easy task.

3. Methodology

In terms of data type, this study is classified as a semi-experimental and post-event study in the field of positive financial research, implemented using the neural network analysis method and related techniques. In terms of results, this study can be classified as applied research, and in terms of implementation, it is correlational and algorithmic. First, based on a set of analytical procedures using software such as MATLAB and WEKA, the evaluation processes of artificial neural network algorithms are examined. Then, by comparing the selected algorithms, the most effective type of analysis is determined from the perspective of evaluating financial ratios to predict the probability of fraud in companies. It should be noted that the analytical implementation in this process is based on three steps: extracting financial ratios, evaluating extracted ratios, and finally applying neural networking to the evaluated algorithms.

4. Result

In this study, a systematic review of the literature related to the research field over recent years was first conducted to select financial ratios appropriate for evaluating fraud in capital market companies. Then, based on quadratic programming (QP), the basic algorithm for evaluating the accuracy of corporate fraud was determined using firm-year observations by minimizing the gap between predicted and actual data obtained from the identified financial ratios. Through the adaptive neural fuzzy inference system (ANFIS) test and the cross-validation process (k-fold), it was determined that the unsupervised learning algorithm, which incorporates evaluation parameters based on a meta-heuristic approach and provides higher prediction accuracy, was selected as the foundation for the algorithms in this study. To create a reference index for assessing fraud probability based on financial ratios, the decile method was used to identify which companies with a coefficient of $D > 1$ could be distinguished between financially healthy firms and those with fraud probability, under the constant return to scale (CRS) and variable return to scale (VRS)

evaluation scales. The results indicate that companies with fraud probability were concentrated in four deciles, suggesting that two algorithms, genetic and bee colony selection, were used to further evaluate prediction accuracy. Finally, it was found that the bee colony algorithm had a higher accuracy coefficient in predicting fraud accuracy probability compared to the genetic algorithm. It was also found that the ratio of net profit to sales is the most important criterion for evaluating the accuracy of fraud prediction in the companies under study.

5. Discussion

In interpreting the results, it should first be noted that the bee colony algorithm performs better in solving complex problems due to its multi-process optimization through collective intelligence. This algorithm can be more effective in financial decision-making at the capital market level because it shows higher convergence power and accuracy in predicting the probability of fraud in the companies under study compared to the genetic algorithm. In addition, the coefficients obtained from the bee colony algorithm indicate more effective optimization of financial ratios in predicting the probability of fraud. Financial decision-makers can therefore use this algorithm for more accurate evaluations based on financial ratios, enabling them to identify companies with a probability of fraud and avoid purchasing their shares when forming a portfolio.

6. Conclusion

Given the importance of fraud in financial decision-making, investors are advised to reduce the risk of predicting financial fraud in companies by improving their level of technical analysis training when forming a stock portfolio. The most significant analytical techniques are related to prediction methods based on unsupervised learning algorithms. Focusing on this set of algorithms provides decision-makers with deeper insights by enabling them to recognize the true nature of the data and, without manipulating or remapping it, identify predictable patterns, structures, and relationships of financial fraud in companies.



الگوی تشخیصِ قلب مالی تحت اجرای ارزیابی الگوریتم‌های مطلوبیت شبکه عصبی مصنوعی

- مرضیه پورساعدی ^{id} دانشجوی دکتری، گروه حسابداری، واحد بروجرد، دانشگاه آزاد اسلامی، بروجرد، ایران
- محمود همت‌فر ^{id} * دانشیار، گروه حسابداری، واحد بروجرد، دانشگاه آزاد اسلامی، بروجرد، ایران
- سید عنایت‌اله علوی ^{id} استادیار، گروه کامپیوتر، دانشکده مهندسی، دانشگاه شهید چمران اهواز، اهواز، ایران
- رؤیا نصیرزاده ^{id} استادیار، گروه آمار، دانشکده علوم، دانشگاه فسا، فسا، ایران

چکیده

هدف این مطالعه، مدل‌سازی تشخیص قلب مالی شرکت‌ها تحت اجرای ارزیابی الگوریتم‌های مطلوبیت شبکه عصبی مصنوعی می‌باشد. در این مطالعه با استفاده از فرآیندهای برنامه‌ریزی مجذوری «QP» در الگوریتم‌های شبکه عصبی مصنوعی، تلاش شده است تا طی چندین مرحله بر اساس داده‌های زمانی ۱۳۹۲ تا ۱۴۰۱، اقدام به تعیین الگوریتم پایه در وهله اول و انتخاب پارامترهای تکنیکال شبکه عصبی مصنوعی در وهله دوم گردد. سپس با توسعه یک مدل تشخیصی بر اساس دو مقیاس آزمون و کنترل، الگوریتم‌هایی فرابتکاری که بالاترین ضرایب دقت در پیش‌بینی صحت قلب مالی را دارند، در سطح شرکت‌های بازار سرمایه موردبررسی قرار گیرند. لذا بر اساس فرآیند نمونه‌گیری سیستماتیک، تعداد ۹۵ شرکت بورس اوراق بهادار انتخاب شدند تا بر اساس ۹۵۰ مشاهده (سال-شرکت)، حد فاصل شرکت‌های دارای سلامت مالی با شرکت‌های دارای احتمال قلب مالی از طریق دهک‌بندی تعیین گردد و شرکت‌های قرارگرفته در دهک‌های دارای قلب مالی، از طریق پارامترهای مطلوبیت شبکه عصبی مصنوعی موردبررسی قرار گیرند. نتایج مطالعه نشان داد، الگوریتم یادگیری بدون نظارت که شامل مجموعه

پارامترهای ارزیابی مبتنی بر الگوریتم فرا ابتکاری است، از صحت پیش‌بینی‌های مبتنی بر داده‌های برآورده شده بالاتری برخوردار می‌باشد. همچنین نتایج ناشی از صحت پیش‌بینی تقلب‌های مالی شرکت‌های دهک‌بندی شده بر اساس دو الگوریتم انتخابی ژنتیک و کلونی زنبور عسل نشان می‌دهد، الگوریتم کلونی زنبور عسل از ضریب دقت بالاتری در پیش‌بینی صحت احتمال تقلب شرکت‌های مورد بررسی برخوردار است. همچنین مشخص شد، نسبت سود خالص به فروش مهم‌ترین معیار ارزیابی دقت پیش‌بینی احتمال تقلب در شرکت‌های مورد بررسی می‌باشد.

کلیدواژه‌ها: تقلب مالی، شبکه عصبی مصنوعی، برنامه‌ریزی مجذوری.

۱- مقدمه

بازارهای مالی غالباً کارآمد، با در اختیار قراردادن فرصت‌های ارزیابی مناسب برای تصمیم‌گیری مطلوب، سطح متوازی از جریان گردش اطلاعات عملکردی شرکت‌ها را زمینه‌سازی می‌نماید تا جذابیت‌های ورود به این بازار، در مقایسه با سایر بازارها همچون بازارهای پولی به‌طور متداوم حفظ گردد (Ran et al., 2023). به‌ویژه پیاده‌سازی چنین فرآیندهایی می‌تواند در زمره ضرورت‌هایی استراتژیک مالی در اقتصاد کشورهای در حال توسعه‌ای تلقی گردد که به دلیل وجود محدودیت‌های تأمین منابع، با چالش‌های جدی مواجه هستند (Lisbinski & Burnquist, 2024)؛ اما ارجاع به تجارب علمی و واقعیت عملکردی بازارهای مالی این کشورها، همواره نشان‌دهنده‌ی وجود رانت‌هایی از اطلاعات برای گروه‌های خاصی از تصمیم‌گیرندگان مالی می‌باشد که بدون داشتن نیاز به تخصص یا شیوه‌های ارزیابی تکنیکال و تحلیل‌های بنیادی، سطحی از عدم تقارن‌پذیری اطلاعات را به وجود می‌آورند (Lusardi & Mitchell, 2014). لذا همان‌طور که نظریه‌پردازانی همچون Bockel-Rickermann et al. (2023) و Spann (2013) اذعان نموده‌اند، نیاز به تدوین الگوهای مدل‌سازی شده‌ی تشخیص سیستماتیک تقلب، می‌تواند تاحدی به دستیابی چنین توازی برای تصمیم‌گیری مالی متوازن مؤثر باشد. لذا از آنجایی که تقلب زائیده‌ی تفکر فرصت‌طلبانه‌ای، به دلیل وجود شکاف‌های قانونی و استانداردهای اثرگذار مالی و گزارشگری است که باعث می‌شود تا با تضییع فراگیر حقوق سرمایه‌گذاران و سهامداران، منافع ناشی از دستکاری در صورت‌های مالی و ساختگی‌سازی اطلاعات، صرفاً به گروهی از اقلیت‌های دارای احاطه بر شرایط بازارهای مالی اختصاص یابد، توسعه مدل‌سازی سیستماتیک بر روی این پدیده منفی مالی در بازارهای سرمایه کشورهای در حال توسعه می‌بایست به‌عنوان یک اهرم مؤثر در تصمیم‌گیری‌های مالی متوازن مدنظر قرار گیرد (Zhou et al., 2024).

براین اساس تقلب در گزارشگری را بدون شک می‌بایست یک مسئله جدی در نظام مالی کشورها تلقی نمود که گستره‌ای از ریسک‌های مبتنی بر عدم قطعیت را برای

سرمایه گذاران و تصمیم گیرندگان مالی به همراه دارد (du Toit, 2024). لذا ارجاع به برخی از تجارب شرکت های تجاری در عرصه های بازارهای مالی سراسر جهان مانند شرکت توشیای ژاپن در سال ۲۰۱۵ (Foster, 2015)، قهوه لاکین چین در سال ۲۰۱۹ (Sender et al., 2020) و وایرکارد آلمان در سال ۲۰۲۰ (Teichmann et al., 2023) نشان دهنده ی رسوایی مالی ناشی از تقلب، به واسطه تمرکز بیش از اندازه بر تداوم فعالیت های حسابرسی بدون توسعه ی الگوهای تشخیص سیستماتیک تقلب مالی می تواند تلقی شود؛ به عبارت دیگر، تغییر پارادایم های سنتی و کلاسیک در عرصه ی تشخیص شیوه های تقلب از طریق الگوریتم های مطلوبیت شبکه عصبی مصنوعی، امروزه به یک نیاز تکنیکال جهت ارتقاء تصمیم گیری های مالی در نظام بازار سرمایه بدل شده است که به واسطه ی مدل سازی های اقتضایی در شناسایی عملکردهای مالی متقلبانه ی شرکت ها می تواند به پویایی بالاتر فرآیند گردش سیستماتیک اطلاعات کمک نماید (Motie & Raahemi, 2024). چراکه با پیچیده تر شدن سیستم های مالی و افزایش احتمال پنهان سازی عملکردهای منفی شرکت ها، به واسطه ی ظهور کارکردهای وسیع هوش مصنوعی از طریق مدل زبانی بزرگ «Large language Models» (LLM) یا توسعه ی جی پی تی «Generative Pre – input Transformer» (GPT) در بستر یادگیری ماشینی «Machine Learning» (ML) از یک سو و ظهور تبادلات ناشی از دارایی های دیجیتال مثل توکن های غیرقابل معاوضه «Non – Fungible Token» (NFT) یا انواع رمز ارزها از سویی دیگر، نیاز به الگوریتم های مدل سازی شده ی شبکه عصبی برای تشخیص سیستماتیک تقلب را می تواند حائز اهمیت نماید.

درواقع ظهور چنین فناوری های مالی به عنوان یک شمشیر دولبه در بازار سرمایه، ضمن اینکه می تواند به شفافیت بیشتر کمک کند، درعین حال گسترده ی تقلب های مالی ماحصل دیگری از این تحولات دیجیتال محسوب شود که در این راستا نیاز به مدل سازی های سیستماتیک شناسایی تقلب شرکت ها را بیش از پیش به یک ابزار استراتژیک در شیوه های تحلیل مالی بدل نموده است (Shoetan & Familoni, 2024).

شیوه‌هایی که به دلیل ضعف در نظارت‌های ساختاری یا عدم توسعه‌ی ابزارهای مالی تشخیصی مناسب برای کنترل ریسک فاکتورهای مؤثر بر تقلب مالی شرکت‌ها، باعث تحمیل هزینه‌های زیادی بر سرمایه‌گذاران به‌ویژه در کشورهای در حال توسعه شده است (ملکی کاکلر و همکاران، ۱۴۰۰). به‌طور مصداقی تر اختلال‌هایی از نوع دخالت نهادی یا حاکمیتی در سطح فرآیندی اجرای نظارت اثربخش بر انعکاس شفاف عملکردهای مالی از یک سو و نبود آموزش‌های مالی بروز شده در بستر نرم‌افزارهای مالی توسعه‌یافته در بازار سرمایه از سویی دیگر، همگی به بخشی از بدنه‌ی نهیف سرمایه‌گذاری در نظام مالی بازارهای در حال توسعه بدل شده است. با ارجاع به وقایع بازار سرمایه ایران، در طی سال‌های ۱۳۹۸ تا ۱۴۰۰ می‌توان دریافت که به دلیل ضعف نهادی یا خود خواسته‌ی دولتی، مبنی بر کاهش تورم از طریق جذب سرمایه‌های نقدی عامه‌ی مردم در بورس اوراق بهادار، دستکاری‌های گسترده و متقلبانه‌ی شرکت‌ها و عدم پوشش تمام واقعیت‌های عملکردی به استفاده‌کنندگان اطلاعات کم برخوردار از نفوذ، به ایجاد یک ناترازی مالی گسترده در این بازار منتهی گردید و ریزش شدید قیمت سهام خریداری‌شده را، برای سرمایه‌گذاران ناآگاه از جریان تحلیلی، به همراه داشت (دادگر و همکاران، ۱۴۰۲).

لذا این مطالعه با درک اهمیت شیوه‌های تحلیلی تقلب حتی باوجود پژوهش‌های زیادی که در بازار سرمایه ایران از نظر تحلیل‌های الگوریتمی انجام شده است، به دنبال مدل‌سازی تشخیص سیستماتیک تقلب مالی، تحت اجرای ارزیابی اثربخش الگوریتم‌های مطلوبیت شبکه عصبی مصنوعی می‌باشد. درواقع تفاوت این مطالعه با پژوهش‌هایی همچون آشاروزنیا و همکاران (۱۴۰۲)؛ تشدید و همکاران (۱۳۹۸) و بهرامی و همکاران (۱۴۰۰) که از طریق الگوریتم‌های مختلف شبکه عصبی، اقدام به ارزیابی زمینه‌های تقلب مالی شرکت‌ها نموده‌اند، این است که مطالعه‌ی حاضر تلاش دارد ابتدا با تعیین معیارهای اصلی تأثیرگذار بر شناسایی رویه‌های تقلب مالی شرکت‌های بورس اوراق بهادار تهران در فرآیند بخش کیفی، غربالگری محتوایی سیستماتیک را از پژوهش‌های گذشته باهدف ایجاد یکپارچگی بالاتر انجام دهد تا بر اساس برنامه‌ریزی مجذوری

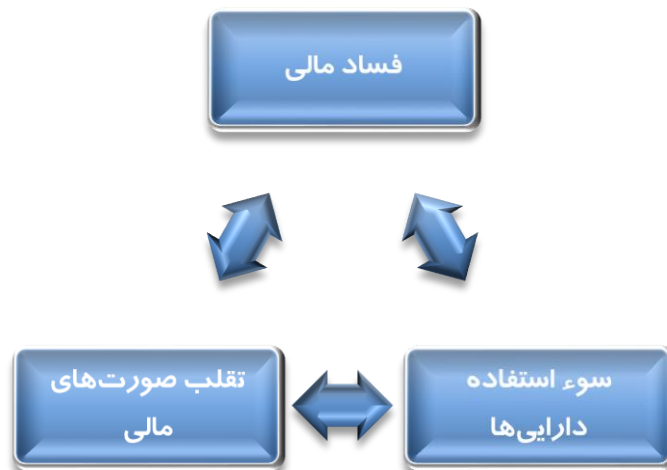
«Quadratic Programming (QP)» در الگوریتم‌های عصبی، با ایجاد زیرمجموعه‌های تصادفی و یا چندبخشی ($k - \text{fold}$) در تفکیک داده‌های واقعی از داده‌های قابل پیش‌بینی، نسبت به انتخاب مهم‌ترین الگوریتم اقدام شود تا با تعیین الگوریتم‌های مرجع، مدلی به‌واسطه‌ی مقیاس‌های کنترل و آزمون به وجود بیاید تا تعیین‌کننده‌ی مبنای تفاوت بین شرکت‌های متقلب با شرکت‌های دیگر در نظر گرفته شود.

۲- مبانی نظری

افشاء اطلاعات شفاف مالی به‌عنوان یک منبع مهم در تصمیم‌گیری‌های اقتصادی ذینفعان تلقی می‌شود که می‌تواند به توازن بازارهای سرمایه‌گذاری کمک شایان توجهی نماید (Uras, 2020)؛ اما طی تحولات ساختاری نظری و پارادایمی عرصه‌ی شفافیت‌های مالی، تقلب به‌عنوان یک مبنای رفتاری و عملکردی به معضلی برای شفافیت‌های مالی شرکت‌ها بدل شده است (Sargiacomo et al., 2024). برای شناخت از تعریف رسمی تقلب مالی، بهترین مرجع، بند ۴ بخش ۲۴۰ استانداردهای حسابرسی ایران می‌باشد که در تعریف این پدیده مالی اشاره به اقدام تعمّدی شرکت‌ها توسط مدیران اجرایی، ارکان راهبری، کارکنان، یا اشخاص ثالث دارد که باعث تصاحب منافع غیر مشروعی می‌گردد که تضییع گسترده‌ی سایر ذینفعان را به همراه دارد (حسن‌پور و همکاران، ۱۳۹۹)؛ اما نکته‌ی مهم این است که در بند ۹ همین استاندارد، بین تقلب و اشتباه تفاوت قائل شده‌اند و تنها ویژگی متمایزکننده تقلب از اشتباه را قصد و نیت تلقی نمودند که اساساً رسیدن به مرز مشخصی از این دو مبنا برای رعایت حقوق ذینفعان کار راحتی نمی‌باشد (اعتمادی و عبدلی، ۱۳۹۶). لذا تقلب مالی را می‌بایست کارکردی منفی بر تعامل بین شرکت‌ها با سایر استفاده‌کنندگان از اطلاعات تلقی نمود که به دلیل ضعف استانداردها و قوانین تجاری، فرصت‌هایی برای سوءاستفاده در اختیار اداره‌کنندگان شرکت‌ها قرار می‌گیرد که این مسئله می‌تواند تا مرز سقوط و بروز بحران‌های مالی را برای شرکت‌ها به همراه داشته باشد (Smaili et al., 2020). ارجاع به بیانیه شماره ۸۲ هیئت استانداردهای حسابرسی در مورد تقلب نشان می‌دهد که حسابرسان به‌عنوان یکی از ارکان نهادی، ملزم به بررسی تمامی جوانب

عملکردی شرکت‌ها هستند تا بتوانند مانع از تضییع حقوق سهامداران و سرمایه‌گذارانی گردند که به دلیل عدم آگاهی از قوانین و ساختارهای استنادی، معمولاً حقوقشان نادیده گرفته می‌شود (Dezoort & Lee, 1998). از طرف دیگر بر اساس گزارش انجمن بازرسان رسمی تقلب «The Association of Fraud Examiners» (ACFE)، می‌توان بروز تقلب حرفه‌ای را برای شناخت بهتر از ماهیت و نیت آن به سه دسته در قالب شکل (۱) ارائه نمود.

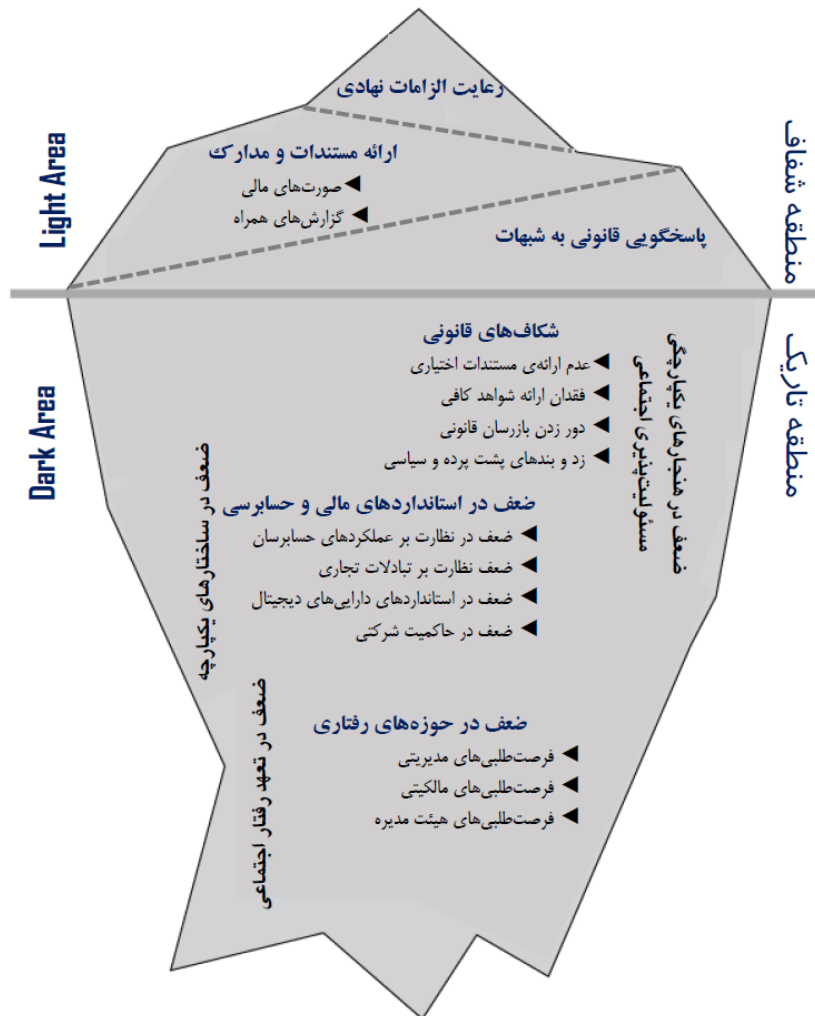
شکل ۱. انواع تقلب بر مبنای رویکرد انجمن بازرسان رسمی تقلب



بر اساس این رویکرد، می‌توان دریافت تقلب حرفه‌ای ماهیت مستقلی در ارزیابی حقوقی ندارد، بلکه کارکردی وابسته به سه عنصر فساد، سوءاستفاده از دارایی‌ها و دستکاری در حساب‌ها برای گمراهی استفاده‌کنندگان از صورت‌های مالی می‌باشد. فساد مالی را می‌توان تقلبی در سطح فردی و واحد سازمانی تلقی نمود. سوءاستفاده از دارایی‌ها را می‌توان رفتاری جمعی در بهره‌برداری نادرست و نامشروع از دارایی‌های شرکت‌ها تلقی نمود و درنهایت تقلب صورت‌های مالی را دستکاری ناشی از منفعت‌طلبی ساختاری شرکت‌ها باهدف فریب استفاده‌کنندگان از صورت‌های مالی تعریف نمود (تشددی و همکاران، ۱۳۹۸). به لحاظ تأثیرگذاری منفی بر تصمیم‌های استفاده‌کنندگان از اطلاعات،

می‌بایست وقوع تقلب مبتنی بر عدم شفافیت در صورت‌های مالی را مهم‌ترین خروجی در تعاملات بین شرکت‌ها با ذینفعان تلقی نمود که می‌تواند به بروز شکاف‌های عمیق نمایندگی و ریسک‌هایی همچون سقوط قیمت سهام و ورشکستگی شرکت‌ها منتج گردد. به‌منظور شناخت بهتر ماهیت تقلب، نظریه‌های چندانی برای تفهیم آن در کارکردهای بازارهای مالی توسعه‌نیافته است، اما دو نظریه قله کوه یخی «Iceberg Theory» و نظریه سوسک حمام «Cockroach Theory» احتمالاً سرآمدترین نظریه‌های مرتبط با این پدیده منفی در بازارهای مالی باشند. طبق نظریه قله کوه یخی که در شکل (۲) مشاهده می‌شود، شواهد ارائه‌شده مالی به‌صورت گزارشگری در دید همگان قرار دارد که رعایت الزاماتی است که بر اساس استانداردهای حسابداری می‌تواند مبنای تصمیم‌گیری‌های مالی قرار گیرد که بر اساس اسناد و مدارک قابل استناد به نهادهای ناظری همچون بازرسان قانونی و حسابرسان ارائه می‌شود؛ اما نیمه‌ی پایینی کوه یخی که قالب مشاهده نیست، اصطلاحاً «منطقه تاریکی» (Dark Area) شناخته می‌شود که به‌واسطه‌ی وجود زمینه‌های رفتاری می‌تواند کارکردهای تقلب‌آمیزی تلقی گردد که به دلیل شکاف‌های قانونی، عدم نظارت‌های نهادی و فرصت‌طلبی‌های ساختاری شرکت‌ها به ضعف در حاکمیت شرکتی منتج می‌گردد و باعث شکاف عمیقی در شناخت از واقعیت‌ها، بین شرکت با سهامداران و سرمایه‌گذاران می‌گردد (Jiang & Cui, 2019). لذا طبق این نظریه، ناظران قانونی و حسابرسان در راستای وظایف محوله و مسئولانه‌ای که دارند می‌بایست به سراغ ناشناخته‌هایی بروند که الزاماً از طریق اسناد و مدارک ارائه‌شده قابل ارزیابی برای ارتقاء شفافیت‌های اطلاعاتی تلقی نمی‌شوند (Francis & Justine, 2023).

شکل ۲. فرآیندهای نظریه قله کوه یخی



بنابراین این نوع از تقلب‌ها به دلیل ضعف در قانون‌گذاری، تأمین‌کننده‌ی منافع نامشروعی هستند که باعث تضییع گسترده‌ی منافع سهامداران و سرمایه‌گذارانی می‌شوند که از دسترسی و نفوذ کافی به عملکردهای منفعت‌طلبانه‌ی شرکت‌ها برخوردار نیستند. از طرف دیگر نظریه سوسک حمام که توسط Blakley (2009) مطرح گردید، بر عادی شدن فرآیند تقلب مالی به صورت یک رفتار تکرارپذیر اشاره دارد. این نظریه با تشبیه به مضمون

رفتار سوسک در ورود به محیط فاضلاب، اذعان می‌نماید که به دلیل ویژگی‌های بیولوژیک این حشره، بزنگاهی برای زیستن می‌تواند به افزایش آن‌ها در یک زیستگاه بدل شود. لذا تقلب در محیطی که نظارتی وجود نداشته باشد، می‌تواند به تکثیر رفتارهای مشابهی در یک ساختار منتج شود که حتی در صورت مشاهده، امکان مقابله با آن به دلیل گستردگی فساد وجود ندارد.

لذا همان‌طور که نظریه‌های مرتبط با تقلب بر آن تأکید دارند، غالباً ماهیت این پدیده در سایه اتفاق می‌افتد جایکه دسترسی به اسناد اثبات‌کننده رفتارهای ناهنجار مالی معمولاً دشوار است، چراکه تقلب آمیخته با نیاتی است که غالباً ارکان قدرت از آن برای رساندن به اهداف خود بهره می‌برند. ازاین‌رو محتمل‌ترین حالت در تفسیر بروز تقلب، پذیرش این ادراک نسبت به پدیده تقلب در ساختارهای شرکت‌ها است که بایستی همواره به صورت یک پیش‌فرض متداوم در اجرای فرآیندهای مالی آن را در نظر گرفت تا بتوان بر اساس مکانیزم‌هایی که امکان کشف تقلب را برای تصمیم‌گیرندگان مالی مشهود می‌نماید، تأثیرگذاری منفی آن را بر ذی‌نفعان کاهش داد. لذا با اتکا به این توضیحات و ماهیت مطالعه، سوال‌های پژوهش به ترتیب زیر ارائه می‌شود:

- سؤال اول پژوهش (معیارهای ارزیابی بروز تقلب مالی در شرکت‌های بازار سرمایه کدام‌اند؟)
- سؤال دوم پژوهش (الگوریتم پایه شبکه عصبی مصنوعی در ارزیابی بروز تقلب مالی در شرکت‌های بازار سرمایه کدام‌اند؟)
- سؤال سوم پژوهش (مطلوب‌ترین الگوریتم‌های پایه شبکه عصبی مصنوعی بر اساس پارامترهای فرا ابتکاری کدام‌اند؟)

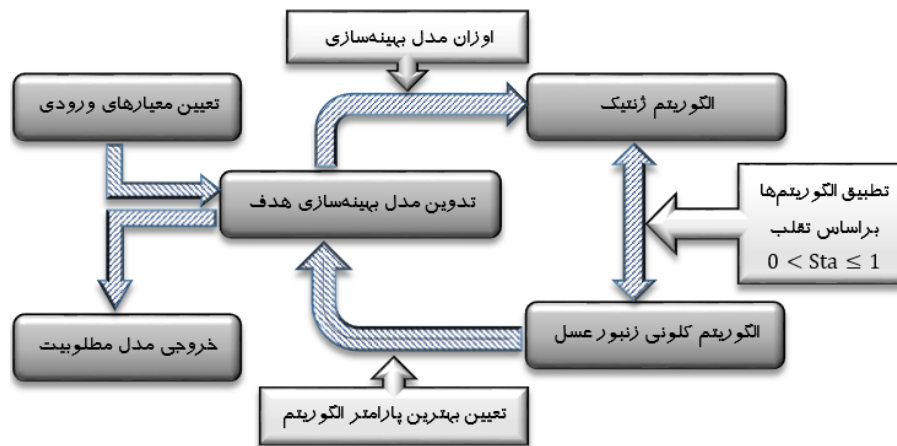
همان‌طور که مشاهده می‌شود، سؤال اول پژوهش بر اساس یک فرآیند نظام‌مند در غربالگری سیستماتیک محتوایی، اقدام به تعیین معیارهای ارزیابی بروز تقلب مالی شرکت‌ها می‌نماید. سپس از طریق مجموع برنامه‌ریزی مجذوری «QP» در الگوریتم‌های عصبی، با ایجاد زیرمجموعه‌های تصادفی و یا چندبخشی (k-fold)، الگوریتم پایه برای

پاسخ به سؤال دوم پژوهش تعیین می‌شود تا درنهایت با تشکیل مقیاس‌های کنترل و آزمون در الگوریتم فرا ابتکاری به‌عنوان مبنای تحلیلی الگوریتم یادگیری بدون نظارت، اقدام به انتخاب الگوریتم‌های تکنیکال برای پیش‌بینی صحت زمینه‌های تقلب مالی شرکت‌ها می‌شود.

۳- روش شناسی

این مطالعه از نظر نوع داده در دسته مطالعه‌های نیمه تجربی و پس‌رویدادی در حوزه تحقیقات اثباتی دانش مالی قرار می‌گیرد که با استفاده از روش تحلیل شبکه عصبی و سایر تحلیل‌های وابسته، پیاده‌سازی می‌شود. لذا از نظر نتیجه، این مطالعه را می‌توان در دسته پژوهش‌های کاربردی طبقه‌بندی نمود که کارکرد آن به لحاظ اجرا، از نوع همبستگی و الگوریتمی می‌باشد. لذا ابتدا بر اساس مجموع دستورهای تحلیلی از طریق نرم‌افزارهایی همچون متلب و نرم‌افزار WEKA تلاش می‌شود تا فرآیندهای ارزیابی الگوریتم‌های شبکه عصبی مصنوعی مورد بررسی قرار گیرد و سپس با مقایسه الگوریتم‌های انتخاب‌شده، بهترین نوع تحلیل از منظر ارزیابی نسبت‌های مالی جهت تعیین پیش‌بینی احتمال تقلب شرکت‌ها مشخص شود. لازم به ذکر است که پیاده‌سازی تحلیلی در این فرآیند بر اساس سه گام تحلیلی است که شامل؛ استخراج نسبت‌های مالی، ارزیابی نسبت‌های مالی استخراج‌شده و درنهایت شبکه‌سازی عصبی الگوریتم‌های مورد ارزیابی می‌باشد. لذا با استناد به توضیح‌های داده‌شده می‌توان فرآیندهای تشریح شده در پیاده‌سازی تحلیلی پژوهش را به ترتیب زیر ارائه نمود:

شکل ۳. فرآیند پیاده‌سازی تحلیلی



این چارچوب شاکله سه مرحله فوق را برای ارائه یک الگوی تشخیص تقلب در شرایط عدم قطعیت را نشان می‌دهد که می‌تواند با ارزیابی نسبت‌های مالی به بررسی این مسئله پردازد که چه سازوکاری در تشخیص تقلب‌های مالی شرکت‌ها می‌تواند به آگاهی‌های تصمیم‌گیرندگان مالی کمک نماید.

جامعه آماری و بازه زمانی پژوهش

واحد موردبررسی در این مطالعه، شرکت‌های بازار سرمایه‌ای هستند که در بازه ۱۳۹۲ تا ۱۴۰۱ داده‌های مرتبط با هریک از نسبت‌های مالی را می‌توان بر اساس شرایط زیر، از آنها استخراج نمود.

- در وهله اول شرکت‌هایی در این شرایط قرار می‌گیرند که تاریخ ورود آن به بورس اوراق بهادار قبل بازه زمانی مطالعه باشد و تا پایان دوره موردبررسی، جزء شرکت‌های بورس اوراق بهادار قلمداد شوند.
- در وهله دوم شرکت‌های مورد مطالعه به‌عنوان واحد استخراج داده‌ها می‌بایست از سال مالی منتهی به پایان اسفندماه برخوردار باشند.
- همچنین شرکت‌های انتخابی به‌عنوان واحد استخراج نسبت‌ها، تغییر فعالیت در بازه سال

مالی تعیین شده را نداشته باشد.

- شرکت‌های انتخاب شده، جزء شرکت‌های سرمایه‌گذاری و صندوق‌های بازنشستگی یا بیمه‌ای نباشند.

- طول وقفه انجام معاملات شرکت‌های موردبررسی، بیش از ۶ ماه نباشد.

با مدنظر قراردادن معیارهای غربالگری حذف سیستماتیک، تعداد ۹۵ شرکت (۹۵۰ سال- شرکت) از ۲۵ صنعت مختلف به عنوان جامعه غربالگری شده، انتخاب گردیدند.

متغیرهای پژوهش

در این مطالعه به منظور یکپارچگی بیشتر تعیین معیارهای ارزیابی احتمال تقلب شرکت‌های بازار سرمایه، از فرآیند غربالگری نظام‌مند پژوهش‌های مشابه در بازه زمانی ۴ سال گذشته در داخل و خارج بهره برده شده است. لذا با ارجاع به پایگاه‌های اطلاعاتی مرجع و تعیین کلیدواژگان مشخص طبق جدول (۱)، تلاش شد تا پژوهش‌های مشابه با ماهیت موردبررسی این مطالعه مورد غربالگری هدفمند قرار گیرند.

جدول ۱. جستجوی کلمات کلیدی در انتخاب پژوهش‌های مشابه

کلمات کلیدی جستجو در پژوهش‌های خارجی	کلمات کلیدی جستجو در پژوهش‌های داخلی
Financial Fraud	تقلب مالی
Neural Network Algorithms for Fraud Prediction	الگوریتم‌های شبکه عصبی پیش‌بینی تقلب
Financial Fraud Risks	ریسک‌های تقلب مالی
Fraudulent Reporting	گزارشگری متقلبانه
Data Mining to Detect Fraud	داده‌کاوی کشف تقلب

بر اساس این کلمات کلیدی، مطابق با شکل (۴) اقدام به ارزیابی دو مرحله‌ی پژوهش‌های اولیه شناسایی شده در بازه زمانی ۲۰۲۰ تا ۲۰۲۴ و ۱۳۹۹ تا ۱۴۰۳ می‌شود. درواقع در غربالگری اول ماهیت محتوایی پژوهش‌ها و نزدیکی به نوع تدوین سؤالات مدنظر قرار گرفت و در غربالگری دوم نیز، فرآیندهای تحلیلی مورد استفاده بررسی شد.

شکل ۴. ارزیابی دو مرحله‌ای پژوهش‌های اولیه شناسایی شده



با تعیین ۷ پژوهش ناشی از غربالگری نتیجه‌ی ارزیابی دو مرحله‌ی اولیه پژوهش‌های شناسایی‌شده، در ادامه با مرور نظام‌مند متون پژوهش‌های انتخاب‌شده، اقدام به تعیین معیارهای مؤثر در پیش‌بینی احتمال تقلب مالی شرکت‌ها می‌شود که نتایج آن در جدول (۲) ارائه شده است.

جدول ۲. غربالگری نظام‌مند شناسایی معیارهای مؤثر در پیش‌بینی احتمال تقلب مالی

معیارهای ارزیابی	Zhou et al (2024)	Motie and Raahemi (2024)	du Toit (2024)	Zhang et al (2022)	Craja et al (2020)	جعفری و همکاران (۱۴۰۲)	ملکی کاکر و همکاران (۱۴۰۰)	جمع فراوانی
نسبت جاری	✓	-	✓	✓	-	-	✓	۴
نسبت تمرکز سرمایه (وام)	-	✓	-	✓	-	-	-	۳
نسبت دارایی جاری به کل دارایی‌ها	✓	✓	✓	-	✓	✓	-	۵
نسبت انعطاف‌پذیری سرمایه	✓	-	-	✓	-	✓	-	۳
نسبت دارایی ثابت به کل دارایی‌ها	✓	-	✓	-	✓	✓	-	۴
نسبت سود انباشته به کل دارایی‌ها		✓	✓	✓			✓	۴
نسبت وجوه مالکانه	-	✓	-	-	-	-	✓	۲
نسبت سرمایه در گردش به کل	✓	✓	-	✓	✓	-	-	۴

معیارهای ارزیابی	Zhou et al (2024)	Motie and Raahemi (2024)	du Toit (2024)	Zhang et al (2022)	Craja et al (2020)	جعفری و همکاران (۱۴۰۲)	ملکی کاکر و همکاران (۱۴۰۰)	جمع فراوانی
دارایی‌ها								
نسبت کل بدهی به کل دارایی‌ها	✓	-	✓	-	-	✓	✓	۴
نسبت دوره وصول مطالبات	-	✓	-	-	✓	-	-	۲
نسبت پوشش هزینه مالی	✓	-	-	✓	-	-	-	۲
نسبت سود ناخالص به کل دارایی‌ها	✓	-	✓	✓	-	✓	-	۴
نسبت سود هر سهم	-	✓	-	-	✓	-	-	۲
نسبت کل بدهی به حقوق صاحبان سهام	-	✓	✓	-	✓	-	✓	۴
نسبت نوسان (انحراف معیار) سودآوری	✓	-	-	✓	-	✓	-	۳
نسبت سود انباشته به کل دارایی‌ها	✓	-	✓	✓	-	✓	-	۴
نسبت فروش به کل دارایی‌ها	✓	✓	-	-	✓	-	✓	۴
نسبت دوره بازپرداخت بدهی	-	✓	-	-	-	✓	-	۲
نسبت بدهی بلندمدت به حقوق صاحبان سهام	✓	✓	-	✓	✓	-	✓	۵
نسبت بهای تمام‌شده کالای فروش رفته به کل فروش	-	✓	✓	-	✓	✓	-	۴
نسبت سود خالص به کل دارایی‌ها	✓	-	✓	✓	-	✓	-	۴
نسبت استقلال مالی	✓	-	-	-	-	-	✓	۲
نسبت سود ناخالص به فروش	-	✓	-	✓	✓	-	✓	۴
نسبت منابع استقراری به دارایی‌های غیر جاری	-	✓	-	✓	-	-	-	
نسبت سود عملیاتی به فروش	✓	✓	-	✓	-	✓	✓	۴
نسبت آبی	-	-	✓	-	✓	-	-	۲
نسبت سود خالص به حقوق صاحبان	✓	-	✓	-	✓	✓	-	۴

معیارهای ارزیابی	Zhou et al (2024)	Motie and Raahemi (2024)	du Toit (2024)	Zhang et al (2022)	Craja et al (2020)	جعفری و همکاران (۱۴۰۲)	ملکی کاکر و همکاران (۱۴۰۰)	جمع فراوانی
سهام								
نسبت هزینه‌های مالی به کل بدهی‌ها	-	✓	✓	✓	✓	-	✓	۵
نسبت گردش دارایی‌ها	✓	-	-	-	-	✓	-	۲
نسبت سود خالص به فروش	✓	-	✓	-	✓	✓	-	۴
نسبت هزینه‌های عملیاتی به فروش	-	✓	✓	✓	-	✓	-	۴
نسبت خالص سرمایه در گردش به کل بدهی‌ها	✓	-	-	-	-	-	✓	۲
نسبت حساب‌های دریافتی به فروش	✓	-	✓	✓	-	-	✓	۴
نسبت وام اخذشده به وجود مالکانه	-	✓	-	-	-	✓	-	۲
نسبت موجودی‌ها به فروش	-	✓	✓	-	✓	-	✓	۴
نسبت گردش دارایی ثابت	✓	-	-	-	-	✓	-	۲
نسبت موجودی نقد به کل دارایی‌ها	✓	-	✓	✓	-	✓	-	۴
نسبت گردش سرمایه‌گذاری	-	✓	-	-	-	-	✓	۲
نسبت سود خالص به سود ناخالص	✓	✓	-	✓	✓	✓	-	۵
نسبت جریان وجه نقد عملیاتی به کل دارایی‌ها	-	✓	-	✓	✓	✓	-	۴
نسبت بازده دارایی‌ها	✓	-	✓	-	✓	-	-	۳

لذا همان‌طور که مشاهده می‌شود، ۲۴ نسبت مالی برای ارزیابی پیش‌بینی تقلب مالی شرکت‌های بورس اوراق بهادار تهران انتخاب شدند تا از طریق الگوریتم‌های شبکه عصبی موردبررسی قرار گیرند.

یافته‌های پژوهش

باتوجه به تعیین معیارهای ارزیابی تقلب شرکت‌ها در بخش کیفی که شامل ۲۴ نسبت مالی می‌باشد، در ادامه مجموعه بر اساس برنامه‌ریزی مجذوری «QP» که در مطالعه‌هایی با توالی تکنیک‌های یادگیری ماشین‌بردار در الگوریتم‌های شبکه عصبی اجرا می‌شود و از روش‌های لاگرانژ افزوده به عنوان مبنای حل بهینه در مسائل استفاده می‌کند، بهره برده می‌شود. برای این منظور ابتدا می‌بایست بر اساس سه مبنای پایه‌ی

- الگوریتم‌های یادگیری نظارت‌شده یا نظارتی
«Supervised Learning Algorithms»

- الگوریتم‌های یادگیری بدون نظارت
«Unsupervised Learning Algorithms»

- الگوریتم‌های تقویتی «Reinforcement learning algorithms»

از طریق روش آموزش (Train)، نسبت به واریسی اعتبار (CV) به عنوان یک روش توسعه یافته و مورد پذیرش برای آنالیز صحت پیش‌بینی در شبکه عصبی برای انتخاب یکی از سه الگوریتم یادشده، اقدام شود تا امکان ارزیابی مناسب را هم‌راستا با داده‌های مطالعه ایجاد نماید. لذا به طور خلاصه ابتدا هر سه الگوریتم فوق توضیح داده می‌شود تا در ادامه بر اساس زیرمجموعه‌های تصادفی و یا چندبخشی (k – fold) در تفکیک داده‌های واقعی از داده‌های قابل پیش‌بینی، نسبت به انتخاب مهم‌ترین الگوریتم در این مطالعه اقدام شود.

▪ الگوریتم‌های یادگیری نظارت‌شده یا نظارتی

در این الگوریتم معمولاً داده‌هایی که در محیط آزمایشگاهی نگاشت می‌شوند استفاده می‌شود تا از طریق تکامل داده‌ها، نسبت به ارزیابی تغییرات آن اقدام شود. در یادگیری نظارت‌شده، الگوریتم با مجموعه‌ای از داده‌های ورودی همراه با خروجی‌ها یا برجسب‌های موردنظر مربوطه ارائه می‌شود. هدف یادگیری یک تابع نگاشت است که بتواند خروجی ورودی‌های جدید و نادیده را به دقت پیش‌بینی کند.

■ الگوریتم‌های یادگیری بدون نظارت

الگوریتم‌های یادگیری بدون نظارت، دسته‌ای از الگوریتم‌های عصبی، تلقی می‌شوند که با تمرکز بر روی ماهیت واقعی داده‌ها و بدون برچسب‌زنی یا نگاشت آن‌ها در پی کشف الگوها، ساختارها و روابطی هستند که الزاماً ممکن است شواهدی برای آن در گذشته وجود نداشته باشد. این الگوریتم برخلاف یادگیری تحت نظارت که در آن الگوریتم از داده‌های برچسب‌گذاری شده یاد می‌گیرد، یادگیری بدون نظارت بر یافتن ساختارها و الگوهای ذاتی درون خود داده تمرکز می‌کند.

■ الگوریتم‌های تقویتی

الگوریتم‌های یادگیری تقویتی دسته‌ای از الگوریتم‌های عصبی هستند که از طریق تعامل با یک محیط برای به حداکثر رساندن سیگنال عملکردی آن‌ها موردبررسی قرار می‌گیرند. لذا با ترکیب با متغیرهای ثانویه، تلاش می‌شود تا ماهیت درونی داده‌های موردبررسی از طریق آزمون و خطا مشخص شوند.

لذا با تعیین معیارهای ارزیابی احتمال تقلب مالی که از طریق فرآیند غربالگری محتوایی سیستماتیک تعیین گردید، اقدام به تفکیک ارزش واقعی با ارزش فازی می‌شود. برای این منظور با تعیین مقادیر فازی و واقعی بر اساس مجموع نسبت‌های مالی طبق ۹۵۰ مشاهده (سال-شرکت) از طریق دو معیار « R^2 » و «MSE» درصد پیش‌بینی صحیح مقادیر شاخص‌های مالی نسبت به مقدار واقعی تعیین می‌شود بر اساس این نتیجه مشخص شده است، هر قدر مجموع شاخص‌های مالی به ۱/۰۰ نزدیک‌تر باشند، می‌تواند در ارزیابی الگوریتم‌های شبکه عصبی مورد استفاده قرار گیرد. لذا طبق نتیجه جدول (۳) شاخص‌های تقلب ما بر اساس دو معیار « R^2 » و «MSE» نشان‌دهنده‌ی تفاوت مقادیر واقعی و بر اساس ۹۵۰ مشاهده می‌باشد.

جدول ۳. مقادیر واقعی و پیش‌بینی‌شده در خصوص پیش‌بینی نسبت‌های مالی

مشاهده	واقعی	الگوریتم شبکه عصبی	
		R ²	MSE
۱	۱۱۱.۰	۵۵۵.۰	۳۶.۰
۲	۰۵۳.۰	۷۴۵.۰	۴۵۳.۰
۳	۰۴۱.۰	۷۶۹.۰	۴۸۵.۰
۴	۰۱.۰	۷۹۸.۰	۴۴۶.۰
۵	۱۶۶.۰	۸۸۲.۰	۵۸۲.۰
۶	۲۶۶.۰	۵۰۸.۰	۴۳.۰
۷	۱۵۱.۰	۷۹۵.۰	۳۰۹.۰
۸	۱۰۴.۰	۸۷۱.۰	۱۹۴.۰
۹	۱۰۴.۰	۸۹۳.۰	۱۰۲.۰
۱۰	۴۸۵.۰	۸۹۵.۰	۲۸۲.۰
⋮	...	...	...
۹۵۰	۰۴۸.۰	۹۴۱.۰	۰۳.۰

با تعیین این مقادیر بر اساس ۹۵۰ مشاهده انجام‌شده، جهت تعیین درصد پیش‌بینی صحیح نسبت‌های مالی برای ارزیابی احتمال تقلب نسبت به مقدار واقعی، می‌بایست از کدهای دستوری نروفازی (ANFIS) در پایتون به‌عنوان مبنای ارزیابی الگوریتم‌های شبکه عصبی بهره برده می‌شد که اطلاعات آن در شکل (۵) ارائه شده است.

شکل ۵. تعیین درصد تناسبی پیش‌بینی صحیح مقادیر شاخص مالی تعیین تقلب نسبت به مقدار

واقعی

array : array_like of rank N

Input array

pad_width : {sequence, array_like, int}

Number of values padded to the edges of each axis. ((before_1, after_1), ... (before_N, after_N)) unique pad widths for each axis. ((before, after),) yields same before and after pad for each axis. (pad,) or int is a shortcut for before = after = pad width for all axes.

mode : str or function

One of the following string values or a user supplied function.

‘constant’

Pads with a constant value.

‘edge’

Pads with the edge values of array.

‘linear_ramp’

Pads with the linear ramp between end_value and the array edge value.

‘maximum’

Pads with the maximum value of all or part of the vector along each axis.

‘mean’

Pads with the mean value of all or part of the vector along each axis.

‘median’

Pads with the median value of all or part of the vector along each axis.

‘minimum’

Pads with the minimum value of all or part of the vector along each axis.

‘reflect’

Pads with the reflection of the vector mirrored on the first and last values of the vector along each axis.

‘symmetric’

Pads with the reflection of the vector mirrored along the edge of the array.

‘wrap’

Pads with the wrap of the vector along the axis. The first values are used to pad the end and the end values are used to pad the beginning.

<function>

Padding function, see Notes.

stat_length : sequence or int, optional

Used in ‘maximum’, ‘mean’, ‘median’, and ‘minimum’. Number of values at edge of each axis used to calculate the statistic value.

((before_1, after_1), ... (before_N, after_N)) unique statistic lengths for each axis.

((before, after),) yields same before and after statistic lengths for each axis.

(stat_length,) or int is a shortcut for before = after = statistic length for all axes.

Default is **None**, to use the entire axis.

constant_values : sequence or int, optional

Used in 'constant'. The values to set the padded values for each axis.

((before_1, after_1), ... (before_N, after_N)) unique pad constants for each axis.

((before, after),) yields same before and after constants for each axis.

(constant,) or int is a shortcut for before = after = constant for all axes.

Default is 0.

end_values : sequence or int, optional

Used in 'linear_ramp'. The values used for the ending value of the linear_ramp and that will form the edge of the padded array.

((before_1, after_1), ... (before_N, after_N)) unique end values for each axis.

((before, after),) yields same before and after end values for each axis.

(constant,) or int is a shortcut for before = after = end value for all axes.

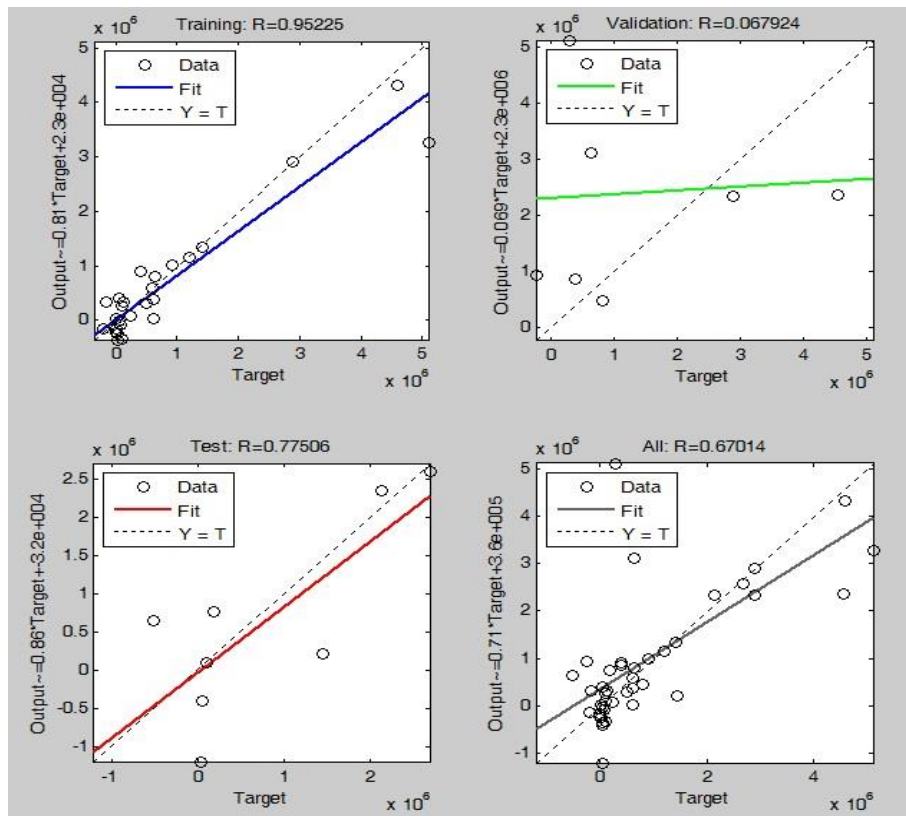
Default is 0.

reflect_type : {'even', 'odd'}, optional

Used in 'reflect', and 'symmetric'. The 'even' style is the default with an unaltered reflection around the edge value. For the 'odd' style, the extended part of the array is created by subtracting the reflected values from two times the edge value.

نتایج حاصل نشان می‌دهد، طبق دو مقیاس R^2 و «MSE»، تقریباً ۶۶٪ درصد پیش‌بینی نسبت‌های تعیین‌شده، به‌درستی احتمال تقلب‌های مالی شرکت‌ها را شناسایی می‌کنند و درصد خطای احتمالی بر اساس تعداد مشاهدات، برابر با ۰/۰۶ است. فرآیند کسب‌شده در شکل (۶) نشان‌دهنده صحت درصد پیش‌بینی صحیح نسبت‌های مالی تعیین‌شده برای ارزیابی احتمال تقلب مالی شرکت‌ها می‌باشد.

شکل ۶. درصد پیش‌بینی صحیح مقادیر نسبت‌های مالی ارزیابی تقلب نسبت به مقدار واقعی



در ادامه و بر اساس منطق استنتاج فازی عصبی انطباق‌پذیر «Adaptive Neuro – Fuzzy Inference System» (ANFIS) می‌بایست، منطقی جهت انجام آزمون واریسی اعتباری برای تعیین بهترین مقادیر ایجاد نمود. درواقع این آزمون که اصطلاحاً به آن روش آموزش گفته می‌شود، اقدام به ارزیابی مطلوبیت فرآیندهای تحلیلی شبکه عصبی فازی برای تعیین اعتبارسنجی انتخاب الگوریتم‌ها می‌نماید. لذا از طریق از روش آموزش (Train) به وسیله‌ی معیار واریسی اعتبار (CV) اقدام کرد. بر این اساس معیارهای MMC، f – measure، recall، precision و accuracy مبنای تعیین مطلوبیت الگوریتم‌های شبکه عصبی قرار می‌گیرند. بر اساس مراحل زیر و بسط معادلات ریاضی اقدام به انجام تحلیل روابط برای تعیین مقادیر مشاهده‌شده تا مقادیر

پیش‌بینی شده را در بازه ۱ تا ۱- نشان می‌دهد.

اعتبارسنجی ضربدری یا اصطلاحاً k -fold یکی از عمومی‌ترین انواع اعتبارسنجی در تحلیل‌های فازی است که با استفاده از گام‌های زیر انجام می‌شود. مجموعه داده اصلی را به k زیرمجموعه مساوی تقسیم می‌کنیم. هر زیرمجموعه یک k -fold نامیده می‌شود که برای تعریف فولدها به صورت $f_1, f_2, \dots, f_k$ نام گذاری می‌شود.

رابطه $\text{For } i = 1 \text{ to } i = k(1)$

فولد f_i را به عنوان مجموعه اعتبارسنجی محسوب می‌شود و همه $k - 1$ فولد باقی مانده را در مجموعه آموزشی اعتبارسنجی ضربدری «Cross validating training set» برآورد می‌شود.

سپس بر اساس یادگیری ماشین بُردار با استفاده از مجموعه داده‌های قرار گرفته شده در موقعیت فازی سلسله مراتبی؛ ویکور و تاپسیس که از طریق نرم افزار WEKA ترتیب داده می‌شوند، میزان صحت هر کدام از آن‌ها در هر سه حالت فوق مورد بررسی قرار می‌گیرد. در این پژوهش، از روش اعتبارسنجی k -باره (k -تایی) استفاده شد. در این روش کل داده‌ها به K دسته تقریباً مساوی تقسیم می‌شوند. ($K - 1$) دسته برای آموزش مدل و دسته باقیمانده برای آزمایش مدل به کار می‌رود. به این ترتیب به تعداد K مرتبه مدل آموزش و آزمایش مورد محاسبه قرار می‌گیرد (He et al., 2019). لذا برای تقسیم‌بندی کل داده‌ها به مجموعه داده‌های آموزش (سه حالت مورد نظر) و آزمایش (تعاریف پیش فرض بر اساس داده‌ها) به روش اعتبارسنجی ۱۰ باره استفاده شده است. لذا بر اساس دستور Stratified cross-validation ۲۴ نسبت مالی ارزیابی احتمال تقلب مالی طبق نتیجه جدول (۴) تعیین کننده الگوریتم پایه بر اساس روش آموزش و معیار اعتبار واری می‌باشد.

جدول ۴. اعتبارسنجی گزاره‌های پژوهش

مقادیر	MMC	f-measure	recall	precision	accuracy	Rank
الگوریتم‌های یادگیری نظارت شده یا نظارتی	۷۱/۲۹	۵۲/۱۹	۷۰/۵۴	۶۶/۴۵	۸۹/۶۶	3 rd
الگوریتم‌های یادگیری بدون نظارت	۸۳/۸۷	۷۰/۱۶	۶۵/۴۷	۸۹/۱۰	۹۲/۰۹	1 st
الگوریتم‌های تقویتی	۶۹/۳۰	۵۷/۱۷	۸۳/۱۱	۷۰/۳۶	۸۴/۸۱	2 nd

همان‌طور که مشاهده می‌شود، مقدار الگوریتم‌های یادگیری بدون نظارت معیار متناسب‌تری از استنتاج فازی عصبی انطباق‌پذیر جهت سنجش نسبت‌های تعیین‌شده ارزیابی احتمال تقلب مالی محسوب می‌شود. لذا با توجه به اینکه الگوریتم فرا ابتکاری، مهم‌ترین مبنای این نوع از فرآیندهای تحلیلی شبکه عصبی می‌باشد، در ادامه تلاش می‌شود تا نسبت به انتخاب بهترین الگوریتم فرا ابتکاری در پیش‌بینی تقلب مالی شرکت‌ها اقدام شود. لذا بر اساس تعاریف عملیاتی هریک از نسبت‌های مالی مؤثر در پیش‌بینی تقلب مالی، مدلی به ترتیب زیر با ضرایب الگوریتم فرا ابتکاری که پایایی آن در پژوهش Hidayattullah et al. (2020) مورد تأیید قرار گرفته است، به ترتیب زیر ارائه نمود:

$$Sta(\sigma) = \sum_{i=1}^n \int_{i=1}^n \varepsilon_i \ln \prod_{j=1}^n Activ index \quad \text{رابطه (۲)}$$

در این رابطه؛ $Sta(\sigma)$ ضریب تقلب مالی است. $\sum_{i=1}^n Activ index$ مجموع نسبت‌های مالی شناسایی‌شده جهت ارزیابی تقلب شرکت‌ها می‌باشد. « $Activ index$ » اشاره به صحت داده‌های موردبررسی در بازه زمانی « t » می‌باشد. ε_i ضریب محاسبه‌ی نسبت‌های ارزیابی تقلب مالی « $Activ index$ » می‌باشد. $\prod_{j=1}^n Activ index$ ادغام عملکرد مجموع ۲۴ نسبت ارزیابی تقلب مالی می‌باشد. در ادامه می‌بایست به پیروی از مطالعه Ngai et al. (2011)، منابع تجمعی مالی شرکت‌های i در سال t ، با عنوان $Sta_{i,t}$ ، به صورت رابطه (۳) تعریف شده است:

$$Sta_{i,t} = \text{Ln}[\sigma_{i,t} + \sum_1^t (1 - \gamma)^t \sigma_{i,t-t}] \quad \text{رابطه (۳)}$$

که در آن:

$Sta_{i,t}$ شاخص تقلب ناشی از تجميع منابع مالی شرکت‌های i در سال t ^۱ است و γ مجموعه نسبت‌های مالی ارزیابی تقلب شرکت‌ها می‌باشد^۲. با تعیین مدل برای ورودی داده‌ها به الگوریتم فرا ابتکاری می‌بایست با تعیین دو بازه ۰ و ۱، مجموع معیارهای سنجش را بر مبنای تعیین تقلب یا عدم تقلب مالی شرکت‌ها محاسبه نمود. برای این منظور از شاخص ترکیبی «Composite Index» (CI) استفاده می‌شود.

برای پیاده‌سازی چنین فرآیندی، لازم بود تا نسبت به تعیین اوزان هریک از نسبت‌های قرار گرفته در شاخص معیارهای شاخص ایجاد شده «CI» به گونه‌ای اقدام شود تا امکان ایجاد یک معیار مشخص برای ارزیابی تقلب مالی شرکت‌ها در قالب یک الگوی شبکه عصبی ممکن باشد. براین اساس، این فرآیند با ترکیب معیارهای استخراج شده در مرحله اول، امکان برآزش مطلوبیت الگوریتم‌های منتخب را برای تشخیص تقلب می‌دهد؛ بنابراین ۲۴ نسبت مالی ارزیابی تقلب انتخاب شده، مبنای تعیین تفاوت بین شرکت‌های متقلب با شرکت‌های دیگر قرار می‌گیرد. لذا رابطه (۴) برای مشترک کردن شاخص «CI» استفاده می‌شود.

$$Y = F(x_1, x_2, \dots, x_n) \quad (۴) \text{ رابطه}$$

که در این رابطه؛

Y متغیر وابسته یا نماد $Sta(\sigma)$ مبنی بر سنجش تقلب مالی است، تفاوت بین شرکت‌های دارای احتمال تقلب با شرکت‌های دیگر را نشان می‌دهد. به طوری که اگر نسبت‌های مالی تعیین شده نشان دهنده سلامت مالی بالاتر باشند، شاخص ترکیبی (CI) برابر یک و در شرایطی که نسبت‌های مالی محاسبه شده، احتمال تقلب را نشان دهند، در این صورت شاخص ترکیبی (CI) برابر صفر منظور می‌شود. به علاوه بایستی بیان شود، طبق روش تحلیل تشخیصی، x_i نشان دهنده متغیرهای توضیحی می‌باشد که هریک از پارامترهای ترکیب

۱ حرف t یک کلمه یونانی است که اشاره t سال در بررسی‌های داده‌های شرکت‌ها دارد.

۲ این مطالعه برای سنجش مقادیر γ و t از فروض گریلیشز (۱۹۸۴) مقدار ثابت $\gamma = 0.4$ و دوره زمانی $t = 10$ را در رابطه (۱) استفاده نموده است.

خطی را برای شناسایی شرکت‌های متقلب از طریق رابطه (۵) پیش‌بینی می‌کند:

$$D = V_1X_1 + V_2X_2 + \dots + V_iX_i \quad (۵) \text{ رابطه}$$

که در این رابطه؛ D تابع تشخیصی؛ V وزن هر معیار در تابع؛ X متغیرهای توضیحی در تابع و i تعداد متغیرهای توضیحی می‌باشد.

لذا بر اساس این روابط، می‌بایست با توجه به نسبت‌های مالی تعیین شده جهت تفکیک شرکت‌ها بر اساس مشاهدات (سال-شرکت) به دو گروه احتمال تقلب و سلامت مالی، ضرایب استاندارد با استفاده از روش تحلیل تشخیصی برآورد و در شاخص ترکیبی تعبیه شوند. لذا می‌بایست شاخص ترکیبی را بر اساس ۲۴ نسبت مالی شناسایی شده برای سنجش احتمال شرکت‌های دارای تقلب مالی، از طریق مجموع مشاهدات (شرکت-سال)، دهک‌بندی نمود تا تفاوت بین شرکت‌ها مشخص گردد؛ به عبارت دیگر می‌بایست از طریق دهک‌بندی مشاهدات (سال-شرکت) بر اساس خروجی داده کاوی نسبت‌های مالی، تفاوت بین شرکت‌ها از منظر تقلب و سلامت مالی مشخص شوند. طبق این الگو، مشاهدات با امتیاز پایین تر از نظر نسبت‌های سلامت مالی در دهک اول و مشاهدات با بالاترین سطح امتیاز نسبت‌های مالی تعیین شده، در دهک آخر قرار می‌گیرند. در ادامه نیز می‌بایست تمامی رتبه‌های دهک‌های تعیین شده در شاخص «CI» بر کل نسبت‌های مالی تقسیم گردد تا قرار گرفتن مشاهدات در رتبه‌های یک تا ده تعیین‌کننده‌ی تفاوت مبنای احتمال تقلب با سلامت مالی باشد. به طوری که مقادیر بزرگ‌تر (کوچک‌تر) بیانگر میزان سلامت مالی بالاتر (پایین‌تر) در مشاهدات مورد بررسی تلقی می‌گردد. استفاده از این شاخص ترکیبی، همچنین چولگی ناشی از محاسبه‌ی هریک از نسبت‌های مالی قابل سنجش را کاهش داده و معیار دقیق‌تری را برای آزمون فراهم می‌آورد.

جدول ۵. محاسبه شاخصه ترکیبی تعیین تقلب از سلامت مالی

تکرار مشاهدات	مجموع شاخص	$D \Rightarrow 1392$	$D \Rightarrow 1393$	$D \Rightarrow 1394$	$D \Rightarrow 1395$	$D \Rightarrow 1396$	$D \Rightarrow 1397$	$D \Rightarrow 1398$	$D \Rightarrow 1399$	$D \Rightarrow 1400$	$D \Rightarrow 1401$	(MAX θ)	دهک‌ها
۴۳۵	۱/۰۰۰	۰/۰۹۷	۰/۲۷۳	۰/۲۱۷	۰/۲۲۱	۰/۲۲۳	۰/۰۹۱۸	۰/۱۸۷	۰/۱۶۲۷	۰/۱۰۶	۰/۸۰۷	خفنی	دهک اول
۱۰۹	۱/۰۷۶	۱/۰۰	۱/۰۰	۰/۱۱۶	۰/۸۰۲	۰/۴۰۳	۱/۰۰	۰/۴۷۱	۰/۴۴۲۲	۱/۰۰	۰/۹۸۳	مطلوب	دهک دوم
۸۷	۱/۳۳۲	۰/۴۲۴	۰/۰۹۸	۱/۰۰	۱/۰۰	۰/۵۴۹	۰/۷۹۴	۰/۹۲۲	۰/۷۴۹۵	۱/۰۰	۰/۵۵۶	مطلوب	دهک سوم
۱۴	۰/۸۹۸	۰/۸۴۷	۱/۰۰	۰/۵۴۲	۰/۴۱۱	۰/۶۲۵	۰/۱۱۱	۰/۲۳۹	۰/۱۳۲	۰/۲۸۵	۰/۳۱۷	بحرانی	دهک چهارم
۱۲۳	۱/۰۰۲	۰/۳۵۲	۰/۳۲۴	۰/۳۱۲	۱/۰۰	۰/۷۱۲	۰/۹۸۵	۰/۸۶۴	۱/۰۰	۰/۳۴۹	۰/۳۲۹	مطلوب	دهک پنجم
۳۱	۰/۰۷۴	۰/۹۱۱	۰/۲۲۲	۰/۳۰۵	۰/۳۲۶	۰/۳۳۷	۰/۲۴۵	۰/۲۱۸	۰/۴۲۴	۰/۱۹۶	۰/۱۲۹	بحرانی	دهک ششم
۶۵	۱/۸۳۹	۰/۷۴۲	۱/۰۰	۱/۰۰	۰/۹۲۸	۰/۷۴۳	۱/۰۰	۱/۰۰	۱/۰۰	۰/۶۳۵	۰/۷۶۳	عالی	دهک هفتم
۲۶	۰/۷۸۷	۰/۲۲۹	۰/۳۴۳	۰/۷۲۹	۰/۲۲۲	۰/۲۱۰	۰/۴۱۴	۰/۵۳۷	۰/۳۰۱	۰/۵۴۶	۰/۲۹۱	بحرانی	دهک هشتم
۴۴	۱/۴۱۷	۱/۰۰	۰/۹۳۶	۰/۸۳۲	۰/۴۵۵	۰/۵۲۴	۰/۳۷۵	۱/۰۰	۰/۲۸۳	۱/۰۰	۱/۰۰	عالی	دهک نهم
۱۶	۰/۴۸۸	۰/۹۵۵	۰/۶۱۷	۰/۳۸۲	۰/۱۱۹	۰/۵۲۷	۰/۷۲۹	۰/۳۳۲	۰/۲۳۵	۰/۱۸۲	۰/۰۹۸	بحرانی	دهک دهم

نتایج این جدول نشان می‌دهد، به دلیل اینکه مشاهدات ارزیابی شده، حکایت از قرار گرفتن دهک‌های چهارم؛ ششم؛ هشتم و دهم برابر با « $D > 1$ » با تعداد مجموع مشاهدات ۸۷ احتمال دارای ریسک تقلب نسبت به سایر مشاهدات در دهک‌های دیگر با ۸۶۳ مشاهده می‌باشد، در ادامه می‌بایست تفاوت بین شرکت‌های دارای سلامت مالی با شرکت‌های دارای احتمال تقلب را از طریق آزمون‌های باکس، لامبدای ویلکز و تشخیص کلونی، از نظر اعتبار مدل

موردبررسی قرار داد که نتایج آن در جدول (۶) مشخص شده است.

جدول ۶. نتایج آزمون باکس و لامبدای ویلکز و تابع تحلیل تشخیصی کانونی

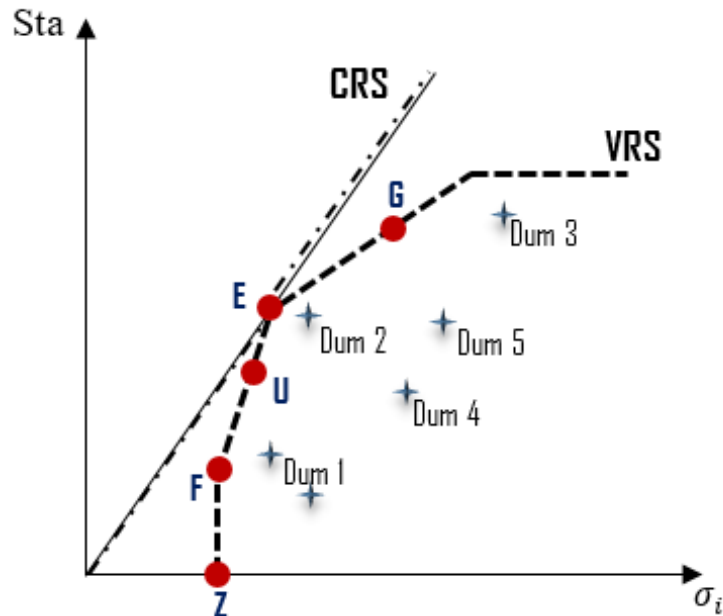
نوع آزمون	مقدار آماره		سطح معناداری
باکس	۶۰/۳۰۹		۰/۰۰۰
لامبدای ویلکز	لامبدا	Chi-Square	۰/۰۰۰
	۰/۳۱۹	۱۹۳۷/۱۱	
تعداد مشاهدات (سال-شرکت) دارای احتمال سلامت مالی مطلوب یا متوسط از مجموع ۹۵۰ مشاهده			۸۶۳
تعداد مشاهدات (سال-شرکت) دارای احتمال تقلب مالی از مجموع ۹۵۰ مشاهده			۸۷
درصد پیش‌بینی صحیح مشاهدات (سال-شرکت) دارای احتمال سلامت مالی			۹۰/۸۴
درصد پیش‌بینی صحیح مشاهدات (سال-شرکت) دارای احتمال تقلب مالی			۹/۱۶
درصد پیش‌بینی صحیح احتمال تقلب مالی بر اساس تعداد مشاهدات			۰/۸۱۶
ضریب همبستگی کلونی			۰/۸۱۶

همان‌طور که نتایج نشان می‌دهد، ماتریس کواریانس دو گروه تفاوت معناداری با یکدیگر دارند. همچنین مقدار لامبدای ویلکز نیز بیانگر این است که نسبت مجذورات درون گروهی به کل مجموع مجذورات کوچک بوده و میانگین دو گروه با یکدیگر تفاوت دارند. این مقادیر از قدرت تشخیص بالای تابع حکایت داشته و نشان می‌دهد که مدل تحلیل تشخیصی برآورد شده از لحاظ آماری معنادار است. به علاوه مقدار ضریب همبستگی تابع تحلیل تشخیص کلونی که برابر با ۰/۸۱۶ می‌باشد، نشان می‌دهد، صحت درصد تفاوت بین سلامت مالی شرکت‌ها با شرکت‌های دارای احتمال تقلب برابر با ۸۱/۶ درصد می‌باشد و این به معنای آن است که معیارهای شناسایی شده، از قابلیت مطلوبی برای ارزیابی شرکت‌های دارای احتمال تقلب برخوردار هستند، چراکه صحت پیش‌بینی صورت گرفته بیشتر از ۵۰ درصد می‌باشد. درواقع از آنجایی که، مبنای تعیین درصد احتمال تقلب با سلامت مالی، مشاهدات مطالعه می‌باشد، لذا بر اساس دهک‌بندی مشخص شد، از مجموع ۹۵۰ مشاهده، تعداد مشاهدات دارای احتمال سلامت مالی مطلوب یا متوسط برابر با ۸۶۳ مورد ($\theta = 863/950$) با درصد پیش‌بینی ۹۰/۸۴٪ و تعداد مشاهدات دارای

احتمال تقلب مالی برابر با ۸۷ مورد ($\theta = 87/950$) با درصد پیش‌بینی ۹/۱۶٪ هستند که تفاضل یا حدها فصل پیش‌بینی در دو گروه ($\theta = 90.84\% - 9.16\%$) حکایت از تعیین ۸۱/۶۸٪ صحت پیش‌بینی انجام شده بر اساس الگوی شاخص ترکیبی و داده کاوی مبتنی بر دهک‌بندی شرکت‌ها می‌باشد.

بنابراین می‌توان استنباط نمود، اگر مقادیر نسبت‌های مالی در شاخص مرجع مساوی یا بیش از ۱ ($\sigma \leq 1$) باشند، این به معنای آن است که احتمال سلامت مالی شرکت‌های موردبررسی بر اساس مشاهدات (سال-شرکت) بالاتر از حد متوسط بازار سرمایه است و این موضوع نشان‌دهنده‌ی پایین بودن درجه آسیب‌پذیری در برابر ریسک‌های احتمالی همچون ریسک ورشکستگی می‌باشد که به آن عدد ۱ تعلق می‌گیرد؛ اما اگر مقادیر نسبت‌های مالی شرکت‌های موردبررسی بر اساس مشاهدات (سال-شرکت) کوچک‌تر از ۰ ($\sigma > 0$) باشند، این به معنای آن است که احتمال سلامت مالی شرکت‌های موردبررسی بر اساس مشاهدات (سال-شرکت) پایین‌تر از حد متوسط بازار سرمایه است و این موضوع نشان‌دهنده‌ی بالا بودن درجه احتمال تقلب مالی شرکت‌ها و افزایش ریسک‌های احتمالی همچون ریسک ورشکستگی می‌باشد که به آن عدد ۰ تعلق می‌گیرد. لذا بر اساس شکل (۸) بر اساس دو محور ورودی و خروجی مدل، دو مبنای بازده ثابت نسبت مقیاس (CRS) و بازده متغیر نسبت به مقیاس (VRS) مبنای تفکیک بین شرکت‌ها از نظر سلامت مالی با شرکت‌های دارای احتمال تقلب بر اساس مشاهدات (سال-شرکت) تلقی می‌شوند.

شکل ۸. مبنای شرکت‌ها از نظر احتمال تقلب با سلامت مالی



طبق شکل (۸) می‌توان دریافت با در نظر گرفتن مشاهده (سال-شرکت) i که در نقطه F فعالیت می‌کند، پایین بودن احتمال سلامت مالی، تحت بازده ثابت نسبت به مقیاس (CRS) با فاصله بین نقاط E و F (EF) نشان داده می‌شود. از طرف دیگر، تفاوت بین EF و UF ، یعنی EU ، نشان‌دهنده احتمال سلامت مالی بالاتر در نقطه G است؛ به عبارت دیگر، هر قدر نسبت‌های مالی ارزیابی احتمال تقلب در بازده متغیر نسبت به مقیاس (VRS) به نقطه‌ی بهینه حرکت کنند، این به معنای آن است که سطح مطلوبیت سلامت مالی بالاتر است، درحالی‌که، قرار گرفتن مشاهده (سال-شرکت) در حدفاصل نقاط « F » و « U » می‌تواند به عنوان پایین بودن مبنای سلامت مالی و افزایش احتمال تقلب مالی شرکت‌های مورد بررسی قلمداد گردد. لذا با توجه به توضیح‌های داده‌شده، داده‌های دهک‌های بحرانی که مبنای تعیین احتمال شرکت‌های متقلب هستند، می‌بایست برای تعیین دو الگوریتم مطلوب از میان الگوریتم‌های فرا ابتکاری انتخاب شوند؛ به عبارت دیگر، شرکت‌های قرار گرفته در دهک‌های با ضریب « $D > 1$ » مبنای بررسی پارامترهای الگوریتم فرا ابتکاری

می‌باشد. برای این منظور بر اساس توابع برآزش مختلف، هریک از نزدیک‌ترین الگوریتم‌های فرا ابتکاری با مقادیر بحرانی دهک‌های تعیین شده، نسبت به انتخاب اقدام نمود. نکته‌ی قابل توجه این است که باتوجه به اینکه هدف مطالعه در این بخش، دستیابی به مؤثرترین نسبت تعیین تقلب مالی شرکت‌ها بر اساس الگوریتم‌های شبکه عصبی یادگیری بدون نظارت می‌باشد که به‌عنوان پایه ارزیابی مبتنی بر صحت پیش‌بینی احتمال تقلب مالی شرکت‌ها در بخش اول تحلیل‌ها انتخاب شد، در این بخش از مجموعه‌ی الگوریتم فرا ابتکاری بهره برده می‌شود تا متناسب‌ترین الگوریتم‌های اجرایی برای ارزیابی احتمال تقلب بتواند مورد استفاده تصمیم گیرندگان مالی قرار گیرد.

برای این منظور می‌بایست، ابتدا نسبت به ایجاد تعداد «M» جواب تصادفی و فازی بر اساس دو معیار ارزش واقعی ایجاد نمود. تا بتوان نسبت به محاسبه تابع هدف و تعیین بهترین و بدترین جواب (X_{best}^R) اقدام لازم را انجام داد. سپس می‌بایست نسبت به تقسیم بازه‌های تعریف شده به صورت مساوی برای تعیین مطلوب‌ترین منطقه انجام تحلیل اقدام نمود. براین اساس اگر N جمعیت کل باشد، k برابر است با یک چهارم N. لذا $X_{best}^R = X_i^{gp}$ و $i = 1, 2, \dots, k$ و $p = 1, 2, 3, 4, \dots, X_i^{gp}$ به شرطی این معادله برقرار است که در فضای جواب، بهترین الگوریتم بالاتر از ۱ قرار بگیرد. لذا می‌بایست جمعیت اولیه (الگوریتم‌های بهینه‌سازی فرا ابتکاری) را در تابع تعریف شده قرار داد تا $j = 1, 2, \dots, N; X_j^I$ به‌طوری که هریک از الگوریتم‌های برآزش شده بر اساس مقیاس‌ها از طریق تابع $X_{best}^R = X_1^I \Rightarrow X_{best}^R = X_7^I$ مشخص می‌نمایند، کدام الگوریتم‌ها امکان دستیابی به مطلوبیت یکپارچه را دارند. لذا طبق جدول (۷) اقدام به ارزیابی می‌شود.

جدول ۷. مبانی انتخاب بهترین الگوریتم فرا ابتکاری

مبنای	الگوریتم ژنتیک*	الگوریتم گرگ خاکستری	الگوریتم بهینه‌سازی وال	الگوریتم کلونی زنبور عسل*	الگوریتم ازدحام ذرات
میانگین	۰/۰۷۲۸۳	۰/۰۲۴۲۲	۰/۰۳۳۹۲	۰/۰۸۸۷۶	۰/۰۲۹۷۱
بهترین نتیجه	۱/۱۸۷۱	۰/۰۶۱۰۹	۰/۰۷۳۶۹	۱/۲۰۰۸	۰/۰۶۸۲۴

مبنای	الگوریتم ژنتیک*	الگوریتم گرگ خاکستری	الگوریتم بهینه‌سازی وال	الگوریتم کلونی زنبور عسل*	الگوریتم ازدحام ذرات
بدترین نتیجه	۰/۰۶۳۵۴	۱/۲۲۰۱	۰/۳۰۹۱۱	۰/۰۵۹۲۸	۱/۲۸۴۱۳
انحراف معیار	۰/۰۱۵۵۷	۱/۵۶۸۱	۱/۲۱۲۷	۰/۰۱۹۱۱	۱/۶۳۰۴

بر اساس نتایج کسب‌شده، مشخص شد، دور الگوریتم ژنتیک و الگوریتم کلونی زنبور عسل، به دلیل کسب ضریب بالاتر از ۱ در معیار بهترین نتیجه و ضریب پایین‌تر از ۱ در معیار بدترین نتیجه، از کارکرد مطلوب‌تری برای تعیین صحت پیش‌بینی تقلب مالی شرکت‌های مورد بررسی برخوردار هستند. لذا در ادامه ابتدا هر الگوریتم به صورت مجزا ارزیابی می‌شود و سپس بر اساس دو معیار تعیین‌شده‌ی بازده ثابت نسبت مقیاس (CRS) و بازده متغیر نسبت به مقیاس (VRS) اقدام به انتخاب بهترین الگوریتم فرا ابتکاری برای تعیین قابلیت بالاتر نسبت‌های مالی در پیش‌بینی تقلب مالی می‌شود.

الگوریتم ژنتیک

برای اجرای الگوریتم ژنتیک از طریق آموزش داده‌ها اقدام به دسته‌بندی شرکت‌های قرار گرفته در دهک احتمال تقلب می‌گردد. لذا تابع هدفی که الگوریتم ژنتیک تلاش دارد تا با بهینه‌سازی آن، به تعیین ضرایب پیش‌بینی دست یابد طبق رابطه (۶) تا (۸) ارائه می‌شود.

$$L(B.X) = \frac{1}{1+e^{XB^t}} \quad \text{رابطه (۶)}$$

$$r = (B.X) = \begin{cases} 1 & \text{if } L(B.X) > 0.5 \\ 0 & \text{if } L(B.X) \leq 0.5 \end{cases} \quad \text{رابطه (۷)}$$

$$\Omega = (B.X.Y) = \{\sum_{i=1}^n (Y - \Gamma_i[B.X])\} / N \quad \text{رابطه (۸)}$$

در این روابط، تابع «L» تابعی از مبنای لجستیک می‌باشد که به صورت ماتریسی محاسبه می‌شود. همچنین تابع « Γ » تعیین‌کننده‌ی اختصاص خروجی به صفر یا یک می‌باشد و

در نهایت تابع « Ω » مبنای هزینه‌ای است که الگوریتم به دنبال بهینه‌یابی آن است؛ بنابراین، الگوریتم ژنتیک با تابع هزینه فوق به دنبال کاهش مجموع اختلاف خطای دسته‌بندی در داده‌های آموزشی شرکت‌هایی است که در دهک‌بندی‌های انجام‌شده، از احتمال تقلب مالی برخوردار می‌باشند. لذا در جدول (۸) خروجی‌های ناشی از داده‌های آموزش مربوط به این الگوریتم ارائه شده است.

جدول ۸. خروجی داده‌های آموزش مربوط به الگوریتم ژنتیک

معیارهای ارزیابی صحت داده‌های آموزش	دهک چهارم	دهک ششم	دهک هشتم	دهک دهم
درصد طبقه‌بندی صحیح داده‌های آموزش	۶۰/۱۸	۶۳/۳۷	۶۵/۱۰	۵۹/۴۲
درصد طبقه‌بندی داده‌های اعتبارسنجی	۷۱/۲۳	۷۳/۸۱	۷۵/۶۴	۷۰/۳۸
درصد طبقه‌بندی صحیح داده‌های آزمون	۴۹/۵۱	۴۱/۰۳	۵۰/۵۳	۴۲/۳۳
درصد طبقه‌بندی کل داده‌ها	۷۴/۲۹	۷۵/۲۹	۷۶/۲۱	۷۲/۲۸

جدول (۸) نشان می‌دهد، اجرای داده‌های آموزش ناشی از الگوریتم ژنتیک برای صحت پیش‌بینی تقلب مالی شرکت‌های دهک‌بندی شده مشخص می‌نماید که تمامی دهک‌ها از درصد بالایی برای پیش‌بینی صحت تقلب مالی شرکت‌ها برخوردار هستند. لذا در ادامه و بر اساس پارامترهایی همچون نرخ همبندی^۱ و نرخ جهش^۲ توابع الگوریتم ژنتیک در پیش‌بینی صحیح احتمال تقلب بر اساس طبقه‌بندی صحیح داده‌های آموزش مورد محاسبه قرار می‌گیرد.

1 Crossover Percentage

2 Mutation Percentage

شکل ۹. نتایج ناشی از پیش‌بینی احتمال تقلب بر اساس بر اساس طبقه‌بندی صحیح داده‌های آموزش

Output Class	1	163 38.72%	333 61.78%	74.11% 23.06%
	2	206 38.22%	258 62.28%	71.23% 24.06%
		76.94% 23.44%	74% 21%	73.51% 21.34%
		1	2	
		Target Class		

پارامترهای نهایی به‌دست‌آمده نشان‌دهنده‌ی تعداد ۴۲۱ مشاهده از مجموع ۹۶۰ مشاهده برای صحت پیش‌بینی تقلب در شرکت‌های موردبررسی می‌باشد که با نرخ همبندی بین ۷۰٪ تا ۷۵٪ و نرخ جهش بین ۲۱٪ تا ۲۴٪ می‌توان اقدام به صحت پیش‌بینی تقلب بر اساس نسبت‌های مالی انتخاب شده نمایند. لذا این نتایج نشان می‌دهد، بالاترین صحت پیش‌بینی مربوط به شرکت‌های دارای تقلب ۶۲/۲۸٪ می‌باشد که درصد احتمالی داده‌های کنترل برابر با ۳۸/۷۲٪ می‌باشد. لذا تعداد ۲۵۸ شرکت قرارگرفته در دهک‌بندی‌های دارای احتمال بالا در تقلب، با صحت پیش‌بینی ۶۲/۲۸٪ می‌تواند به‌درستی توسط نسبت‌های مالی تعیین‌شده، شناسایی شوند.

جدول ۹. پارامترهای آماری تابع هدف در روش الگوریتم ژنتیک

ضریب تغییرات	انحراف معیار	متوسط	کمینه	بیشینه	پارامتر
۰/۰۱۷۷	۰/۰۳۲۰۱	۱/۳۲۷	۱/۲۰۱	۱/۶۶۱	CRS
۰/۲۰۶	۱/۵۸۳	۱۳/۳۶۵	۱۲/۴۳۷	۱۵/۱۰۱	VRS

ضرایب ارائه‌شده، نشان از تأیید پارامترهای هریک از نسبت‌های مالی مؤثر در پیش‌بینی

صحت تقلب مالی شرکت‌ها بر اساس الگوریتم ژنتیک می‌باشد.

الگوریتم کلونی زنبورعسل

این الگوریتم بر اساس موقعیت اصطلاحی شناسایی کرده گیاهان توسط زنبورعسل، به دنبال بهینه‌ترین راه‌حل در انتخاب یک پاسخ محتمل می‌باشد. لذا این الگوریتم از طریق رابطه (۹) قابل بسط در این مطالعه می‌باشد.

$$p_i = \frac{fit_i}{\sum_{n=1}^{sn} fit_n} \quad \text{رابطه (۹)}$$

در این رابطه؛ « fit_i » مطلوبیت الگوریتمی بر اساس ارزیابی راه‌حل i تعیین می‌شود، که n ارزیابی متناسب با مقدار شاهد منبع غذایی در موقعیت « i » می‌باشد. « sn » تعداد منابع غذایی است که برابر با تعداد زنبورهای کارگر «BN» است. لذا در جدول (۱۰) خروجی‌های ناشی از داده‌های آموزش مربوط به این الگوریتم ارائه شده است.

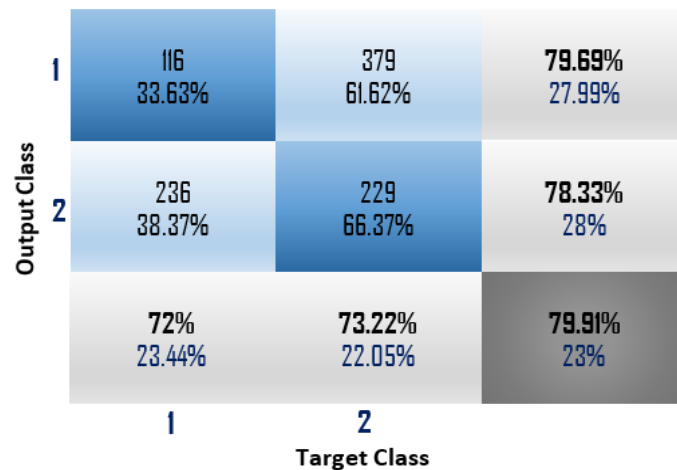
جدول ۱۰. خروجی داده‌های آموزش مربوط به الگوریتم کلونی زنبورعسل

معیارهای ارزیابی صحت داده‌های آموزش	دهک چهارم	دهک ششم	دهک هشتم	دهک دهم
درصد طبقه‌بندی صحیح داده‌های آموزش	۵۶/۲۲	۶۲/۶۶	۶۵/۳۵	۵۰/۴۶
درصد طبقه‌بندی داده‌های اعتبارسنجی	۵۹/۰۳	۶۴/۰۷	۶۲/۲۹	۵۴/۹۸
درصد طبقه‌بندی صحیح داده‌های آزمون	۳۹/۱۷	۵۸/۲۲	۶۷/۷۸	۴۸/۸۲
درصد طبقه‌بندی کل داده‌ها	۵۸/۱۱	۶۹/۱۸	۶۶/۱۹	۵۷/۷۳

این جدول نشان می‌دهد، اجرای داده‌های آموزش ناشی از الگوریتم کلونی زنبورعسل برای صحت پیش‌بینی تقلب مالی شرکت‌های دهک‌بندی شده مشخص می‌نماید که تمامی دهک‌ها از درصد متناسبی برای پیش‌بینی صحت تقلب مالی شرکت‌ها برخوردار هستند. لذا در ادامه و بر اساس پارامترهایی همچون حد فراموشی «Abandonment Limit» و حداکثر ضریب شتاب «Acceleration Coefficient Upper» توابع الگوریتم کلونی زنبورعسل در پیش‌بینی صحیح احتمال تقلب بر اساس طبقه‌بندی صحیح داده‌های آموزش

و داده‌های کنترل مورد محاسبه قرار می‌گیرد.

شکل ۱۰. نتایج ناشی از پیش‌بینی احتمال تقلب بر اساس بر اساس طبقه‌بندی صحیح داده‌های آموزش



پارامترهای نهایی به دست آمده نشان‌دهنده‌ی تعداد ۳۴۵ مشاهده از مجموع ۹۶۰ مشاهده برای صحت پیش‌بینی تقلب در شرکت‌های مورد بررسی می‌باشد که با حد فراموشی بین ۷۸٪ تا ۷۹٪ و حداکثر ضریب شتاب بین ۲۳٪ تا ۲۸٪ می‌توان اقدام به صحت پیش‌بینی تقلب بر اساس نسبت‌های مالی انتخاب شده نمایند. لذا این نتایج نشان می‌دهد، بالاترین صحت پیش‌بینی مربوط به شرکت‌های دارای تقلب ۶۶/۳۷٪ می‌باشد که درصد احتمالی داده‌های آزمون کنترل برابر با ۳۳/۶۳٪ می‌باشد. لذا تعداد ۲۲۹ شرکت قرار گرفته در دهک‌بندی‌های دارای احتمال بالا در تقلب، با صحت پیش‌بینی ۶۶/۳۷٪ می‌تواند به درستی توسط نسبت‌های مالی تعیین شده، شناسایی شوند.

جدول ۱۱. پارامترهای آماری تابع هدف در روش الگوریتم کلونی زنبور عسل

پارامتر	پیشینه	کمینه	متوسط	انحراف معیار	ضریب تغییرات
CRS	۱/۵۲۷	۱/۲۹۱	۱/۳۱۹	۰/۰۴۳۱	۰/۰۱۹۸
VRS	۱۴/۴۴۳	۱۳/۵۶۱	۱۳/۷۳۸	۱/۷۱۲	۰/۲۲۸

ضرایب ارائه‌شده، نشان از تأیید پارامترهای هریک از نسبت‌های مالی مؤثر در پیش‌بینی صحت تقلب مالی شرکت‌ها بر اساس الگوریتم کلونی زنبورعسل می‌باشد. لذا نتایج به‌دست آمده از هر دو الگوریتم، نشان‌دهنده توازن نسبی صحت برابر پیش‌بینی احتمال تقلب مالی شرکت‌های دهک‌بندی‌شده می‌باشد. لذا در گام آخر می‌بایست بر اساس دو مقیاس ایجادشده، مشخص نمود، کدام نسبت مالی و کدام الگوریتم از ظرفیت‌های بالاتری بر کمک به تصمیم‌گیرندگان مالی برای پیش‌بینی صحیح تقلب در شرکت‌های موردبررسی برخوردار هستند. برای این منظور از آزمون ویلکا کسون و ارزیابی ضرایب هریک از نسبت‌های مالی بر در الگوریتم‌های فرا ابتکاری بهره برده می‌شود.

جدول ۱۲. نتایج ارزیابی الگوریتم‌های فرا ابتکاری بر اساس نسبت‌های ارزیابی تقلب مالی

ردیف	معیارهای شناسایی شده	Code	مقیاس	الگوریتم ژنتیک	الگوریتم کلونی زنبورعسل
۱	نسبت جاری	σ_1	CRS ۲۱۴.۲	۲/۴۸۷	۴/۴۳۸
			VRS ۴۲۴.۱۰		
۲	نسبت دارایی جاری به کل دارایی‌ها	σ_2	CRS ۹۶۹.۱	۲/۶۵۹	۳/۲۸۱
			VRS ۷۰۷.۱۱		
۳	نسبت دارایی ثابت به کل دارایی‌ها	σ_3	CRS ۲۶۹.۲	۳/۲۸۵	۵/۰۲۸
			VRS ۸۲۵.۱۲		
۴	نسبت سود انباشته به کل دارایی‌ها	σ_4	CRS ۵۸۵.۱	۲/۵۲۷	۴/۳۳۲
			VRS ۲۶۸.۱۲		
۵	نسبت سرمایه در گردش به کل دارایی‌ها	σ_5	CRS ۷۶۱.۱	۳/۱۴۴	۳/۱۹۲
			VRS ۴۶۳.۱۰		
۶	نسبت کل بدهی به کل دارایی‌ها	σ_6	CRS ۸۵۲.۱	۳/۲۳۲	۴/۱۷۲
			VRS ۸۶۱.۱۱		
۷	نسبت سود ناخالص به کل دارایی‌ها	σ_7	CRS ۰۸۳.۲	۳/۸۷۴	۴/۰۴۵
			VRS ۵۷۶.۱۱		
۸	نسبت کل بدهی به حقوق صاحبان سهام	σ_8	CRS ۹۹۴.۱	۲/۹۵۷	۳/۲۲۸
			VRS ۰۱۲.۱۳		
۹	نسبت سود انباشته به کل	σ_9	CRS ۰۷۶.۲	۳/۴۱۴	۴/۱۹۲

ردیف	معیارهای شناسایی شده	Code	مقیاس	الگوریتم ژنتیک	الگوریتم کلونی زنبور عسل
	دارایی‌ها		۵۳۱.۱ VRS		
۱۰	نسبت فروش به کل دارایی‌ها	σ_{10}	۱۹۳.۱۲ CRS ۴۶۶.۱ VRS	۴/۰۱۹	۵/۳۸۲
۱۱	نسبت بدهی بلندمدت به حقوق صاحبان سهام	σ_{11}	۴۸۹.۱۰ CRS ۷۳۸.۱ VRS	۳/۲۹۱	۴/۸۸۷
۱۲	نسبت بهای تمام شده کالای فروش رفته به کل فروش	σ_{12}	۰۹۱.۲ CRS ۴۵۵.۱۳ VRS	۲/۷۸۸	۳/۱۹۵
۱۳	نسبت سود خالص به کل دارایی‌ها	σ_{13}	۲۶۳.۲ CRS ۵۷۴.۱۱ VRS	۳/۰۶۳	۳/۸۱۶
۱۴	نسبت سود ناخالص به فروش	σ_{14}	۲۷۸.۲ CRS ۲۸۴.۱۲ VRS	۴/۱۱۰	۶/۰۰۸
۱۵	نسبت سود عملیاتی به فروش	σ_{15}	۵۴۶.۲ CRS ۵۸۸.۱۲ VRS	۲/۵۶۴	۳/۲۲۸
۱۶	نسبت سود خالص به حقوق صاحبان سهام	σ_{16}	۱۹۸.۱ CRS ۴۹۶.۲ VRS	۳/۹۸۵	۵/۷۸۶
۱۷	نسبت هزینه‌های مالی به کل بدهی‌ها	σ_{17}	۱۵.۲ CRS ۹۷.۱۲ VRS	۳/۰۹۱	۴/۱۱۱
۱۸	نسبت سود خالص به فروش	σ_{18}	۲۲۴.۲ CRS ۵۶۵.۱۲ VRS	۴/۱۶۲	۶/۰۷۱
۱۹	نسبت هزینه‌های عملیاتی به فروش	σ_{19}	۳۷۷.۱ CRS ۷۹۵.۱۲ VRS	۳/۱۷۱	۴/۲۱۷
۲۰	نسبت حساب‌های دریافتی به فروش	σ_{20}	۹۱۴.۱ CRS ۱۵۴.۱۲ VRS	۳/۲۸۴	۴/۹۲۱
۲۱	نسبت موجودی‌ها به فروش	σ_{21}	۹۱۲.۱ CRS ۳۵.۱۳ VRS	۳/۴۰۵	۴/۲۲۵
۲۲	نسبت موجودی نقد به کل دارایی‌ها	σ_{22}	۳۴۵.۱ CRS ۶۳۷.۱۱ VRS	۴/۱۲۸	۶/۰۴۶
۲۳	نسبت سود خالص به سود ناخالص	σ_{23}	۱۶.۱ CRS ۳۱۵.۱۲ VRS	۲/۸۱۶	۳/۳۲۱

الگوی تشخیصِ تقلب مالی تحتِ اجرای ارزیابی الگوریتم‌های ...؛ پورساعدی و همکاران | ۱۲۵

ردیف	معیارهای شناسایی شده	Code	مقیاس	الگوریتم ژنتیک	الگوریتم کلونی زنبور عسل
۲۴	نسبت جریان وجه نقد عملیاتی به کل دارایی‌ها	σ_{24}	CRS	۲/۷۳۱	۴/۸۷۱
			VRS	۰/۱۸.۱۳	

همان‌طور که دو مقیاس بازده ثابت نسبت مقیاس (CRS) و بازده متغیر نسبت به مقیاس (VRS) نشان می‌دهد، تمامی ضرایب بالاتر از ۱ می‌باشد که نشان می‌دهد، هر دو الگوریتم از سطح بهینگی لازم در پیش‌بینی احتمال تقلب مالی شرکت‌ها برخوردار می‌باشند؛ اما ضرایب الگوریتم‌های فرا ابتکاری نشان می‌دهد، کلونی زنبور عسل نسبت به الگوریتم ژنتیک از دقت بالاتری برخوردار می‌باشد. همچنین مشخص شد، نسبت سود خالص به فروش « σ_{18} » با ضریب « $۶/۰۷۱$ » بالاترین نسبت تأثیرگذار در پیش‌بینی صحت مالی شرکت‌های دهک‌بندی شده در هر دو الگوریتم کلونی زنبور عسل و الگوریتم ژنتیک را دارد. در ادامه جهت صحت این نتایج از آزمون ویلکا کسون بهره برده می‌شود تا بر اساس مجموع ضرایب نسبت‌های مالی مقایسه‌ی دو الگوریتم مورد توجه در این مطالعه، از حیث دقت آزمون طبق جدول (۱۳) مورد بررسی قرار گیرد.

جدول ۱۳. تفاوت ارزیابی الگوریتم‌های فرا ابتکاری شبکه‌عصبی بر اساس نسبت‌های ارزیابی تقلب

الگوریتم	تعداد ارزیابی تابع هدف	مقدار تابع کلی	میانگین ضرایب	مجموع ضریب	نتیجه ارزیابی
الگوریتم ژنتیک	۲۹۰۰۰	-۰/۰۰۰۲۹۵	۱۱/۱۹	۰/۳۰۱	Bee > GEN
الگوریتم کلونی زنبور عسل	۲۹۰۰۰	۰/۰۰۰۰۵۶	۱۲/۵۵	۰/۳۶۶	
		Z		-۰/۳۸۲	
		Asymp. Sig. (2-tailed)		۰/۰۰۰	

همان‌طور که بر اساس نتایج آزمون ویلکا کسون مشاهده می‌شود، ضریب معنی‌داری

به دست آمده با توجه به اینکه کوچک تر از ۰/۰۵ می باشد، می تواند بیان کننده ی این مسئله باشد که الگوریتم کلونی زنبور عسل نسبت به الگوریتم ژنتیک، به لحاظ آماره الگوریتمی، با دقت بالاتری در پیش بینی صحت احتمال تقلب شرکت های مورد بررسی همراه می باشد.

بحث و نتیجه گیری

هدف این مطالعه، مدل سازی تشخیص تقلب مالی شرکت ها تحت ارزیابی الگوریتم های مطلوبیت شبکه عصبی مصنوعی می باشد. در این مطالعه ابتدا با مرور نظام مند ادبیات مرتبط با حوزه پژوهش در طی چند سال گذشته، تلاش شد تا نسبت های مالی متناسب با ارزیابی تقلب شرکت های بازار سرمایه انتخاب شدند. سپس بر اساس برنامه ریزی مجذوری «QP» تلاش شد تا با حذف فاصله بین داده های ناشی از پیش بینی و واقعی حاصل از نسبت های مالی شناسایی شده، الگوریتم پایه ارزیابی صحت تقلب شرکت ها بر اساس مشاهدات (سال-شرکت) تعیین گردد. لذا طی آزمون منطق استنتاج فازی عصبی انطباق پذیر (ANFIS) و فرآیند اعتبارسنجی ضربدری یا اصطلاحاً «k-fold»، مشخص شد، الگوریتم یادگیری بدون نظارت که شامل مجموعه پارامترهای ارزیابی مبتنی بر الگوریتم فرا ابتکاری است و از صحت پیش بینی های مبتنی بر داده های برآورده شده بالاتری برخوردار می باشد، به عنوان پایه ی الگوریتم های این مطالعه انتخاب شد. در ادامه به منظور ایجاد شاخص مرجع، بر اساس نسبت های مالی ارزیابی احتمال تقلب شرکت ها از فرآیند دهک بندی بهره برده شد تا مشخص شود، چه شرکت هایی از دهک های دارای ضریب « $D > 1$ » می توانند در مقیاس های ارزیابی آزمون و کنترل، یعنی بازده ثابت نسبت مقیاس (CRS) و بازده متغیر نسبت به مقیاس (VRS) به تفکیک بین شرکت ها از نظر سلامت مالی با شرکت های دارای احتمال تقلب کمک نمایند. نتایج از قرار گرفتن شرکت های دارای احتمال تقلب در ۴ دهک حکایت دارد که برای ارزیابی های بیشتر صحت پیش بینی تقلب های مالی شرکت های دهک بندی شده از دو الگوریتم انتخابی ژنتیک و کلونی زنبور عسل بهره برده شد.

در نهایت نیز مشخص شد، الگوریتم کلونی زنبورعسل نسبت به الگوریتم ژنتیک، از ضریب دقت بالاتری در پیش‌بینی صحت احتمال تقلب شرکت‌های موردبررسی برخوردار است. همچنین مشخص شد، نسبت سود خالص به فروش مهم‌ترین معیار ارزیابی دقت پیش‌بینی احتمال تقلب در شرکت‌های موردبررسی می‌باشد.

در تفسیر نتایج کسب‌شده، ابتدا می‌بایست اشاره گردد که الگوریتم کلونی زنبورعسل به دلیل برخورداری از بهینه‌سازی فرآیندی چندگانه به واسطه‌ی هوش جمعی، از کارکردهای بالاتری در حل یک مسئله پیچیده برخوردار می‌باشد. درواقع این الگوریتم به دلیل اینکه از قدرت همگرایی و دقت بالاتری در پیش‌بینی احتمال تقلب شرکت‌های موردبررسی نسبت به الگوریتم ژنتیک، برخوردار می‌باشد می‌تواند کارایی اثربخش‌تری در تصمیم‌گیری‌های مالی در سطح بازار سرمایه داشته باشد. به‌علاوه ضرایب کسب‌شده در الگوریتم کلونی زنبورعسل، حکایت از بهینه‌سازی اثربخش‌تر نسبت‌های مالی در پیش‌بینی احتمال تقلب شرکت‌ها دارد. چراکه تصمیم‌گیرندگان مالی می‌توانند از این الگوریتم برای ارزیابی‌های دقیق‌تر مبتنی بر نسبت‌های مالی بهره ببرند تا بتوانند، شرکت‌های دارای احتمال تقلب را شناسایی کنند و در تشکیل پرتفوی سهام این شرکت‌ها را خریداری نکنند. از طرف دیگر مشخص گردید، شرکت‌های دارای نسبت سود خالص به فروش پایین‌تر می‌تواند معیار مناسبی در پیش‌بینی ضرایب الگوریتم کلونی زنبورعسل تلقی شود. چراکه این نسبت مالی این امکان را در ارزیابی تقلب شرکت‌ها مهیا می‌نماید که نشان‌دهنده‌ی این مسئله است که به دلیل عدم کنترل هزینه‌ها، شرکت‌های ظرفیت‌های مناسبی برای بالاتر بردن سود خالص از محل فروش را ندارند و این مسئله به دلیل اینکه واکنش منفی سرمایه‌گذاران را به همراه دارد، احتمال افزایش تقلب از طریق پوشش و دستکاری در حساب را در شرکت‌ها می‌تواند افزایش دهد. لذا کاهش این نسبت می‌تواند در ضریب الگوریتم کلونی زنبورعسل، مبنایی برای پیش‌بینی دقت تقلب در شرکت‌های موردبررسی تلقی گردد. همچنین ارجاع به بازده ثابت نسبت به مقیاس (CRS) و بازده متغیر نسبت به

مقیاس (VRS) نشان می‌دهد، کارکردهای عملیاتی شرکت‌های دارای احتمال تقلب بالاتر در یک بازار رقابتی، توانمندی‌های کسب سودهای خالص شرکت از محل فروش که بر اساس نوع مدیریت و بازده کارایی دارایی‌ها می‌تواند حادث شود را به تدریج از دست می‌دهند. لذا نکته‌ای که از مبانی تعیین ارزیابی تطبیقی بین الگوریتم‌ها یعنی بازده ثابت نسبت مقیاس (CRS) و بازده متغیر نسبت به مقیاس (VRS) می‌توان استنباط نمود این است که، شرکت‌های دارای احتمال تقلب مالی بالاتر، با بازدهی سرمایه پایین‌تری مواجه هستند که در بلندمدت می‌تواند ضمن تشدید تعهدات و نسبت اهرمی بالاتر برای شرکت، شکاف بین بازده مورد انتظار با بازده واقعی را نیز عمیق‌تر نماید. نتایج این مطالعه به لحاظ ماهیت اجرای الگوریتم‌های فرا ابتکاری و مقایسه بین اجزای آن اگرچه در پژوهش‌های گذشته مورد توجه نبوده است تا بتوان مقایسه‌ی دقیقی بین نتایج کسب‌شده با نتایج آن‌ها انجام داد، اما می‌توان به لحاظ ماهیت مفهومی و هم‌راستا با اهمیت به کارگیری الگوریتم‌های فرا ابتکاری نتایج این مطالعه را با پژوهش‌هایی همچون تشدید و همکاران (۱۳۹۸)؛ Zhou et al. (2024) و du Toit (2024) مورد مقایسه قرار داد.

بر اساس اهمیت مسئله تقلب در تصمیم‌گیری‌های مالی، به سرمایه‌گذاران توصیه می‌شود تا از طریق ارتقاء سطح آموزش تحلیل‌های تکنیکال در تشکیل سبد سهام، درصد ریسک‌های احتمالی پیش‌بینی تقلب مالی شرکت‌ها را کاهش دهند. مهم‌ترین شناخت تکنیک‌های تحلیلی می‌تواند مربوط به روش‌های پیش‌بینی باشد که بر پایه‌ی الگوریتم‌های یادگیری بدون نظارت ایجادشده باشد، چراکه تمرکز بر روی مجموع الگوریتم‌های این روش، شناخت بیشتری را در تصمیم‌گیرندگان ایجاد می‌نماید که با تمرکز بر روی ماهیت واقعی داده‌ها و بدون دستکاری یا نگاهت آن‌ها، مجموعه‌ای از الگوها، ساختارها و روابط قابل پیش‌بینی از تقلب مالی شرکت‌ها را می‌توانند شناسایی نمایند. همچنین با توجه به اینکه غالباً بخش مهمی از علل اثرگذار بر تقلب مالی شرکت‌ها، شکاف‌های قانونی هستند که امکان دستکاری و ساختگی‌سازی حساب‌ها را

به شرکت‌ها می‌دهند، ضرورت توسعه ظرفیت‌های بازرسی‌های قانونی دوره‌ای مبنی بر نظارت اثربخش‌تر بر شفافیت‌های مالی شرکت‌ها به حدی از ضرورت برخوردار است که غالباً بخش مهمی از عملکردهای متقلبانه می‌تواند به واسطه‌ی فقدان استانداردها یا نظارت‌های قانونی توسط شرکت‌ها به وجود بیاید.

تعارض منافع

تعارض منافع ندارم.

ORCID

Marzieh Poursaedi

Mahmood Hematfar

Seyed Enayatallah

Alavi

Roya Nasirzadeh



<https://orcid.org/0009-0006-3489-6513>



<https://orcid.org/0000-0002-3976-9322>



<https://orcid.org/0000-0003-0495-5704>



<https://orcid.org/0000-0002-5302-3789>

منابع

- آشاروزنیا، کتایون، پورآقاجان، عباسعلی، نسل موسوی، سیدحسین. (۱۴۰۲). رویکردی نوین به پیش‌بینی تقلب در صورت‌های مالی (مقایسه مدل‌های سنتی و شبیه‌سازی و مدرن)، تحقیقات حسابداری و حسابرسی، ۱۵(۵۷): ۱۴۱-۱۶۰. <https://doi.org/10.22034/IAAR.2023.172761>
- اعتمادی، حسین، عبدلی، لیلا. (۱۳۹۶). کیفیت حسابرسی و تقلب در صورت‌های مالی، دانش حسابداری مالی، ۴(۴): ۲۳-۴۳. https://jfak.journals.ikiu.ac.ir/article_1308.html?lang=fa
- بهرامی، آسو، نوروش، ایرج، راد، عباس، محمدملقرنی، عطاله. (۱۴۰۰). تقلب در صورت‌های مالی و تکنیک‌های نوین مورد استفاده جهت کشف آن، مطالعات حسابداری و حسابرسی، ۱۰(۳۸): ۱۰۵-۱۱۸. <https://doi.org/10.22034/iaas.2021.13454>
- تشدیدی، الهه، سپاسی، سحر، اعتمادی، حسین، آذر، عادل. (۱۳۹۸). ارائه رویکردی نوین در پیش‌بینی و کشف تقلب صورت‌های مالی با استفاده از الگوریتم زنبور عسل، مجله دانش حسابداری، ۱۰(۳): ۱۳۹-۱۶۷. <https://doi.org/10.22103/jak.2019.13616.2927>
- جعفری، محمد، رضایی، ندا، سلگی، محمد. (۱۴۰۲). استفاده از شبکه عصبی مصنوعی پیشخور چند لایه در تشخیص گزارشگری مالی متقلبانه در شرکت‌های پذیرفته شده در بورس اوراق بهادار تهران، مطالعات حسابرسی مظهر، ۱(۲): ۹۳-۱۱۳. https://masj.ihu.ac.ir/article_208557.html
- حسن‌پور، داود، ولیان، حسن، صفری‌گرایلی، مهدی، طهماسبی‌زاده، رضا. (۱۳۹۹). بکارگیری مدل‌های تئوری راف توسعه‌یافته (ERST) و تحلیل تفسیری-ساختاری (ISM) و درخت تصمیم (CART) برای کمک به حسابرسان جهت شناخت تقلب در صورت‌های مالی شرکت‌های پذیرفته شده در بورس اوراق بهادار ایران، دانش سرمایه‌گذاری، ۹(۳۳): ۱۷۹-۲۰۸. http://www.jik-ifea.ir/article_15665.html?lang=en
- دادگر، یداله، درگاهی، حسن، قلی‌زاده، سعید. (۱۴۰۲). نقش احساسات سرمایه‌گذاران و رفتار دولت در نوسانات بازار بورس اوراق بهادار تهران: رویکرد اقتصاد رفتاری، نظریه‌های کاربردی اقتصاد، ۱۰(۱): ۱۹۱-۲۱۴. <https://doi.org/10.22034/econj.2023.53000.3091>

ملکی کاکلر، حسن، بحری ثالث، جمال، جبارزاده کنگرلویی، سعید، آشتاب، علی. (۱۴۰۰).
کارایی مدل‌های آماری و الگوهای یادگیری ماشین در پیش‌بینی گزارشگری مالی
متقلبانه، نشریه اقتصاد مالی، ۱۵(۵۴): ۲۶۷-۲۹۲. <https://doi.org/20.1001.1.25383833.1400.15.54.11.2>

References

- Blakley, M. (2009). Fraud detection using a database platform. Available online at: <http://www.slideshare.net/mblakley>
- Bockel-Rickermann, Ch., Verdonck, T., & Verbeke, W. (2023). Fraud analytics: A decade of research: Organizing challenges and solutions in the field. *Expert Systems with Applications*, 232(2), 65-89. <https://doi.org/10.1016/j.eswa.2023.120605>
- Craja, P., Kim, A., & Lessmann, S. (2020). Deep learning for detecting financial statement fraud, *Decision Support Systems*, 139(2), 47-71. <https://doi.org/10.1016/j.dss.2020.113421>
- Dezoort, F. T., & Lee, Th. A. (1998). The Impact of SAS No. 82 on Perceptions of External Auditor Responsibility for Fraud Detection. *International Journal of Auditing*, 2(2), 167-182. <https://doi.org/10.1111/1099-1123.00037>
- du Toit, E. (2024). The red flags of financial statement fraud: a case study. *Journal of Financial Crime*, 31(2), 311-321. <https://doi.org/10.1108/JFC-02-2023-0028>
- Foster, M. (2015). Toshiba scandal sheds harsh light on Japan's corporate governance. <https://www.theguardian.com/business/2015/jul/21/toshiba-scandal-sheds-harsh-light-on-japans-corporate-governance>
- Francis, S. A., & Justine, C. (2023). Corruption and National Development: The Consequential Effects on the Sustainability in Ghana. *American Journal of Multidisciplinary Research and Innovation*, 2(3), 32-40. <https://doi.org/10.54536/ajmri.v2i3.1434>
- He, H., Shang, Y., Yang, X., Di, Y., Lin, J., Zhu, Y., ... & Chai, Z. (2019). Constructing an associative memory system using spiking neural network. *Frontiers in neuroscience*, 13, 650.
- Hidayattullah, S., Surjandari, I., & Laoh, E. (2020). Financial Statement Fraud Detection in Indonesia Listed Companies using Machine Learning based on Meta-Heuristic Optimization. *International Workshop on Big Data and Information Security (IWBIS)*, Depok, Indonesia, 79-84. <https://doi.org/10.1109/IWBIS50925.2020.9255563>

- Jiang, W., & Cui, H. (2019). The View on Curbing Financial Fraud of Listed Companies, *Proceedings of the 2019 International Conference on Economic Management and Cultural Industry (ICEMCI 2019)*, 1-21. <https://doi.org/10.2991/aebmr.k.191217.069>
- Lisbinski, F.C., & Burnquist, H.L. (2024). Institutions and financial development: Comparative analysis of developed and developing economies, *Economia*, 25(2), 347-376. <https://doi.org/10.1108/ECON-11-2023-0201>
- Lusardi, A., & Mitchell, O. S. (2014). The Economic Importance of Financial Literacy: Theory and Evidence. *Journal of Economic Literature*, 52(1), 5-44. <https://doi.org/10.1257/jel.52.1.5>
- Motie, S., & Raahemi, B. (2024). Financial fraud detection using graph neural networks: A systematic review. *Expert Systems with Applications*, 240(2), 101-125. <https://doi.org/10.1016/j.eswa.2023.122156>
- Ngai, E.W.T., Hu, Y., Wong, Y. H., Chen, Y., & Sun, X. (2011). The application of data mining techniques in financial fraud detection: A classification framework and an academic review of literature. *Decision Support Systems*, 50(3), 559-569. <https://doi.org/10.1016/j.dss.2010.08.006>
- Ran, R., Wang, L., & Fu, L. (2023). Exploring the Existence of Efficiency in the Financial Markets. *Advances in Economics Management and Political Sciences*, 13(1), 335-342. <https://doi.org/10.54254/2754-1169/13/20230741>
- Sargiacomo, M., Everett, J., Ianni, L., & D'Andreamatteo, A. (2024). Auditing for fraud and corruption: A public-interest-based definition and analysis, *The British Accounting Review*, 56(2), 1-23. <https://doi.org/10.1016/j.bar.2024.101355>
- Sender, H., Weinland, D., & Hume, N. (2020). Luckin Coffee scandal catches and world's most powerful investors. <https://www.ft.com/content/59191f9d-f2ce-4dda-8756-283aae6e9e83>
- Shoetan, Ph. O., & FAMILONI, B. T. (2024). Transforming Fintech Fraud Detection with Advanced Artificial Intelligence Algorithms, *Finance & Accounting Research Journal*, 6(4), 602-625. <https://doi.org/10.51594/farj.v6i4.1036>
- Smaili, N., Arroyo, P., & Issa, F.A. (2022). The dark side of blockholder control: evidence from financial statement fraud cases. *Journal of Financial Crime*, 29(3), 816-835. <https://doi.org/10.1108/JFC-05-2021-0113>
- Spann, D. D. (2013). *Fraud Analytics: Strategies and Methods for Detection and Prevention*, Wiley Corporate F&A, 3rd Edition.

- Teichmann, F., Boticiu, S. R., & Sergi, B. (2023). Wirecard scandal. A commentary on the biggest accounting fraud in Germany's post-war history. *Journal of Financial Crime*, 2(3), 37-56. <https://doi.org/10.1108/JFC-12-2022-0301>
- Uras, B. R. (2020). Finance and development: Rethinking the role of financial transparency. *Journal of Banking & Finance*, 111(1), 45-61. <https://doi.org/10.1016/j.jbankfin.2019.105721>
- Zhang, X., Su, J., & Zhou, L. (2022). Fraud identification of financial reports based on neural network algorithm, 40th Chinese Control Conference (CCC), Shanghai, China, pp. 8045-8049. <https://doi.org/10.23919/CCC52363.2021.9549887>
- Zhou, Y., Xiao, Zh., Gao, R., & Wang, Ch. (2024). Using data-driven methods to detect financial statement fraud in the real scenario. *International Journal of Accounting Information Systems*, 54(2), 1-31. <https://doi.org/10.1016/j.accinf.2024.100693>

References [In Persian]

- Asharoonia, K., Pouraghajan, A. A., & Naslmosavi, S. (2023). A New Approach to Predicting Fraud in Financial Statements (Compare models and simulations traditional and modern). *Accounting and Auditing Research*, 15(57), 141-160. [In Persian]
- Bahrami, A., Noravesh, I., Raad, A., & Mohammadi Molgharni, A. (2021). Financial Statements Fraud and new techniques used to detect it. *Accounting and Auditing Studies*, 10(38), 105-118. [In Persian]
- Dadgar, Y., Dargahi, H., & Gholizadeh, S. (2023). The Role of Investor Sentiment and Government Behaviour in Volatility of Tehran Stock Exchange Market: A Behavioural Economics Approach. *Quarterly Journal of Applied Theories of Economics*, 10(1), 191-214. [In Persian]
- Etemadi, H., & Abdoli, L. (2018). Audit Quality and Financial statement fraud. *Financial Accounting Knowledge*, 4(4), 23-43. [In Persian]
- Hasanpoor, D., Valiyan, H., Safari griyly, M., & Tahmasbizadeh, R. (2020). Applying Rough Developed theoretical Models (ERST), Interpretation-Structural Analysis (ISM) and Decision Tree (CART) for Help Auditors to Identify Fraud in the Financial Statements of Companies Listed on the Stock Exchange of Iran. *Journal of Investment Knowledge*, 9(33), 179-208. [In Persian]
- Jafari, M., Rezaei, N., & Selgi, M. (2023). Using a Multilayer Feedforward Artificial Neural Network in Detecting Fraudulent Financial Reporting in Companies Listed on the Tehran Stock Exchange, *Motahar Audit Studies*, 1(2), 93-113. [In Persian]

- Malekikalger, H., Bahrisales, J., Jabarzadeh Kangarlouie, S., & Ashtab, A. (2021). The effectiveness of statistical models and machine learning patterns in predicting fraudulent financial reporting. *Journal of Financial Economics*, 15(54), 267-292. [In Persian]
- Tashdidi, E., Sepasi, S., Etemadi, H., & Azar, A. (2019). New Approach to Predicting and Detecting Financial Statement Fraud, Using the Bee Colony. *Journal of Accounting Knowledge*, 10(3), 139-167. [In Persian]

استناد به این مقاله: پورساعدی، مرضیه، همت‌فر، محمود، علوی، سید عنایت‌اله، نصیرزاده، رؤیا. (۱۴۰۴).
الگوی تشخیص تقلب مالی تحت اجرای ارزیابی الگوریتم‌های مطلوبیت شبکه عصبی مصنوعی، *مطالعات تجربی حسابداری مالی*، ۲۲(۸۷)، ۸۳-۱۳۴.
DOI: 10.22054/qjma.2025.82826.2632



Empirical Studies in Financial Accounting is licensed under a Creative Commons Attribution-NonCommercial 4.0 International License.