



RESEARCH ARTICLE

A Comparative Study of XGBoost and Artificial Neural Networks for Earnings Management Prediction

Maedeh Moayeri*, Mohammad Arabmazar Yazdi, Vahid Menati

Department of Accounting, Faculty of Management and Accounting, Shahid Beheshti University, Tehran, Iran

How to cite this article:

Moayeri, M., Arabmar Yazdi, M. and Menati, V. (2026). A Comparative Study of XGBoost and Artificial Neural Networks for Earnings Management Prediction in Tehran Stock Exchange. *Iranian Journal of Accounting, Auditing and Finance*, 10(2), 69-87. doi: 10.22067/ijaaf.2026.47746.1579
https://ijaaf.um.ac.ir/article_47746.html

ARTICLE INFO

Article History

Received: 2025-08-09

Accepted: 2025-10-28

Published online: 2026-04-12

Keywords:

Artificial Neural Network (ANN), Earnings Management, Kasznik Model, Machine Learning, XGBoost

Abstract

The present study was conducted to compare the accuracy of ANN and XGBoost algorithms in predicting earnings management in listed companies of the Tehran Stock Exchange. Earnings management is one way for managers to mislead their stakeholders, which can result in financial losses; therefore, accurate detection methods are important for earnings management and beyond statistical models. The present study used 2016–2021 quarterly financial data from 103 publicly traded companies in basic metals, automotive, chemical, food and pharmaceutical producing industries (5076 year–firm). Discretionary accruals were used and calculated by the Kasznik model to capture earnings management and split them into three groups: Increasing Accruals (+1), Decreasing Accruals (-1), and Near-Zero Accruals (0). A Confusion Matrix was used to perform the model evaluation. The results showed that the XGBoost algorithm is significantly superior to the ANN with an overall accuracy of 98.4%. With very few errors, all earnings management categories XGBoost had near-perfect results. In contrast, ANN demonstrated significant weaknesses, leading to an overall accuracy of 63.1%. The study found that the model performance of XGBoost can further predict earnings management with more accuracy, thereby securing a process for financial institutions. Inspired by the relatively rare occurrence of applying the XGBoost model, used for trilateral classification of increasing, decreasing and near-zero earnings management in emerging markets, including the Tehran Stock Exchange. This research contributes significantly to the literature by demonstrating the superior predictive power of ensemble learning methods over traditional neural networks in detecting financial misconduct, which offers a robust, high-accuracy tool for regulators and investors to enhance market transparency and reduce financial risk.

<https://doi.org/10.22067/ijaaf.2026.47746.1579>



NUMBER OF REFERENCES
48

NUMBER OF FIGURES
4

NUMBER OF TABLES
5

Homepage: <https://ijaaf.um.ac.ir>
 E-Issn: 2717-4131
 P-Issn: 2588-6142

*Corresponding Author: Maedeh Moayeri
 Email: maedehmoayeri153@gmail.com
 Tel: 09167381642
 ORCID: 0009-0009-9308-9748

1. Introduction

Earnings management is essentially the manipulation of reported earnings accruals to achieve specific objectives (Saber et al., 2023). In addition, earnings management occurs when managers use judgment in financial reporting and transaction structuring to alter financial reports, either to mislead some stakeholders about the company's underlying economic performance or to influence contractual outcomes that depend on reported accounting numbers (Healy and Wahlen, 1999). Understanding earnings management is possible in two ways: firstly, as opportunistic managerial behavior, it maximizes the apparent utility of management's performance for profit-based rewards; secondly, earnings management provides management with flexibility to use it for the benefit of related parties, in terms of anticipating unexpected events, for themselves and the company (Scott, 2015).

Accordingly, managers' motivations for earnings management can be divided into two distinct and opposing categories (Ahmadpoor and Montazeri, 2011):

1. **Efficient Earnings Management:** The objective of which is to enhance the quality of information preparation to help users better understand the company's profitability and financial position.
2. **Opportunistic Earnings Management:** Where managers undertake these actions solely to maximize their own interests, rather than those of investors.

If earnings management is efficient (good), earnings quality is expected to improve; conversely, if earnings management is opportunistic (bad), it may negatively impact earnings quality. Opportunistic behavior, by manipulating earnings, reduces the reliability of information and impairs earnings quality. The transmission of confidential information about future profitability opportunities to the market through intervention in the earnings measurement process affects earnings quality and consequently adds to the company's market value (Kordestani and Tatli, 2014).

In this research, opportunistic earnings management is of particular interest, as it can cause financial losses to financial statement users by creating misleading information. Therefore, detecting earnings management can prevent a significant portion of potential losses for users.

Accounting research over the years has proposed numerous models for identifying earnings management through financial statement items. These efforts began with seminal models such as the Healy Model (1985) and the DeAngelo Model (1996). The approach reached a turning point with the introduction of the Jones Model (1991) and subsequently the Modified Jones Model by Dechow et al. (1995), which attempted to control for non-discretionary factors. However, more sophisticated models, including the Kasznik Model (1999), were developed due to challenges arising from the insufficient control of firm operating performance. The Kasznik model enhanced the power to separate discretionary and non-discretionary components by adding the operating cash flow variable (Sun and Rath, 2010; Siekelova, 2021).

Accounting is the science of applying data, and in some cases, analyzing a large volume of data and making decisions based on it, which can be difficult and time-consuming. Given the increasing complexity of processes, the vast amount of information, and the intricate analysis required. Consequently, accounting researchers who need to automate these processes, like those in other fields, should explore the possibility of utilizing artificial intelligence tools. Machine learning is a subfield of artificial intelligence that focuses on developing algorithms capable of learning from data and making predictions or performing actions without explicit programming. The advantage of using these algorithms in accounting data analysis is saving time and analysis costs, and developing methods to increase the usefulness based on existing information and records. With the development of financial markets and the ever-increasing volume of information, participants in these markets are seeking simple and cost-effective tools that can enable them to make accurate and rapid predictions about the future market situation.

Although statistical models possess the ability to make acceptable predictions in analyzing financial and accounting issues, the limiting assumptions of some of these models (ambiguity in the relationships between input and output variables, collinearity error, lack of observations, etc.) reduce their effectiveness (Esfahanipour et al., 2016).

Moving beyond the established superiority of artificial intelligence methods over conventional statistical methods, we encounter a wide variety of machine learning algorithms. These algorithms themselves can exhibit varying levels of performance relative to each other, depending on their specifically defined objective. Therefore, researchers in this field are obligated to investigate two aspects: first, the feasibility of applying these algorithms to achieve the desired goal, and then to compare their accuracy and efficiency with each other to achieve the main objective of their research, which is to meet the needs of financial market participants. In this research, two powerful and novel machine learning algorithms were utilized (XGBoost and ANN) to predict earnings management in five industries: basic metals, automotive and parts manufacturing, chemical products, food products, and pharmaceutical materials and products, within the Tehran Stock Exchange market. We will then examine and compare their performance in terms of accuracy.

These two models, which are similar yet different in fashion in tackling complex financial data and the domestic literature, suffer from a critical research gap that makes the two of them the natural choices. The main advantage of ANNs is their inherent flexibility to model very non-linear relationships, which is a key requirement, considering that most financial data do not comply with simple linear structures (Artene and Domil, 2025). In contrast, XGBoost provides superior predictive accuracy, very fast run time, automatic handling of sparsity and robustness to different types of noise in the data compared with general linear models and many other conventional ensemble methods (Wiens et al., 2025). Importantly, most likely today, there is no evidence of any such literature that directly examines the comparative performance between XGBoost (as the leading gradient boosting algorithm) and Artificial Neural Network (ANN (a representative of non-linear/deep learning models) on predicting earnings management, in Iranian domestic literature.

Thus, the main purpose of this study is to set a new standard by directly comparing these two powerful algorithms under an adverse trilateral classification scenario (Increasing, Decreasing, and Near-Zero Accruals), with a precise Kasznik model. This comparison addresses a significant local research gap and provides a definitive assessment of the optimal predictive model for earnings management within the Tehran Stock Exchange.

The rest of this paper is organized as follows: We begin with Section 2, which establishes the theoretical foundations, reviews the relevant literature, and poses the research question. Next, Section 3 outlines the research methodology, including data collection and the implementation of machine learning algorithms. Subsequently, Section 4 reports the empirical findings and performance analysis, followed by the concluding remarks in Section 5.

2. Theoretical principles, literature review, and research question

2.1. Theoretical grounding of the prediction model structure

The structure of this prediction model is fundamentally based on the literature derived from Positive Accounting Theory (PAT) and Agency Theory. These frameworks provide the theoretical justification for selecting the financial features used as inputs (predictors) for the machine learning algorithms. The central theoretical premise is that managerial incentives (e.g., maximizing bonuses, avoiding debt covenant violations, mitigating political scrutiny) manifest as changes in financial performance indicators. The study uses the core components of the Kasznik Model, a theoretically justified accruals model, as the input features (X).

2.2. Earnings management

Like any other research field or domain, there is no consensus or agreement on a single, unified theory to explain and predict earnings management. The most debated theories of earnings management include: Prospect Theory, Catering Theory, Agency Theory, Signaling Theory, Big Bath Theory, Behavioral Theory, Game Theory, Shareholders Theory, and Entity Theory.

According to Saghafi and Pouryanasab (2010), all earnings management theories share three commonalities:

1. They all endeavor to explain the relationship between company managers and the group(s) of users or stakeholders concerning earnings management.
2. For this reason, and due to the use of two logical premises, "rational human" and "self-interest," and their product, "utility maximization," all these theories seek to explain the relationship between earnings management and value.
3. All these theories have been developed within the realm of positive accounting's dominance and therefore possess a positive and inductive nature, derived from observation and experience.

Ronen and Yaari (2002) accept "Firm Theory" as the dominant theory and place it above other theories. In this article, Firm Theory is also positioned as the dominant theory for three reasons: 1) It has a longer history than others, at least in other related research fields; 2) Other earnings management theories are, in some way, indebted to it, or in other words, this theory serves as a paradigm for the set of its rival theories; 3) And each of the other theories is, in some way, an approach derived from Firm Theory. As mentioned before, all earnings management theories strive to explain the relationship between insiders and outsiders regarding the phenomenon of earnings management. Therefore, a paradigmatic relationship exists between Firm Theory and its rivals, and thus, rival theories are each positioned within the framework of Firm Theory's approaches.

2.3. Models for measuring earnings management

Various models for detecting and measuring earnings management were proposed by different researchers, which are generally classified into two categories: discretionary accruals models, which are based on the Jones model, and discretionary revenue models, which have recently gained attention from financial researchers. Discretionary accruals, used as an indicator for measuring earnings management, are items over which management has control and can manipulate (accelerate or delay recognition, or omit them) (Salehi and Salehi, 2015).

Rahnamay Roodposhti et al. (2014) evaluated five accruals models and the Keeler discretionary revenue model for predicting earnings management in companies listed on the Tehran Stock Exchange. Their research findings indicated that earnings management models based on discretionary accruals have greater predictive power compared to the discretionary revenue model. Furthermore, the earnings management index derived from accrual-based models is a relevant and influential variable on the company's stock market value, in comparison to the revenue model.

To date, extensive research has been conducted in the field of comparing proposed models for detecting earnings management. Some of these studies in Iran have applied the proposed models in various industries and compared their results; for example, Montazeri and Khodamoradi (2017) investigated the detection power of earnings management models in companies listed on the Tehran Stock Exchange to identify the most suitable model for evaluating earnings management. In this research, the Jones, Modified Jones, and Kasznik models were compared. The findings indicated that among the three models, the Kasznik model yields the highest correlation. Also, according to the results of Rahmani and Bashirimanesh's research (2013), the adjusted Jones model was not significant

in some industries and its adjusted coefficient of determination was also at a low level.

Kasznik model

The Kasznik model is as follows:

$$\text{(Model 1)} \\ \frac{TA_{it}}{A_{it-1}} = \alpha_0 \left(\frac{1}{A_{it-1}} \right) + \beta_1 \left(\frac{\Delta REV_{it} - \Delta REC_{it}}{A_{it-1}} \right) + \beta_2 \left(\frac{PPE_{it}}{A_{it-1}} \right) + \beta_3 \left(\frac{\Delta CFO_{it}}{A_{it-1}} \right) + \varepsilon_{it}$$

Where:

$TA_{i,t}$ = Accruals for company i in year t

$A_{i,t-1}$ = Total book value of assets for company i at the beginning of the period

$\Delta REV_{i,t}$ = Change in operating revenues for company i between year t and $t-1$

$\Delta REC_{i,t}$ = Change in accounts receivable for company i between year t and $t-1$

$PPE_{i,t}$ = Gross property, plant, and equipment for company i at the end of year t

ΔCFO_{it} = Change in operating cash flows between year t and $t-1$.

Kasznik's model adjusts the Jones model in three aspects. The first and most important adjustment is the inclusion of the change in cash from operations as the third independent variable (Salehi and Salehi, 2015). Recent empirical evidence and research show that accruals are negatively correlated with changes in cash flows, which is likely due to the nature of the accounting pattern (Dechow et al., 1995).

The second change from the Jones model is loosening the assumption of exogenous revenues. The timing of revenue recognition is a method often used by managers to manage reported earnings. If, in the model design, revenue is considered an exogenous factor beyond management's discretion, the presented model will not be able to detect earnings management through the timing of revenue recognition, and consequently, the model's power to detect earnings management will be weak (Salehi and Salehi, 2015). Following Dechow et al. (1995), this problem can be mitigated by adjusting the revenue variable by the change in accounts receivable.

The third adjustment involves the structure of the estimated portfolios. The accruals model uses a cross-sectional method instead of a time-series method to estimate the coefficients α_0 , β_1 , β_2 , and β_3 (Salehi and Salehi, 2015). The advantage of the cross-sectional method is that it can control the effects of changes in economic conditions across a wide industry scope on total accruals and, due to the possibility of structural changes, allows for these coefficients to vary over the period (Kasznik, 1999).

2.4. Supervised machine learning

In machine learning methods, algorithms are developed that interpret the pattern of input data, extract the governing rules for that data set, and ultimately use these rules to predict other data.

Supervised learning is a type of machine learning algorithm where the model uses labeled training data and a set of training examples to infer a function (Sarker, 2020). Supervised learning occurs when we have identified and predetermined specific objectives for a particular set of data (Sarker et al., 2020). In essence, learning in this category of algorithms is used to identify patterns or behaviors in a "specified" training dataset (Boutaba et al., 2018).

2.5. Research background

First, relevant domestic research was reviewed to identify existing gaps in previous studies:

Faghani et al. (2016) aimed to predict earnings management using the Artificial Neural Network (ANN) and a hybrid Genetic Algorithm-ANN model. The findings showed that ANN has a high ability to predict earnings management compared to the linear Modified Jones model. However, this study's reliance on the Modified Jones model, which has been shown to have low detection power in certain industries, limits its generalizability (Montazeri and Khodamoradi, 2017).

Salehi and Farokhi Pile Rood (2018) compared the predictive accuracy of earnings management using Neural Networks and Decision Trees against linear models. The results indicated that neural networks and decision trees offer greater accuracy and lower error rates in predicting earnings management compared to linear methods. Given the established superiority of AI tools over traditional linear models, the current research moves beyond this conclusion to find the single best-performing algorithm.

Fereydouni et al. (2020) predicted earnings smoothing using the Relevance Vector Machine (RVM) algorithm. The results showed that the RVM algorithm, in both its linear and non-linear forms, has a high capability in predicting the extent of earnings smoothing, with the non-linear approach performing better. The current research contributes by comparing and examining newer and more powerful machine learning algorithms to achieve a more comprehensive conclusion regarding the optimal model for earnings management prediction.

Maleki Kakler et al. (2021) applied machine learning models (Decision Tree, Support Vector Machine, and Bayesian Network) to develop an effective pattern for predicting fraudulent financial reporting. The results showed that the Decision Tree algorithm was successful, achieving an accuracy of over 92.61%. Similar to others, this study focused on demonstrating the superiority of machine learning over statistical models, whereas our research centers on a head-to-head comparison of two distinct and powerful machine learning algorithms.

Hassani et al. (2024) conducted a study to distinguish between accrual earnings management and real earnings management using machine learning methods (decision trees, SVM, KNN, deep learning) and combining them with relief feature selection and principal component analysis (PCA) methods. The findings showed that in predicting accrual earnings management, the relief-based feature selection model has a better ability than the PCA model, but this superiority was not observed in predicting real earnings management. Also, the results indicated that accrual earnings management can be predicted with higher accuracy than real earnings management. In this study, a binary classification method was used, while the present study uses a triple classification method.

Hamidian and Esmaeili (2024) have a study focused on predicting earnings management using support vector machine and decision tree methods in 170 companies. The results showed that support vector machine and decision tree methods have the ability to predict earnings management of companies and the support vector machine method was introduced as the superior method in predicting earnings management. In this study, the focus is more on the independent variables (attributes) than on the specific model of measurement of the dependent variable (discretionary accruals).

After mentioning the gaps in domestic research in the field of applying machine learning in earnings management prediction, studies that are more closely related to the research objective and methodology can be reviewed, focusing on finding new research opportunities:

Rahman et al. (2021) compared two types of machine learning algorithms, linear classification and tree-based classification (Decision Tree and Random Forest), for predicting earnings manipulation. The results indicated that the performance of the two algorithm types was not significantly different,

but tree-based classifiers achieved the best accuracy in the shortest time.

Hammami and Hendijani Zadeh (2022) employed six new data analysis methods in classification to predict earnings management (both accrual and real). The results showed that the ABC-SVM model (Ant Colony Optimization and Support Vector Machine) performed better than other models in predicting earnings management.

Chen et al. (2022) combined Random Forest with Gradient Boosting in a hybrid machine learning model to predict one-year-ahead earnings changes in direction. The accuracy of the proposed model was found to be between 67.52% and 68.66%, which outperforms the linear models.

Adomat et al. (2023) investigated a new cloud production environment-based earnings management technique using machine-learning algorithms. They observed that their combined Neural Network and Linear Regression model achieved the lowest error rates compared to standard and automated linear regression algorithms.

Ali et al. (2023) produced a state-of-the-art FSF detection model dedicated to the MENA region with publicly available financial data. Based on the empirical results, XGBoost performed substantially higher than all other machine learning models - Logistic Regression, Decision Tree, Support Vector Machine, AdaBoost and Random Forest. Optimization resulted in a final accuracy of 96.05% for FSF detection by the XGBoost algorithm.

Chakri et al. (2023) used four different supervised machine learning models (Linear Regression, K-Nearest Neighbors, Support Vector Regression, and Decision Tree) for earnings prediction. Decision Tree proved to be the best model for performance analysis according to the findings.

In Chen et al. (2025), the multi-class Financial Distress Prediction (FDP) model was developed, which offered a deeper insight into corporate financial distress classification. The study introduced a new definition of the financial status in three classes and applied the Deep Forest algorithm for classification. This is the first multi-class forecast in the financial domain, and thus is a remarkable research.

Huy et al. (2025) addressed the limitations of traditional discretionary accrual models in detecting the nonlinear patterns of earnings management. They evaluated tree-based machine learning models (Decision Trees, Random Forests, and Gradient Boosting Machines). Their findings established that the Gradient Boosting Machine (GBM) significantly outperformed the other models across all key metrics (accuracy, precision, recall, and F1 score), proving to be the most effective tool.

The aim of Vezanones and Severin (2025) was to identify the financial profiles of companies using machine learning and to provide a visual representation of the financial profiles associated with earnings management strategies (upward and downward) and tools (accruals and real activities). The proposed machine learning method performed better than traditional forecasting methods.

Domestic literature on predicting earnings management using machine learning algorithms in the Tehran Stock Exchange is constrained by a reliance on older models (proving only the general superiority of Machine Learning over linear methods), the use of less accurate earnings management models like the Modified Jones model (Montazeri and Khodamoradi, 2017; Rahmani and Bashirimanesh, 2013) and a simplistic binary classification.

This study addresses these limitations by establishing a new benchmark for earnings management prediction on the Tehran Stock Exchange. Our innovation lies in three key areas: 1) Employing the Kasznik model, which is locally validated as more robust; 2) Adopting a challenging trilateral classification (increasing, decreasing, and near-zero earnings management); and 3) Conducting a systematic, head-to-head comparison between the XGBoost and Artificial Neural Network (ANN)

algorithms to identify the optimal model in this field.

2.6. Research question

Which of the Artificial Neural Network (ANN) algorithms and XGBoost has higher accuracy in predicting earnings management?

3. Research methodology

The present research, in terms of its objective, is applied; and in terms of its nature, given that it examines and compares the accuracy of algorithms and uses real data for testing, it falls under descriptive research. Regarding the research method, an ex-post facto approach was used. In ex-post facto research, the investigation begins after an event has occurred. Here, the researcher, essentially not present, does not manipulate the variables and does not know them, but rather conducts causal research to identify these variables and factors that caused the event (Sarmad et al., 2004). In this research, by implementing machine learning algorithms, we aim to improve earnings management prediction and compare the results of the algorithms with each other to ultimately select the best one (in terms of accuracy) for better prediction.

3.1. Data collection method

Many accounting and financial studies on the Tehran Stock Exchange are conducted using the Rahavard Novin database. Rahavard Novin is the most comprehensive software for the Iranian capital market. Since its initial release in 1999, due to its extensive database and advanced tools, it has always been the primary choice for both individual and institutional participants in the capital market (Mosallanejad and Nazemi, 2019). Rahavard Novin software provides advanced tools that allow access to extensive information on companies listed on the Tehran Stock Exchange in an optimal time. In this research, a complete set of required information was extracted from this software, and when necessary, information was verified by direct reference to the Codal website.

3.2. Statistical population and sample

The statistical population of this research consists of classified and audited financial data of active companies listed on the Tehran Stock Exchange during the period from 2016 to 2021.

The statistical sample (present in Table 1) is determined by the systematic exclusion method and includes companies that meet the following conditions:

1. Must be a listed company on the stock exchange.
2. The company's fiscal year must end on March 20.
3. The company's fiscal year must not have changed within the research period (2016 to 2021).
4. Must have been continuously active on the stock exchange from 2016 to 2021.
5. Required financial information, including financial statements and accompanying notes, for calculating research variables must be available for the relevant companies within the 2016-2021 timeframe.
6. Must belong to one of the five industries: Basic Metals, Automobile and Parts Manufacturing, Chemical Products, Food Products, and Pharmaceutical Materials and Products.

The table related to the application of each of the above-mentioned restrictions is as follows:

Table 1. Sample selection process by applying population constraints and conditions

Constraint	Percentage of Deletion	Remained Sample Size
Initial population		5076 year-firm

1. Companies must be listed on the main stock exchange (Tehran Stock Exchange) and not on the over-the-counter market.	$\cong 62\%$	1892 year-firm
2. The companies' fiscal year must end on March 20.	$\cong 45\%$	1032 year-firm
3. The companies' fiscal year must not have changed within the research period.	$\cong 14\%$	880 year-firm
4. Must have been continuously active on the stock exchange from 2016 to 2021.	$\cong 5\%$	836 year-firm
5. Required financial information, including financial statements and accompanying notes, for calculating research variables must be available for the relevant companies within the 2016-2021 timeframe.	$\cong 2\%$	816 year-firm
6. Must belong to one of the five industries: Basic Metals, Automobile and Parts Manufacturing, Chemical Products, Food Products, and Pharmaceutical Materials and Products.	$\cong 49\%$	412 year-firm (Final Sample)

Note. The total number of companies at the beginning of the selection process was 5076 year-firm. The percentages indicate the estimated proportion of deletion at each step relative to the previous sample size.

3.3. Obtaining discretionary accruals

In this stage of the research, the variables of interest (according to the Kasznik model), including the book value of company assets, operating income, net accounts receivable, and operating cash flow, were extracted from the available data in the sample's financial statements and saved in an Excel environment. Gross property, plant, and equipment (PPE) is among the information found in the notes to the financial statements, which was collected by direct reference to the financial statements via the Codal website.

In the Kasznik model, β and α are company-specific parameters estimated using OLS regression.

Given the total accruals (TA), which comprise normal accruals (NA) and abnormal accruals (AA), the sum of the four terms on the right side of the Kasznik model (Model 1) represents normal accruals, and $\varepsilon_{i,t}$ (the residual) indicates the abnormal (discretionary) accruals for company i in fiscal year t . Discretionary accruals are recognized as a measure of earnings management, and their level indicates the occurrence or non-occurrence of earnings management (Dechow et al., 1995).

According to Guara et al., total accruals (TA) are calculated using the following formula:

(Formula 1)

$$\text{Total Accruals} = (\text{Net Income} - \text{Operating Cash Flow}) \times (-1)$$

Accordingly, the level of discretionary (abnormal) accruals is measured after estimating the regression coefficients (β and α) for each company separately using OLS regression and measuring $\varepsilon_{i,t}$.

3.4. The role of the Kasznik model in the machine learning framework & data labeling

The Kasznik Model (1999) plays a foundational and dual role in establishing the Supervised Machine Learning framework of this study, which serves as both the measurement tool for the target variable (Y) and the basis for defining the input features (X).

3.4.1. Defining the dependent variable (Y) and labeling

Measurement Tool (Target Definition): The discretionary accruals ($\varepsilon_{i,t}$) obtained in Section 3-3, which serve as the proxy for earnings management, are the core input for data labeling. Data labeling is performed based on the value obtained from the level of discretionary accruals. Based on predefined thresholds, this value is converted into the discrete categorical variable Y (the dependent variable for prediction), classifying the outcome into Increasing Accruals (+1), Decreasing Accruals (-1), or Near-Zero Accruals (0).

3.4.2. Defining the input features (X) and timeliness control

Feature Definition: The model also fundamentally dictates the selection of our input features. The input variables, directly extracted from the sample companies' financial statements (Balance Sheet, Income Statement, and Cash Flow Statement) at the end of year $t-1$, correspond to the components of the Kasznik Model (1999), which are the standard determinants used to estimate Non-Discretionary Accruals.

The machine learning models (XGBoost and ANN) were trained to predict earnings management in year t . Accordingly, all input variables (Features) were exclusively extracted from financial information available at the end of the fiscal year $t-1$ (lagged variables). This rigorous time control ensures that the models use no information that would not have been published and available at the time of the actual prediction (i.e., at the beginning of year t), thereby eliminating the risk of Look-ahead Bias and Data Leakage.

3.5. Examining the general distribution of discretionary accruals

Figure 1 is a histogram illustrating the distribution of data related to discretionary accruals (earnings management) within the examined dataset. A histogram is a type of bar chart for visualizing continuous data, which shows the distribution of values within a dataset (Archambault et al., 2015). In this graph, a probability distribution curve has also been used to more clearly illustrate the data distribution. This graph indicates that the highest data density (1.42) is at a point close to zero (-0.04). According to this graph, the data distribution is asymmetric, non-normal, and has a long tail to the right, which suggests that high values of earnings management are more frequent than low values.

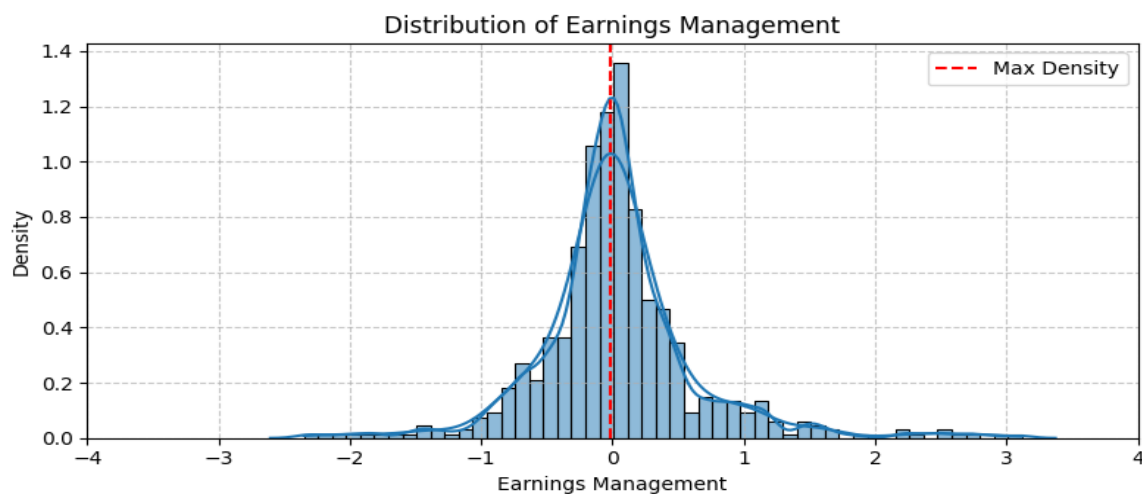


Figure 1. Distribution of data related to discretionary accruals (earnings management)

3.5.1. Determining the first and third quartiles and the median of the data

In this distribution, the median, first quartile, and third quartile values are as follows (Figure 2):

- *Median:* 0.0018487886711916826
- *First Quartile (Q1):* -0.22011942451641645
- *Third Quartile (Q3):* 0.23925656442091495

Based on the obtained values:

- Discretionary accruals whose value is greater than or equal to the third quartile are classified as items with increasing earnings management and labeled as +1.

- Discretionary accruals whose value is less than or equal to the first quartile are classified as items with decreasing earnings management and labeled as -1.
- Discretionary accruals whose value is between the first and third quartiles are classified as items with near-zero earnings management and labeled as 0.

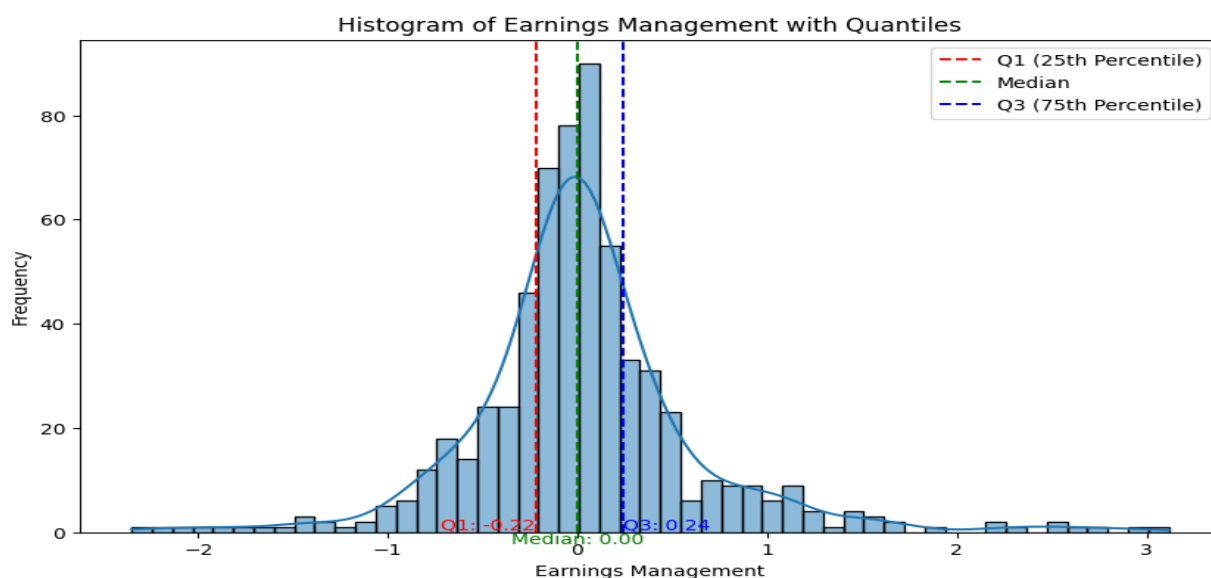


Figure 2. First quartile, third quartile, and median of the discretionary accruals data distribution

3.5.2. Splitting training data into training and test sets

An algorithmic model is built based on training data, and then the model's performance is evaluated using test data (Jain et al., 2022). In fact, at this stage of the research, we focus on developing algorithmic models and adapting them to the data structure, which is considered the main task in this research.

Overfitting is a fundamental problem in supervised machine learning that prevents fully generalizing the model from observed training data to unobserved test data and can arise due to noise, the small size of the training set, and the complexity of the classifier. Despite overfitting (the model works perfectly on the training set and suboptimally on test data), it has some issues dealing with information that is more present in the test set than what was seen during training (Ying, 2019). In order to minimize the probability of overfitting, we take some part of the training dataset as a validation set.

This process is called the hold-out method. We take a label dataset (training) and split it into a training and test set. Next, we train a model on the training set and make predictions for the test set labels. We compute the fraction of correct predictions (via comparison to the pre-predicted labels), which is our estimate of predictive accuracy for the model (Burzykowski et al., 2023).

Splitting the training dataset into training and test sets is considered essential for eliminating or reducing bias in training data within machine learning models. This process is always performed by data scientists and data analysts to prevent algorithm overfitting (Muraina, 2022).

The hold-out method is one of the most common cross-validation methods. The set of cross-validation techniques is a method used to evaluate the predictive performance of a model (Yadav and Shukla, 2016).

This operation is performed on the training set with the following code example:

Split the data into training and testing sets

```
X_train, X_test, y_train, y_test = train_test_split(X, y, test_size=0.2, random_state=42)
```

The percentage for splitting training and test data is not considered significant in any source and can be chosen by the researcher based on the total volume of training data. In this research, 80% of the training dataset is used for model fitting, and 20% of this set is considered for evaluating the model's accuracy (test data). This splitting process was performed using the stratified hold-out method to ensure that the proportional representation of the three earnings management classes is maintained across both the training and test sets, thereby mitigating bias arising from class imbalance.

3.6. Implementation of machine learning algorithms

In this section, the machine learning models of interest for this research, including K-Nearest Neighbors, Support Vector Machine, Decision Tree, and Random Forest, are each fitted separately and their performance evaluated.

Model fitting is performed by estimating the most suitable hyperparameters. Hyperparameter estimation is carried out using GridSearch. This method simply performs an exhaustive search over a specified subset of the training algorithm's hyperparameter space.

Since the hyperparameter space of a machine learning algorithm may include a space with infinite values for some parameters, we must specify a boundary for each hyperparameter to apply Grid Search (Liashchynskiy and Liashchynskiy, 2019). To mitigate the instability of a single Hold-out split and ensure robustness and generalization capability, GridSearch inherently employs k-fold cross-validation (CV) (specifically, we utilized 5-fold cross-validation here) on the training set to select the optimal parameters

XGBOOST Algorithm XGBoost is a machine learning algorithm that utilizes ensemble tree models. Inspired by tree boosting algorithms (composed of decision trees), it uses a combination of multiple Regression Trees as its core model. In this method, each leaf of the tree contains a continuous score. Decision rules within the trees are used to classify these leaves. The final class of each leaf is predicted by aggregating the scores of the corresponding leaves, striving to provide the best prediction for machine learning problems (Rashidpour and Zare Bidaki, 2024).

The XGBoost algorithm has numerous hyperparameters. To achieve the best predictive performance and ensure robustness, the optimal values for the most critical hyperparameters were determined using Grid Search combined with 5-fold Cross-Validation. The three key hyperparameters selected for tuning were the `n_estimators` (number of decision trees), `max_depth` (maximum depth of decision trees), and `learning_rate` (shrinkage rate). Their optimal values are summarized in Table 2.

Table 2. XGBoost best Hyperparameters

Hyperparameter	Value
<code>N_{estimator}</code>	50.000
<code>Learning rate</code>	0.010
<code>Max_{depth}</code>	3.000

ANN Algorithm, an Artificial Neural Network (ANN), is a type of mathematical model that mimics the functioning of the human brain. Their ability lies in extracting patterns from observed data without requiring assumptions about the relationships between variables. They are comprehensive, flexible functions and powerful tools for data analysis and modeling nonlinear relationships with a high degree of accuracy (Sabeti et al., 2023).

ANN encompass various types of models, and in this research, we utilize the Deep Neural Network model. A Deep Neural Network is a multi-layered neural network that possesses several hidden

layers. This algorithm also includes hyperparameters, whose optimal values were obtained using GridSearch and are presented in Table 3.

Table 3. ANN best Hyperparameters

Hyperparameter	Value
Batch _{size}	10
epochs	50
Model _{neurons}	32
Model _{optimizer}	adam

3.7. Programming environment and software implementation

The implementation, training, and testing of the machine learning algorithms (XGBoost and ANN) were performed using the Python 3.9 programming language within the Visual Studio Code (VS Code) environment. The following specialized libraries were utilized for data processing and model execution:

- *Core Libraries:* Scikit-learn (for model implementation, cross-validation, and data splitting), XGBoost (for the Extreme Gradient Boosting model), TensorFlow/Keras (for building and training the Artificial Neural Network), Pandas and NumPy (for data manipulation and algebraic operations).
- *Statistical Libraries:* Statsmodels (utilized for the OLS regression necessary to estimate the Non-Discretionary Accruals in the Kasznik model).

This detailed documentation ensures the reproducibility of our research findings.

4. Research findings

To evaluate the performance of the models, we utilize the Confusion Matrix. A Confusion Matrix is a table that visualizes the performance of a supervised learning algorithm. Each column of the matrix represents the predicted classes, while each row represents the actual classes. All correct predictions are located along the diagonal of the matrix, allowing errors to be easily observed by non-zero values outside the diagonal (Patro and Patra, 2014).

Precision, Recall, F1-score, and Support are metrics extracted from the Confusion Matrix to assess algorithm performance:

- *Precision:* Among all instances that the model predicted to belong to a specific class, what percentage truly belonged to that class? High precision means a low False Positive rate.
- *Recall:* Among all instances that actually belonged to a specific class, what percentage did the model correctly identify? High recall means a low false-negative rate.
- *F1-score:* Represents the harmonic mean of precision and recall. A high F1-score indicates good performance in both precision and recall.
- *Support:* Denotes the actual number of occurrences for each class in the specified dataset.
- *Weighted Average:* reports the model's performance across all classes (Categories) simultaneously, by taking into account the size (Support) of each individual class.

For the XGBoost algorithm, the results, according to the Confusion Matrix, are as follows (Figure 3 and Table 4):

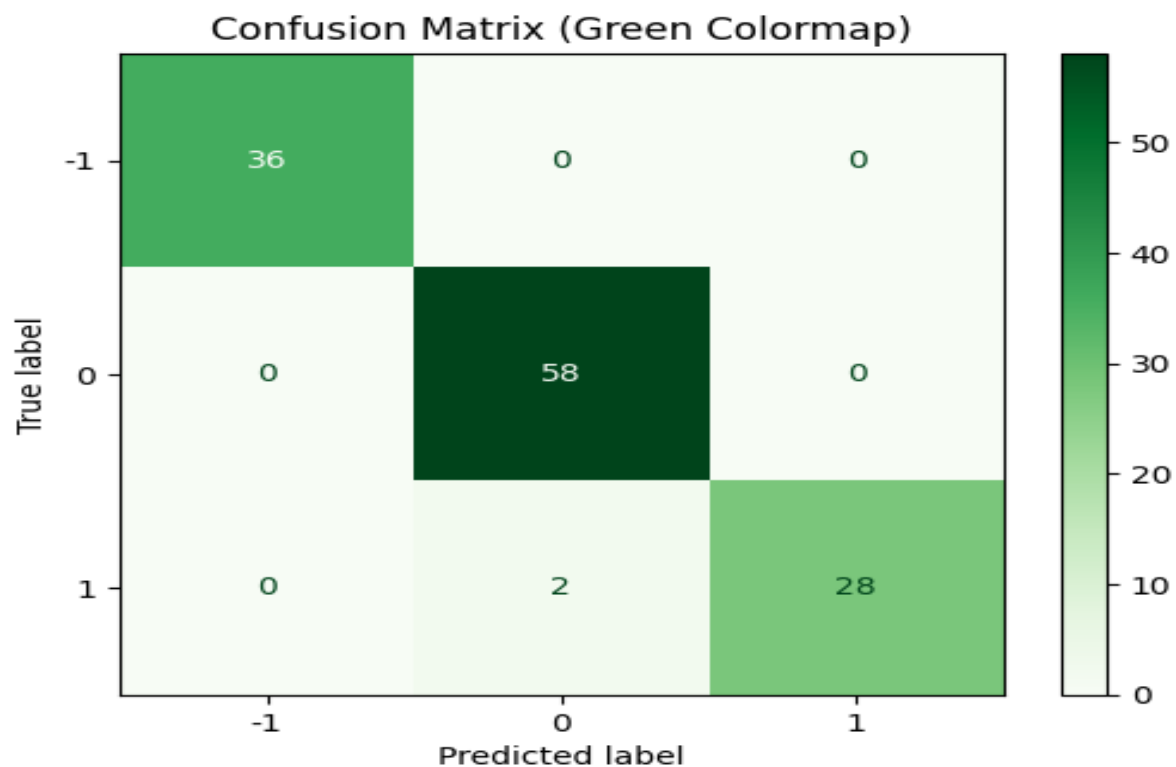


Figure 3. Confusion matrix of the XGBoost algorithm

Table 4. Performance metrics for the XGBoost algorithm (Classification Report)

Class	Precision	Recall	F1-Score	support
Decreasing Accruals (-1)	1.000	1.000	1.000	36
Near-Zero Accruals (0)	0.970	1.000	0.980	58
Increasing Accruals (+1)	1.000	0.930	0.970	30
Weighted Average	0.980	0.980	0.980	124

Note. The reported metrics are based on the testing set (N=124). The overall model accuracy was 0.98.

4.1. Interpretation of XGBoost algorithm results

This model demonstrated perfect performance for Class -1, which correctly identified all 36 instances of Class -1; every time it predicted an instance as Class -1, it was correct.

Furthermore, in identifying Class 0, it correctly found all 58 instances (recall: 1.00). When it predicted an instance as Class 0, it was correct 97% of the time, meaning there were very few false positives for this class (precision: 0.97).

Finally, every time it predicted Class 1, it was correct (precision: 1.00). However, its Recall is slightly lower at 93%, meaning it did not identify 7% of the actual Class 1 instances (i.e., it had false negatives for Class 1).

Similarly, for the ANN algorithm, the results according to the Confusion Matrix are as follows (Figure 4 and Table 5):

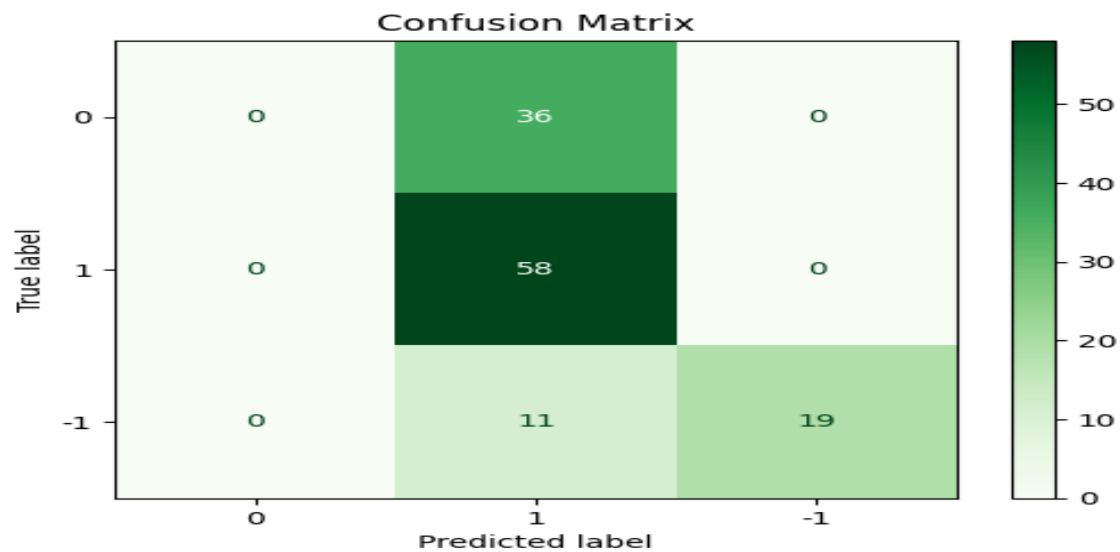


Figure 4. Confusion matrix of the ANN Algorithm

Table 5. Performance metrics for the ANN algorithm (Classification Report)

Class	Precision	Recall	F-Score	support
Decreasing Accruals (-1)	0.000	0.010	0.020	36
Near-Zero Accruals (0)	0.560	1.000	0.720	58
Increasing Accruals (+1)	1.000	0.670	0.800	30
Weighted Average	0.630	0.630	0.630	124

Note. The reported metrics are based on the testing set (N=124). The overall model accuracy was 0.63.

4.2. Interpretation of ANN algorithm results

Recall (Sensitivity): All actual Class 0 instances were correctly classified. *Precision:* Only 55% of the instances predicted as Class 0 by the model were actually Class 0.

All predicted Class 1 instances were actually Class 1. This means that the model has a very high precision in predicting for this class (Precision: 1.00). The recall for Class 1 was at only 63%. That is to say, the model classified 37% of true Class 1 cases as another class (Recall: 0.63).

For Class -1, since there were no instances that belonged to it (highlighted in red, middle column — by default), the Precision and Recall are 0.

5. Conclusion

The primary purpose of this research was to compare the performance of the Artificial Neural Network (ANN) and XGBoost algorithms in predicting earnings management among the companies of five specified industries) listed on the Tehran Stock Exchange. The novelty of this study lies in its use of a trilateral classification of earnings management (Increasing Accruals: +1, Decreasing Accruals: -1, and Near-Zero Accruals: 0) and the application of the Kasznik Model for predicting.

5.1. Methodology and key findings

This study utilized financial data from 103 companies (5076 year-firm) across five industries (basic metals, automotive, chemical, food, and pharmaceutical) over the period 2016 to 2021. The research adopted an applied, descriptive, ex-post facto approach. Discretionary accruals, derived from

the Kasznik model, served as the proxy for earnings management, and the data were strictly labeled into the three classes based on quartiles. To gain a valid prediction of current-year (t) earnings management and to thwart implicit data leakage, a crucial methodological control was instituted by limiting all input variables to information available at the end of the prior fiscal year (t-1). A confusion matrix was used to evaluate the performance of the models.

The findings above clearly demonstrated that the XGBoost algorithm outperforms:

- XGBoost Algorithm: 98.4% overall accuracy demonstrated perfect or near-perfect performance across all three earnings management categories, successfully differentiating between them with very few errors.
- ANN Algorithm: Showed considerable weaknesses, resulting in a significantly lower overall accuracy of 63.1%. The ANN model was completely unsuccessful in identifying the Near-Zero Accruals (Class 0), demonstrating a failure to generalize across all classes.

The study concludes that XGBoost is a superior and more reliable model for predicting earnings management in this context.

5.2. Comparison with past research

This research contributes significantly to the literature by moving beyond the established superiority of machine learning methods over conventional statistical models. Previous studies on the Tehran Stock Exchange have often been limited by using older earnings management models like the Modified Jones model (which has low detection power in some industries) and simpler binary classification. Leveraging the advanced Kasznik model and a complex 3-way classification, this study lays the groundwork for future research and demonstrates that set-based methods, such as XGBoost, surpass traditional neural networks (ANN) in addressing sophisticated financial misbehavior.

5.3. Practical implications and future research suggestions

Practically, the high-accuracy XGBoost model offers a robust tool for regulators and investors to enhance market transparency and reduce financial risk by identifying companies engaging in opportunistic earnings management.

For future research, it is suggested:

1. To extend the application of the XGBoost model to predict other forms of financial misconduct, such as fraudulent financial reporting, and compare its efficacy with other advanced ensemble or deep learning algorithms.
2. To replicate this trilateral classification methodology in other industries to test the generalizability of XGBoost's superior performance.

5.4. Research limitations

The conclusions of this study are limited to the scope of its design; Specifically:

- The findings are restricted to the five specified industries and companies listed on the Tehran Stock Exchange.
- The superior effectiveness of the XGBoost model is only established against the Artificial Neural Network within the context of trilateral classification.
- It is important to note that this study's conclusions focus solely on predictive superiority and do not extend to causal analysis of earnings management.

References

1. Adomat, V., Ehrhardt, J., Kober, C., Ahanpanjeh, M. and Wulfsberg, J. P. (2023). A machine learning approach for revenue management in cloud manufacturing. *Procedia CIRP*, 118, pp. 342-347. <https://doi.org/10.1016/j.procir.2023.06.059>
2. Ahmadpoor, A. and Montazeri, H. (2011). Type of earnings management and the impact of company size, ownership structure and corporate governance on earnings management. *Journal of Accounting Advances*, 3(2), pp. 1-35. <https://doi.org/10.22099/jaa.2012>
3. Ali, A. A., Khedr, A. M., El-Bannany, M. and Kanakkayil, S. (2023). A powerful predicting model for financial statement fraud based on optimized XGBoost ensemble learning technique. *Applied Sciences*, 13(4), A. 2272. <https://doi.org/10.3390/app13042272>
4. Artene, A. E. and Domil, A. E. (2025). Neural networks in accounting: bridging financial forecasting and decision support systems. *Electronics*, 14(5), p. 993. <https://doi.org/10.3390/electronics14050993>
5. Boutaba, R., Salahuddin, M. A., Limam, N., Ayoubi, S., Shahriar, N., Estrada-Solano, F. and Caicedo, O. M. (2018). A comprehensive survey on machine learning for networking: evolution, applications and research opportunities. *Journal of Internet Services and Applications*, 9(1), pp. 1-99. <https://doi.org/10.1186/s13174-018-0087-2>
6. Burzykowski, T., Geubbelmans, M., Rousseau, A. J. and Valkenborg, D. (2023). Validation of machine learning algorithms. *American Journal of Orthodontics and Dentofacial Orthopedics*, 164(2), pp. 295-297. <https://doi.org/10.1016/j.ajodo.2023.05.007>
7. Chakri, P., Pratap, S. and Gouda, S. K. (2023). An exploratory data analysis approach for analyzing financial accounting data using machine learning. *Decision Analytics Journal*, 7(1), A. 100212. <https://doi.org/10.1016/j.dajour.2023.100212>
8. Chen, X., Cho, Y. H., Dou, Y. and Lev, B. (2022). Predicting future earnings changes using machine learning and detailed financial data. *Journal of Accounting Research*, 60(2), pp. 467-515. <https://doi.org/10.1111/1475-679X.12429>
9. Chen, X., Mao, Z. and Wu, C. (2025). Multi-class financial distress prediction based on feature selection and deep forest algorithm. *Computational Economics*, 66(4), pp. 2715-2754. <https://doi.org/10.1007/s10614-024-10761-8>
10. DeAngelo, H., DeAngelo, L. and Skinner, D. J. (1996). Reversal of fortune dividend signaling and the disappearance of sustained earnings growth. *Journal of Financial Economics*, 40(3), 341-371. [https://doi.org/10.1016/0304-405X\(95\)00850-E](https://doi.org/10.1016/0304-405X(95)00850-E)
11. Dechow, P. M., Sloan, R. G. and Sweeney, A. P. (1995). Detecting earnings management. *Accounting Review*, 70(2), pp. 193-225.
12. Esfahanipour, A., Goodarzi, M. and Jahanbin, R. (2016). Analysis and forecasting of IPO underpricing. *Neural Computing and Applications*, 27(3), pp. 651-658. <https://doi.org/10.1007/s00521-015-1884-1>
13. Faghani Makrani, K., Salehnejad, S. H. and Amin, V. (2016). Predicting earnings management based on modified Jones model using artificial neural network model and genetic algorithm. *Financial Engineering and Securities Management (Portfolio Management)*, 7(28), pp. 117-136. <https://dorl.net/dor/20.1001.1.22519165.1395.7.28.7.2>
14. Fereydouni, F., Darabi, R. and Anvari Rostami, A. (2020). Application of artificial intelligence algorithms in predicting earnings smoothing. *Financial Accounting and Auditing Research*, 12(45), pp. 103-134.
15. Hamidian, M. and Esmaeili, M. (2024). Forecasting earnings management using support vector machine and decision tree methods. *Research in Accounting and Economic Sciences*, 40(8), pp. 29-38.

16. Hammami, A. and Hendijani Zadeh, M. (2022). Predicting earnings management through machine learning ensemble classifiers. *Journal of Forecasting*, 41(8), pp. 1639-1660. <https://doi.org/10.1002/for.2885>
17. Hassani, H., Malekian Kallehbasti, E. and Kamyabi, Y. (2024). Explaining the earnings management prediction model using the hybrid of machine learning methods. *Financial Accounting Research*, 16(2), pp. 53-88. <https://doi.org/10.22108/far.2025.143145.2076>
18. Healy, P. M. (1985). The effect of bonus schemes on accounting decisions. *Journal of Accounting and Economics*, 7(1-3), pp. 85-107. [https://doi.org/10.1016/0165-4101\(85\)90029-1](https://doi.org/10.1016/0165-4101(85)90029-1)
19. Healy, P. M. and Wahlen, J. M. (1999). A review of the earnings management literature and its implications for standard setting. *Accounting Horizons*, 13(4), pp. 365-383. <https://doi.org/10.2308/acch.1999.13.4.365>
20. Huy, T. P., Hong, T. P. and Quoc, A. B. N. (2025). Leveraging tree-based machine learning for predicting earnings management. *Journal of International Commerce, Economics and Policy*, 16(02), A. 2550008. <https://doi.org/10.1142/S1793993325500085>
21. Jain, E., Neeraja, J., Banerjee, B. and Ghosh, P. (2022). A diagnostic approach to assess the quality of data splitting in machine learning. ArXiv, Cornell University, Ithaca, New York. <https://doi.org/10.48550/arXiv.2206.11721>
22. Jones, J. J. (1991). Earnings management during import relief investigations. *Journal of Accounting Research*, 29(2), pp. 193-228. <https://doi.org/10.2307/2491047>
23. Kasznik, R. (1999). On the association between voluntary disclosure and earnings management. *Journal of Accounting Research*, 37(1), pp. 57-81. <https://doi.org/10.2307/2491396>
24. Kordestani, G. and Tatli, R. (2014). Identification the efficient and opportunistic earnings management approaches in the earnings quality levels. *Accounting and Auditing Review*, 21(3), pp. 293-312 (In Persian).
25. Liashchynskyi, P., and Liashchynskyi, P. (2019). Grid search, random search, genetic algorithm: a big comparison for NAS. ArXiv, Cornell University, Ithaca, New York. <https://doi.org/10.48550/arXiv.1912.06059>
26. Maleki Kakler, H., Bahri Sales, J., Jabbarzadeh Kangarlooee, S., and Ashtab, A. (2021). Efficiency of statistical models and machine learning algorithms in predicting fraudulent financial reporting. *Financial Economics*, 15(54), pp. 267-292.
27. Montazeri, M., and Khodamoradi, M. (2017). Investigating earnings management models in various industries. First National Conference on New Researches of Iran and the World in Management, Economics, Accounting and Humanities, Shiraz, Iran (In Persian).
28. Mosalla Nejad, A., and Nazemi, A. (2019). Investigating the accuracy of financial information in the Rahavard Novin database. 17th National Accounting Conference of Iran, Qom, Iran (In Persian).
29. Muraina, I. (2022). Ideal dataset splitting ratios in machine learning algorithms: general concerns for data scientists and data analysts. In the 7th International Mardin Artuklu scientific research conference, Adeniran Ogunsanya College of Education, Lagos, Nigeria (pp. 496-504).
30. Patro, V. M., and Patra, M. R. (2014). Augmenting the weighted average with a confusion matrix to enhance classification accuracy. *Transactions on Machine Learning and Artificial Intelligence*, 2(4), pp. 77-91.
31. Rahman, R. A., Masrom, S., Zakaria, N. B., Nurdin, E., and Abd Rahman, A. S. (2021). Prediction of earnings manipulation on Malaysian listed firms: A comparison between linear and tree-based machine learning. *International Journal of Emerging Technology and Advanced Engineering*, 11(8), pp. 111-120. https://doi.org/10.46338/ijetae0821_13
32. Rahmani, A. and Bashirimanesh, N. (2013). Investigating the detection power of earnings

- management models. *Accounting and Auditing Research*, 5(19), pp. 54-73 (In Persian).
33. Rahnama Roodposhti, F., Banimehdi, B., Rezaei, S., and Salehi, A. (2014). Evaluation of the ability and explanation of discretionary accruals models and discretionary revenue model for detecting earnings management. *Journal of Financial Accounting and Auditing*, 6(21), pp.1-22.
 34. Rashidpour, R., and Zare Bidaki, A. M. (2024). Spam page detection using xgboost algorithm. *Iranian Journal of Electrical and Computer Engineering*, 22(4), pp. 287-294 (In Persian).
 35. Ronen, T., and Yaari, V. (2002). On the tension between full revelation and earnings management: a reconsideration of the revelation principle. *Journal of Accounting, Auditing & Finance*, 17(4), pp. 273-294. <https://doi.org/10.1177/0148558X0201700401>
 36. Saberi Gahruei, P., Saberi Gahruei, S., and Daghani, R. (2023). Investigating the difference in the quality of audit and profit management in bankrupt and healthy companies: matched groups method using machine learning. *Professional Auditing Research*, 3(10), pp. 56-71. <https://doi.org/10.22034/jpar.2023.562123.1117>
 37. Saghafi, A., and Pouryanasab, A. (2010). Theory of earnings management. *Accounting and Auditing Research*, 2(6), pp. 34-53 (In Persian).
 38. Salehi, A., and Salehi, B. (2015). A review of earnings management measurement models: discretionary accruals and discretionary revenue. *Knowledge of Accounting and Management Auditing*, 4(16), pp. 103-118 (In Persian).
 39. Salehi, M., and Farokhi Pile Rood, L. (2018). Predicting earnings management using a neural network and a decision tree. *Quarterly Journal of Financial Accounting and Auditing Research*, 10(37), pp. 1-24 (In Persian).
 40. Sarker, I. H., Kayes, A. S. M., Badsha, S., Alqahtani, H., Watters, P., and Ng, A. (2020). Cybersecurity data science: an overview from machine learning perspective. *Journal of Big Data*, 7(1), p. 41. <https://doi.org/10.1186/s40537-020-00318-5>
 41. Sarmad, Z., Bazargan, A., and Hejazi, E. (2004). Research methods in behavioral sciences. *Tehran: Agah Publication*, Tehran, Iran, pp. 7-132 (In Persian).
 42. Scott, W. R. (2015). Financial accounting theory. Seventh Edition. Pearson, United States: Canada Cataloguing, Canada.
 43. Siekelova, A. (2021). Historical development of earnings management models. In SHS Web of Conferences (Vol. 92, p. 02058). EDP Sciences, Les Ulis cedex A, France.
 44. Sun, L., and Rath, S. (2010). Earnings management research: a review of contemporary research methods. *Global Review of Accounting and Finance*, 1(1), pp. 121-135.
 45. Véganzones, D., and Séverin, E. (2025). Earnings management visualization and prediction using machine learning methods. *International Journal of Accounting Information Systems*, 56, A. 100743. <https://doi.org/10.1016/j.accinf.2025.100743>
 46. Wiens, M., Verone-Boyle, A., Henscheid, N., Podichetty, J. T., and Burton, J. (2025). A tutorial and use case example of the eXtreme gradient boosting (XGBoost) artificial intelligence algorithm for drug development applications. *Clinical and Translational Science*, 18(3), A. e70172. <https://doi.org/10.1111/cts.70172>
 47. Yadav, S., and Shukla, S. (2016). Analysis of k-fold cross-validation over hold-out validation on colossal datasets for quality classification. In 2016, IEEE 6th International Conference on Advanced Computing (IACC) (pp. 78-83). IEEE. New York
 48. Ying, X. (2019). An overview of overfitting and its solutions. In Journal of Physics: Conference series (Vol. 1168, p. 022022). IOP Publishing, Bristol, England. <https://doi.org/10.1088/1742-6596/1168/2/022022>