

Strathmore
UNIVERSITY

Predictive modeling of fraudulent transactions for regulated Deposit-Taking SACCOs in Kenya.

Mwangi, Peris Wanjiku

145850

**Submitted in partial fulfillment of the requirements for the Degree of
Bachelor of Business Science in Actuarial Science at Strathmore University**

**Strathmore Institute of Mathematical Sciences
Strathmore University**

Nairobi, Kenya

January 2025

This Research Project is available for Library use on the understanding that it is copyright material and that no quotation from the Research Project may be published without proper acknowledgment.

DECLARATION

I declare that this work has not been previously submitted and approved for the award of a degree by this or any other University. To the best of my knowledge and belief, the Research Project contains no material previously published or written by another person except where due reference is made in the Research Project itself.

© No part of this Research Project may be reproduced without the permission of the author and Strathmore University

Mwangi, Peris Wanjiku.

..... [Name of Candidate]



..... [Signature]

31st January 2025

..... [Date]

This Research Project has been submitted for examination with my approval as the Supervisor.

Dr. Mary Achieng Ochieng.

..... [Name of Supervisor]



..... [Signature]

31st January 2025

..... [Date]

Strathmore Institute of Mathematical Sciences
Strathmore University.

Table of Contents

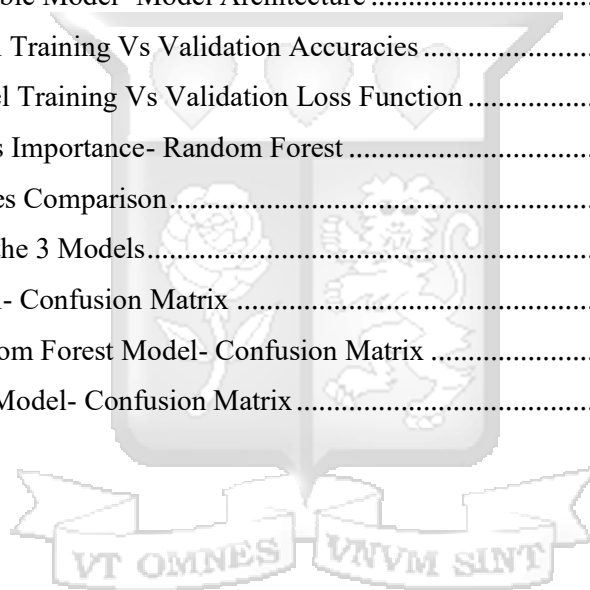
List of Figures	v
List of Tables.....	vi
List of Abbreviations	vii
Abstract	viii
Chapter 1	2
1.0. Introduction.....	2
1.1. Background of the study.....	2
1.2. Research problem.....	5
1.3. Research objectives	6
1.4. Research questions	6
1.5. Significance of the study	6
1.6. Scope of the study	7
Chapter 2.....	8
2. Literature review	8
2.0. Introduction.....	8
2.0. Theoretical literature review.....	8
2.1. Empirical literature review	11
2.2. Proposed model	16
2.3. Gap in literature.....	16
2.4. Summary of the literature review	17
Chapter 3.....	18
3. Methodology.....	18
3.0. Introduction	18
3.1. Data collection.....	18
3.2. Data preprocessing	18
3.3. Model type, technique, and justification.....	19
3.4. Model creation, training, and testing	23
3.5. Model architecture.....	23
3.6. Model optimization	32
3.7. Model evaluation.....	34
Chapter 4.....	37
4. Data analysis and results	37

4.0.	Introduction	37
4.1.	Data preprocessing	37
4.2.	Model training- proposed model	39
4.3.	Model optimization	42
4.4.	Model evaluation.....	43
Chapter 5.....		48
5. Discussion and conclusion		48
5.0.	Introduction	48
5.1.	Summary of findings and conclusion	48
5.2.	Recommendations	50
5.3.	Limitations of the study.....	51
5.4.	References	52



List of Figures

Figure 1: Fraud Triangle Theory	9
Figure 2: Fraud Diamond Theory	10
Figure 3: Routine Activity Theory	11
Figure 4: Bagging Technique.....	20
Figure 5: Boosting Technique.....	20
Figure 6: Ensemble Model.....	23
Figure 11: ReLU Activation Function Graph.....	26
Figure 12: Sigmoid Activation Function Graph.....	27
Figure 13: Multilayer Perceptron- Model Architecture	28
Figure 14: Proposed Ensemble Model- Model Architecture	29
Figure 15: Ensemble Model Training Vs Validation Accuracies	40
Figure 16: Ensemble Model Training Vs Validation Loss Function	41
Figure 17: Hyperparameters Importance- Random Forest	42
Figure 18: Model Accuracies Comparison.....	44
Figure 19: ROC Graph for the 3 Models.....	45
Figure 20: Ensemble Model- Confusion Matrix	46
Figure 21: Optimized Random Forest Model- Confusion Matrix	46
Figure 22: Random Forest Model- Confusion Matrix	47



List of Tables

Table 1: Algorithms Performance Comparison 1	21
Table 2: Algorithms Performance Comparison 2.....	22
Table 3: Algorithms Performance Comparison 3.....	22
Table 4: Base- Model Accuracies	39
Table 5: Ensemble model- Evaluation Metrics.....	43
Table 6: Optimized random forest model- Evaluation metrics	43
Table 7: Unoptimized random forest model- Evaluation metrics	43
Table 8: Summary of Models' Performance.....	48



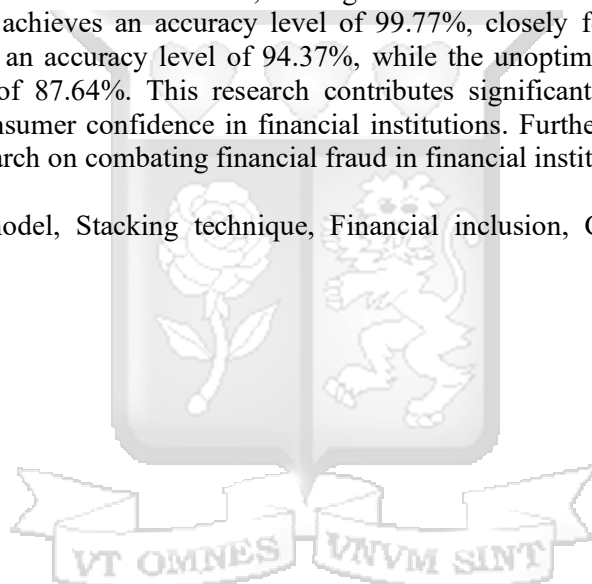
List of Abbreviations

Adam	Adaptive Moment Estimation
API	Application Programming Interface
AUC	Area Under the Curve
BCE	Binary Cross-entropy
BO	Bayesian Optimization
CBK	Central Bank of Kenya
DT-SACCOS	Deposit-Taking SACCOs
EI	Expected Improvement
FN	False Negative
FOSA	Front Office Savings Activities
FP	False Positive
GDP	Gross Domestic Product
GP	Gaussian Process
KYC	Know Your Customer
MLP	Multilayer Perceptron
NWDT-SACCOS	Non-Withdrawable Deposit Taking SACCOs
OOF	Out-of-fold
PwC	Price Waterhouse Coopers
ReLU	Rectified Linear Unit
SACCO	Savings and Credit Co-operatives Society
SASRA	SACCO Societies Regulatory Authority
SMOTE	Synthetic Minority Oversampling Technique
SSFIU	SACCO Societies Fraud Investigation Unit
TN	True Negative
TP	True Positive

Abstract

This study focuses on the issue of transactional fraud within Kenyan SACCOs, and in particular, regulated DT-SACCOs. The primary objective is to develop a comprehensive model that is capable of identifying and flagging fraudulent transactions before they are completed. This paper proposes an ensemble model that employs the stacking technique, which merges two distinct classifiers. Random forest (ML algorithm) forms the base model while multilayer perceptron (DL algorithm) forms the meta-model. Given that the proposed model utilizes more than one algorithm, it is expected to yield more accurate, precise, and reliable predictions in comparison to standalone algorithm-based models. The data used is obtained from one of the top-performing regulated DT-SACCOs in Kenya. However, the name of the SACCO is not disclosed to ensure adherence to data privacy regulations. The data contains more than 900,000 rows and 11 columns- which are primarily transaction attributes such as transaction ID, and amount, among others. Once the model is run, various performance metrics such as accuracy, precision, recall, and F1-score are used to analyze its effectiveness. The performance of the proposed model is also compared to that of a standalone random forest model in order to evaluate the models' relative performance. From the results obtained, it is evident that the ensemble model outperforms both the unoptimized and optimized random forest model, making it suitable for transactional fraud detection in SACCOs. For example, it achieves an accuracy level of 99.77%, closely followed by the optimized random forest model, with an accuracy level of 94.37%, while the unoptimized random forest model attains an accuracy level of 87.64%. This research contributes significantly to enhancing financial inclusion by increasing consumer confidence in financial institutions. Furthermore, this study lays the groundwork for future research on combating financial fraud in financial institutions.

Key Words: Ensemble model, Stacking technique, Financial inclusion, Consumer Confidence, and performance metrics.



Chapter 1

1.0. Introduction

1.1. Background of the study

1.1.1. Financial system in Kenya

The financial system in Kenya, which is recognized as one of the most advanced systems in sub-Saharan Africa, (Ogola, 2016), is commonly divided into two main sectors: banking and non-banking. These are also comprised of different financial entities. The banking sector includes both the Central Bank of Kenya (CBK) and the commercial banks, whereas the non-banking sector comprises insurance firms, capital markets, pension funds, and Saving and Credit Cooperatives (SACCOs), (Ng'el, 1990). Over the years, Kenya's financial system has experienced significant transformations and continuous expansion both in its number of institutions and assets. In 1993, Kenya's financial landscape comprised only 32 commercial banks and 55 non-banking entities, (Ogola, 2016). By 2022, according to (CBK, 2023), the financial sector had evolved to include a total of 38 commercial banks, 1 mortgage finance company, and 14 deposit-taking micro-finance banks, all overseen by the CBK. Within the insurance sector, there were 56 insurance companies, 1 micro-insurance company, 5 re-insurance companies, 210 insurance brokers, and 13,508 insurance agents. Additionally, focusing on SACCOs, the subject of this paper, there were a total of 361 licensed and authorized SACCOs, comprising 176 Deposit-Taking (DT) and 185 Non-Withdrawable Deposit-Taking (NWDT) SACCOs, all regulated by SASRA.

1.1.2. History and development of SACCOs

Over time, the significant expansion of Kenya's financial system has been driven by digital transformation, increased financial literacy, and the need to enhance financial inclusion, all of which contribute positively to the country's economic growth. Of notable significance are SACCOs which have undergone massive growth not just in terms of numbers, but also in terms of members and financial capacity (Rawal, 2021). SACCOs are voluntary associations or cooperative financial institutions owned and governed by their members through an elected board. Their primary purpose is to encourage savings, offer low-interest credit, and provide other financial services to members. The SACCO industry is part of Kenya's cooperative sector, which has its roots in the cooperative movement that started in the country long before independence. (Kabaiku, 2018).

The earliest SACCO dates to the year 1908 when the first cooperative society was formed. Later, after independence, as the number of cooperatives increased, the number of SACCOs equally grew (Motompa, 2016). SACCOs tend to offer a variety of products. Such include credit, savings, and investment options, which all depend on the type of SACCOs at hand. There are usually two broad categorizations of SACCOs in Kenya. On the one hand, we have Deposit-taking SACCOs (DT-SACCOs) which allow members to

make deposits and withdrawals at their own time. This is made possible through the Front Office Savings Activities (FOSA) they offer. On the other hand, we have Non-withdrawable-deposit-taking SACCOs (NWDT-SACCOs) which allow members to make deposits but deny them a chance to withdraw. In this case, withdrawals are only made if the member wishes to exit the SACCO at any time.

As previously noted, the growth of SACCOs holds significance not only to their members but also to the broader economy. For instance, as of 31st December 2021, DT-SACCOs collectively possessed total assets amounting to KES 691.09 billion, while NWDT-SACCOs held assets totaling KES 116.02 billion. These figures represented 5.71% and 0.96% of the nation's Gross Domestic Product (GDP), respectively. The funds were allocated to around 5.99 million SACCO members, including 5.54 million in DT-SACCOs and 460,785 in NWDT-SACCOs. These statistics highlight the significant impact SACCOs have on boosting the country's economy and improving the well-being of individuals, families, and communities, (Opany, 2022). The following year, according to (SASRA, 2022), SACCOs contributed around 6.6% to the GDP, maintaining a similar level as the previous year. This stagnated percentage contribution was attributed to the economic hardships experienced that year, including severe drought, high inflation rates, and high interest rates. However, despite all those economic hardships, the savings and deposits mobilized still grew by 9.84% in 2022 compared to 9.80% in 2021, while the loan advances grew by 11.76% in 2022 compared to 9.67% in 2021.

1.1.3. Regulatory framework

SASRA serves as the primary agency for the government and is responsible for overseeing and regulating SACCO Societies in Kenya. Established under the Sacco Societies Act, No. 14 of 2008, its key role is to issue licenses to SACCO Societies conducting deposit-taking activities (FOSA) in Kenya. Additionally, the authority supervises and regulates both Deposit-Taking (DT) and Non-Withdrawable Deposit Taking (NWDT) SACCO Societies. Furthermore, it bears the responsibility of helping SACCOs address the obstacles they encounter. A few of these challenges include poor market research, the use of outdated technology that prevents them from competing favorably in the financial industry, credit risk that arises from borrowers who default on their loans, poor leadership by the elected board of directors who also tend to lack vision, political interference, lack of innovation, and fraud. To address some of these challenges, for example, fraud, and in particular, employee fraud, SASRA in 2020, established the Sacco Societies Fraud Investigation Unit (SSFIU), which would be responsible for instilling integrity, accountability, transparency, and good governance in the subsector. With the new unit, it is expected that cases of fraud, theft, and embezzlement of members' savings are expected to go down. Consequently, consumer confidence is expected to go up hence increasing membership in SACCOs (SASRA, 2020).

1.1.4. Fraud

According to (ACFE, 2022) fraud refers to actions that are deceptive, usually done for personal or financial gain. It's an act conducted voluntarily, intending to bereave of the other party's property or money by guile, deceit, or other unjust means. (Skalak et al., 2015) defines fraud as the violation of an individual's expectation of fair treatment by others, often resulting in financial loss. Occupational fraud occurs when individuals working for or conducting business with an organization engage in dishonest activities. This type of fraud is further divided into three distinct categories. They include asset misappropriation, corruption, and financial statement fraud. Asset misappropriation takes place when employees of an organization steal or misuse its resources. Corruption, in contrast, involves individuals abusing their influence in business transactions to gain personal benefits or advantages for others, such as friends. Examples of corruption include concealing conflicts of interest, accepting illegal gifts, and offering bribes to secure favorable business outcomes. Lastly, financial statement fraud involves deliberately falsifying financial reports to deceive stakeholders like investors, creditors, or government agencies, (ACFE, 2022). In this paper, however, the focus will be transactional fraud, which can either be committed by employees of a financial institution, or by external perpetrators. A fraudulent transaction is an illegal or unauthorized activity that involves using payment methods or financial systems to gain money, goods, or services without the account holder's consent. This type of transaction often includes identity theft, stolen payment details, or deception, to mislead and cause financial harm to the account holder, business, or financial institution, (Steven, 2024).

Fraud is considered a global issue that all organizations face. It is a complex problem, which if left unattended, yields catastrophic results such as financial losses or even a complete shutdown of businesses. It has been in existence for a long time and continues to adapt to new market changes. For example, with the continuous uptake of technology in most businesses, new forms of fraud have emerged. In reference to PwC's 2020 global economic crime and fraud survey, 49% of participants reported that their companies had fallen victim to fraud or economic crimes. Around 45% claimed to have suffered losses of under a hundred thousand dollars; 30% were preys of losses between a hundred thousand dollars and five million dollars; 6% experienced losses ranging from five million dollars to fifty million dollars; and 3% faced losses exceeding fifty million dollars. This highlights the growing financial impact of fraud. Within organizations, 52% of cases were attributed to internal fraud, while 41% were linked to external sources, (Sánchez-Aguayo et al., 2022). In Kenya, the most affected industry is the financial industry, with some of the contributing factors being, technological advancements, legal factors, personal reasons, and management factors (Beatrice, n.d.). To deal with fraud, several solutions have been proposed and even adopted in some institutions. They include a) the adoption of strong internal control systems, (Wanyama & Reuben, 2022), b) offering ethical training to all employees before they start working in the institution, c) adopting the

Know Your Customer (KYC) policy, whereby private information is obtained and verified before signing up for the contract. Such information includes the customer's driving license, identification details including the full name, date of birth, ID number, mobile number, and PIN & address (Mwithi, 2015), d) reporting irregularities upon occurrence (ACFE, 2022), e) in addition to offering ethical training to the potential employees, conducting due diligence is also encouraged (Bartsiotas & Achamkulangare, n.d.), f) upholding whistleblowing policy, g) conducting fraud training and raising awareness (ACFE, 2022), h) embracing modern technologies and implementing effective security measures. It's important to note that modern technologies is more than just digitizing transactions within financial institutions; it also involves leveraging Artificial Intelligence (AI) to combat fraud. For example, (Ojino & Ndolo, 2023) in their study proposed using knowledge graphs to detect transactional fraud in SACCOs. In addition, the paper proposed the possible use of a model combining data analytics and knowledge graphs as this would easily identify complex patterns.

Despite the implementation of some of the proposed solutions, fraud continues to pose a threat to the Kenyan financial system. Cases of fraud have been continuously reported over the years. Equity Bank, the largest bank in Kenya, recently lost around 2.1 million dollars to fraud. The stolen funds were then transferred to 500 bank and mobile accounts, most of which have been restricted. To recover the lost funds, 19 individuals believed to be involved in the fraud were arrested (Ndege, 2024). In the year 2022, ABSA Bank lost 107 million shillings to fraud. The bank was, however, able to recover 59 million shillings, which was more than half of the lost funds. It also clearly outlined in its report that fraud was a threat to the Kenyan financial sector hence the need to be addressed urgently (ABSA, 2023).

1.2. Research problem

Fraud poses a threat to SACCOs, despite their crucial role in expanding access to financial services, especially among underserved communities where formal services are limited, and in providing credit for small and medium-sized enterprises which are the backbone of the Kenyan economy, (Njoki, 2023). In the last two years, they have lost 118 million shillings to fraud, particularly insider fraud. On the same note, in their 2022 annual report, SASRA revealed that between 2021 and 2022, it had dealt with fraud cases amounting to 233 million shillings (SASRA, 2022). In the same report, SASRA mentioned that in June 2022, they had witnessed the withdrawal of 29.4 million shillings from some of the savings accounts belonging to different members using unregistered mobile phone numbers, after which the money was transferred to different M-Pesa accounts. In another instance, still in the same year, another SACCO reported an illegal withdrawal of 1.6 million shillings in October. This withdrawal was however linked to an internal staff who used unauthorized mobile numbers to make the withdrawal (Alushula, 2023). Continuous reporting of fraud is likely to discourage individuals from participating in the usage of financial products due to fear of potential financial losses, lack of confidence, and lack of trust in financial institutions,

particularly considering the risk-averse nature of many people. This could have a detrimental effect on the government's financial inclusion agenda which entails a situation in which all people who can utilize financial services have access to high-quality financial services that are offered at reasonable prices, conveniently, and with dignity to customers, (Amanja, 2015).

Despite their role in promoting financial inclusion and economic development in Kenya, there is a significant lack of tailored fraud detection models that address their unique operational challenges. The existing research predominantly focuses on fraud detection methods suitable for larger financial institutions such as banks, emphasizing financial statement fraud and credit card fraud, while neglecting the more prevalent issue of transactional fraud in SACCOs. Furthermore, most studies are biased toward European or external economies, limiting their relevance and applicability to the Kenyan context. There is therefore a need to come up with better ways of addressing fraud in the financial system of Kenya. This paper will focus on addressing fraud risk, and in particular, fraudulent transactions, through creating a more comprehensive predictive model for detecting fraudulent transactions, that will enable SACCOs to be more prepared and put necessary measures in place to help minimize the chances of the risk occurring. This will in return have a positive impact on financial inclusion in Kenya.

1.3. Research objectives

- a) To develop a predictive model for detecting fraudulent transactions for regulated deposit-taking SACCOs in Kenya.
- b) To assess the performance of the predictive model.

1.4. Research questions

- a) How can transactional fraud be effectively predicted for regulated deposit-taking SACCOs in Kenya?

1.5. Significance of the study

1.5.1. Contribution to financial inclusion

By addressing the issue of transactional fraud in SACCOs, this study will play a critical role in promoting financial inclusion in Kenya. Fraudulent activities often hinder individuals from engaging with financial institutions due to fear of losing their funds and lack of trust. By developing a predictive model that enhances fraud detection and mitigation, this research aims to restore confidence among current and potential SACCO members. This will encourage more people, especially those in underserved communities, to participate in the financial system, thus supporting the government's financial inclusion agenda and contributing to the overall economic development of the country, (Amanja, 2015).

1.5.2. Benefits to shareholders and financial performance

For shareholders and other stakeholders of SACCOs, this study offers significant benefits as it safeguards their investments from fraud-related losses. Enhanced fraud detection and prevention mechanisms will ensure the stability and profitability of SACCOs, thereby protecting shareholder value. Furthermore, by reducing financial losses due to fraud, SACCOs can allocate more resources towards growth and development, improving their financial performance and competitiveness in the financial sector. This increased financial stability and growth potential will make SACCOs more attractive to investors, thus fueling their expansion and contribution to the economy, (Predictive Analytics, 2023).

1.5.3. Contribution to existing literature and broader financial institutions

Academically, this study will contribute to the existing body of knowledge on fraud detection and financial security in SACCOs and other financial institutions. The development of a comprehensive predictive model will serve as a reference for future research in financial fraud detection, offering insights and methodologies that can be adapted and applied to different contexts. Additionally, other financial institutions can leverage the findings and recommendations from this study to improve their fraud detection systems. By sharing best practices and innovative solutions, the study aims to foster a more secure and resilient financial ecosystem in Kenya and potentially in other regions facing similar challenges.

1.6. Scope of the study

The research paper is solely based on the financial system of Kenya with a specific focus on licensed and authorized deposit-taking SACCOs. DT-SACCOs are selected for use in this paper because transactions are more frequent compared to NWDT-SACCOs, which only allow for the withdrawal of funds upon exit.

By examining both fraudulent and non-fraudulent transactional data, the study aims to identify patterns and anomalies in transactional data and to develop a predictive model for detecting and mitigating fraudulent transactions. In addition to that, the effectiveness of existing models will be evaluated.

Chapter 2

2. Literature review

2.0. Introduction

This chapter will study theories that explain transactional fraud in financial institutions. It will also review past studies done to address the issue of transactional fraud, with a specific focus on the proposed models. The advantages, limitations, and drawbacks of these models will be evaluated to develop a more comprehensive model that will help detect fraudulent transactions.

2.0. Theoretical literature review

Organizations need to be able to detect fraudulent transactions and come up with better methods of mitigating this risk. Preventing fraud helps financial institutions save time and resources. To help address the issue of fraud and halt it, organizations should first understand the underlying issues by identifying the incentives that drive individuals to commit fraud as well as their behavior, (Ruankaew, 2013). Fraud is usually tied to human behavior thus, understanding the motivations, psychological, and personality traits of the perpetrators can provide a new perspective on fraud detection, (Sayal & Singh, 2020). To help understand the relationship between human behavior and fraud, many theories have been developed in the recent past, with the most common ones including the following:

2.0.1. Fraud triangle theory

This theory was developed as part of a PhD study in criminology, by a leading expert in sociology of crime, Dr. Donald R. Cressey, in 1953 (Sujeewa et al., 2018). Cressey aimed at investigating why people committed fraud, and in particular, occupational fraud. This was done through conducting interviews with fraud convicts in the United States. According to the findings, three key elements come into play for occupational fraud: pressure, opportunity, and rationalization (Sánchez-Aguayo et al., 2021). According to Cressey, pressure serves as a driving force that might lead an individual to commit fraud. This pressure can arise from personal issues like financial difficulties, addiction challenges, or factors within the workplace environment, (Sujeewa et al., 2018). (ACFE, 2022) describes pressure as a personal issue, often financial, that forces the victim to engage in fraud. Gambling, drug addiction, personal debt, poor credit, major financial loss, or the influence of peers or family to achieve success are all considered to be different forms of pressures. People may view fraud as their only solution due to reasons like shame, pride, or the desire to prove themselves. The second element, opportunity, refers to the perceived ability to commit fraud without being caught, usually arising from weak internal controls or ineffective monitoring in financial institutions. Rationalization involves the perpetrator justifying their fraudulent actions. According to Cressey, rationalization happens before the fraudulent act is conducted. This is usually the case because the perpetrators often avoid seeing their behavior as wrong or unethical. To maintain a positive self-image,

they rationalize their fraudulent actions in various ways. They might convince themselves they are merely "borrowing" the money intending to repay it later, or they may feel they are underpaid and thus entitled to additional compensation., (ACFE, 2022). It is worth noting that although this theory focuses on occupational fraud, external fraudsters equally follow the same analogy/ rationale when planning to conduct any fraudulent activity including transactional fraud.

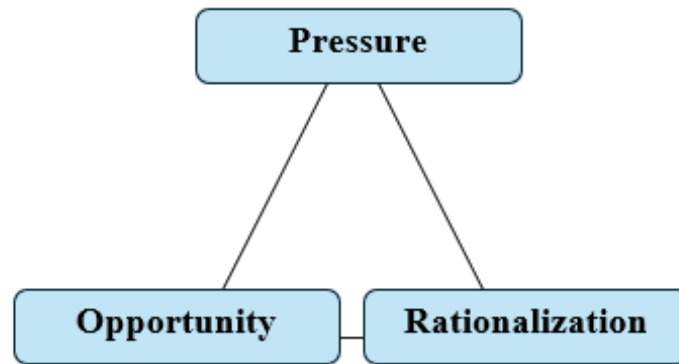
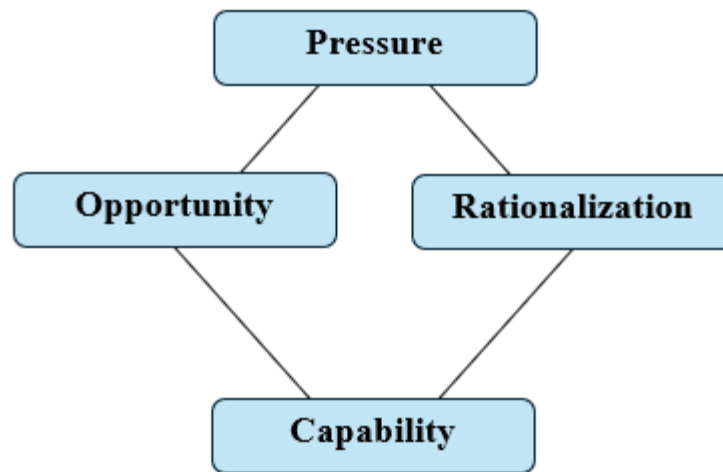


Figure 1: Fraud Triangle Theory
Source: (ACFE, 2022)

2.0.2. Fraud diamond theory

This theory is said to be an improvement of the fraud triangle theory. Developed in the year 2004, by Wolfe and Hermanson, the theory added a fourth element to the fraud triangle theory, which is capability, (Abosede, 2023). The two argued that regardless of how big the opportunity is, how much pressure one is receiving, or even the level of rationalization, the existence of the three elements is not enough to foster the execution of a fraudulent activity. The potential perpetrator must have the capability/ skills/ know-how to fully maximize the identified opportunity, otherwise it will never be a reality. Some of the observable traits possessed by such individuals included the following: a) they occupy a higher position in the organization, b) they are intelligent and equipped with knowledge regarding internal control system functionality, c) they are confident and can equally deal with stress, (Sujeewa et al., 2018). As previously highlighted, the above theory equally relates to external fraudsters who indulge in different forms of transactional fraud such as identity theft, fake payment methods, credit card fraud, etc. This is because they also take into consideration the elements mentioned in the theory. In other words, the theory is not just restricted to corporate fraud.



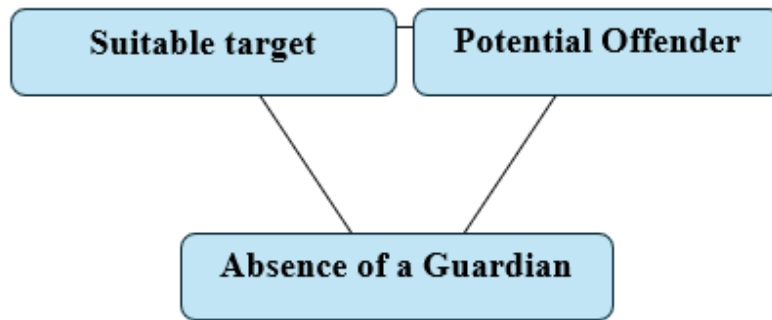
*Figure 2: Fraud Diamond Theory
Source: (Abosedo, 2023)*

2.0.3. Routine activity theory

This theory dates back to the year 1979 when it was established by two gentlemen, Cohen and Felson, to explain the relationship between routine activity and crime rate. The theory revolves around three key factors: a potential offender, a suitable target, and the absence of a capable guardian, (Andresen & Ha, 2017). A potential offender is an individual who is ready to execute a felony, while a suitable target refers to something that might motivate the potential offender to commit a crime. A capable guardian on the other hand refers to something that could have prevented the occurrence of a crime, (Kigerl, 2021). The theory findings indicated that changes in the usual routine when conducting daily activities could have an impact on the probability of an individual being a victim of a crime or even increase the chances of an individual committing a crime. For example, walking too much for an individual increases the chance of encountering a burglar, and for the offender, spending more time with convicts might increase the chance of committing a crime. The guardian in this case would be people who might intervene and stop the burglar, (Andresen & Ha, 2017). Now, in the context of transactional fraud, the same analogy is used. A suitable example would be during the COVID-19 pandemic when most businesses had to operate online, which was a change from the usual way of doing business. Those who were still new to online trading became more exposed to online fraudsters, especially if they had weak monitoring systems, which would indicate the absence of a guardian. A potential offender would be the equivalent of a fraudster who possesses the personal details of an account holder. The fraudster could disguise themselves as the account holder- identity theft and withdraw cash. This theory proposes that financial institutions should take time and closely monitor the behaviors of their customers to ensure that they can detect any anomalies or changes in patterns in how they conduct their transactions. Furthermore, encrypting online transactions through the usage of advanced technologies

would help curb transactional fraud, especially for those new to online business which would be a form of guardian.

Understanding how fraudsters think is essential when it comes to addressing transactional fraud. This forms the basis for the models that have been established to tackle the same.



*Figure 3: Routine Activity Theory
Source:(Abosede, 2023)*

2.1. Empirical literature review

Under this section, the paper will explore past studies that have been done to address transactional fraud. The already existing transactional fraud detection models will be critically evaluated to identify their limitations, which will also form the basis for the proposed model. Some of the already existing approaches that are being used to mitigate transactional fraud in institutions include the following:

2.1.1. Rule-based approach

It is a fraud detection method that relies on a set of "rules" or conditions that, if identified, would indicate transactions that could potentially be fraudulent., (Arya, 2024). The rules are usually created on historical data, expert knowledge, and known fraud patterns. For instance, some of the most common rules used are:

- a) **Location:** Questions asked under this are on whether the transaction is made outside the common area where the user normally makes most of their transactions.
- b) **Frequency:** Is the account being used more frequent than usual? Is there an increment in the average number of transactions made in a day?
- c) **Unknown sender/ receiver:** Are the accounts receiving money new? Do the accounts use the same IP address?

Now, any transaction that does not meet any of the criteria established for identifying fraud, is classified as a non-fraudulent transaction. The model allows the completion of some transactions and denies others based on the rules they satisfy. These rules are typically if-then conditional statements or rules based on decision tables or trees that specify what constitutes suspicious activity. The **IF** section is referred to as the

antecedent, premise, or condition, while the **THEN** section is known as the consequent, conclusion, or action, (J. Kumar & Saxena, 2022). For example,

IF < condition > THEN < action >

IF < condition – 1 > AND < condition – 2 >AND < condition – n > THEN < action >

Rule-based systems play a significant role in fraud detection by providing a structured and deterministic approach to identifying suspicious activities. The limitation of this method is that when adjustments are needed to detect a new type of fraud, they must be made manually, either by modifying the existing algorithms or by developing new ones. As the number of customers and the volume of data grow, the amount of human effort required also increases. Also, with the increment in the number of people adopting new technologies, more fraudulent cases are expected. So, in such cases, the rule-based approach becomes a time-consuming, costly, and unsustainable approach in the future, (Mugoshi, 2021). Another disadvantage of this approach is its tendency to produce false positives, where the test incorrectly indicates the presence of a condition that doesn't exist. For example, flagging a non-fraudulent transaction as fraudulent. The outcome of a transaction relies on the rules and guidelines used to train the algorithm for identifying non-fraudulent transactions. Therefore, with a fixed risk threshold, if a legitimate transaction is mistakenly flagged, it leads to higher rates of false positives. This false positive condition will result in losing genuine customers, who opt to work with competitors, which translates to loss of potential income. Lastly, another drawback of this approach is that over time fraudsters learn to behave to go undetected. Therefore, the static rule-based approach is outperformed by self-learning defense, adopted by fraudsters, (Vanini et al., 2023). Other researchers have attempted to address the model's drawbacks and improve it altogether. For example, (Islam et al., 2024) proposed a rule-based machine learning model to classify financial transactions as fraudulent or non-fraudulent. The new model could now be able to predict credit card fraud, by going beyond the rules. In addition to that, (J. Kumar & Saxena, 2022) in their study, introduced keystroke dynamics, which would be an additional rule to consider when classifying whether a transaction is fraudulent or non-fraudulent. Keystroke dynamics is a behavioral biometric security method that enhances key-based authentication by analyzing a user's typing patterns. It adds an extra layer of security to detect intruders during online transactions. In particular, the user's time difference between pressing and releasing a key when entering the 4-digit pin would be calculated and an acceptance time range would be set. Therefore, if in any case the range is exceeded, the transaction is deemed suspicious and, hence, terminated. The drawback of this dynamic is that in cases where the fraudster can match the holders' time difference between pressing and releasing a key, then the transaction is usually categorized as a false positive, thus making the model unreliable for practical use.

2.1.2. Artificial intelligence (AI) based approach

The most common application of AI is machine learning (ML), which is considered to be a subset of AI that has a crisp on advancing and studying statistical algorithms that can learn from data to perform tasks without explicit instructions. Machine learning offers an advantage over rule-based fraud detection systems, as it can identify patterns in data and adapt to new, unforeseen situations. This allows machine learning models to adjust to evolving fraud techniques, unlike the rigid rule-based approach. Fraud detection typically utilizes two main categories of machine learning models. These are:

- a) **Supervised models-** These models are typically trained on labeled inputs meaning the data is categorized. For instance, a transaction is labeled as either 'fraudulent' or 'non-fraudulent.' Large volumes of such labeled data are used to train the supervised learning model to produce accurate results. The model's accuracy depends heavily on how well-organized the data is. If the labels are incorrect, the supervised algorithms will learn inaccurate information, leading to faulty outputs, (Ali, 2022).
- b) **Unsupervised models-** These on the other hand are built to detect unusual behavior in transactions that have not been detected previously. The data used here is usually not labeled. Unsupervised learning models rely on self-learning to identify hidden patterns within transactions. In this approach, the model independently analyzes the available data and looks for similarities and differences in transaction occurrences. This process aids in detecting fraudulent activities.

Under each of the above two machine learning classifications, there exist different important algorithms, with the most used in financial fraud detection being:

- a) **Logistic regression-** This algorithm is a special linear regression case, where the predicted variable is categorical. It is a popular statistical method that utilizes the logit functions to make predictions about the probability of a binary event occurring, (Ali, 2022). It is a form of a supervised model, thus the data used is labeled. In the case of detecting fraud, after thorough analysis of the features of the transaction, such as location, transaction amount, frequency, age, and gender, the algorithm classifies the transaction as either fraudulent or non-fraudulent, (Mugoshi, 2021). The drawback of this algorithm is over-fitting training data which is not desirable. In addition to that, it does not perform well on non-linear datasets.
- b) **Decision tree-** It is used to solve both regression and classification problems. This algorithm classifies a transaction as either legitimate or fraudulent by recursively partitioning the data into subsets based on the values of different features or appropriate descriptors. Two major steps are involved, with the initial being building a decision tree using the data provided, and the last step entailing the use of decision rules to classify the incoming transactions, (Joshi, 2021). The decision trees built contain nodes and branches. Each internal node of the tree represents a check on an attribute, and each branch indicates the result of that check. Additionally, every terminal node, or

leaf node, holds a class label, (Mytnyk et al., 2023). Although the decision trees are relatively efficient, have a higher accuracy percentage, and can handle large datasets, their drawback is that they require a lot of time to train the model. They are also very sensitive to data changes and tend to over-fit the data as well.

- c) **Random Forest**- Unlike the previous algorithm that contains only one decision tree, random forests contain multiple decision trees which are then combined to make a prediction. Each of the trees explores a different feature of the data given, (Negi, 2021). For example, one tree could check on the frequency, while the other looks at the transaction amount. The multiple decision trees form what is known as a forest, whose strength increases as the number of trees increases. To decide whether a transaction is fraudulent or not, the algorithm utilizes the average of all the predictions made by multiple trees. According to (Han et al., 2020), a random forest algorithm typically outperforms a single decision tree. It is more resistant to noisy data and less likely to overfit due to the use of multiple trees. Additionally, it can effectively manage high-dimensional data. Since each tree in the random forest is constructed independently, the training and prediction processes can be parallelized, accelerating the computationally demanding model training. The advantages of the model (Algorithm) include its high level of accuracy, and ability to analyze large data sets, (Negi, 2021). Its downside, on the other hand, is that it is slower when it comes to training large data sets. Further to that, when used in solving regression problems, it is unable to predict beyond the variety in the training of the data, (Mugoshi, 2021).
- d) **K-Nearest Neighbor** – It is a supervised non-parametric algorithm that groups different data points by analyzing their proximity. Proximity in this case represents the degree to which the data points are comparable to one another, while non-parametric means that the algorithm does not make any assumptions on the underlying data set. It is based on the principle that similar items, sharing the same characteristics are likely to be found in the same class. According to, (Khodabakhshi & Fartash, 2016), K-nearest neighbor algorithm is the best clustering algorithm in which the test sample belongs to the class that has the most votes in the K-nearest neighbor. To assign an item in the new data set to a class, the average distance between that item and all the other items in the training data is also calculated in addition to comparing the characteristics. The shortest distance is where the new data point falls, (Alammar et al., 2022). In fraud detection, this algorithm compares the new transactional data traits with the training data, and specifically, the character traits portrayed by legitimate transactions. If a transaction portrays more similarity to the legitimate transactions in the training data set, it is categorized as non-fraudulent. On the other hand, if the transaction in the new data set portrays a high level of dissimilarities, it is flagged as a fraudulent transaction. One of

the drawbacks of this algorithm is that it is slow when making the classifications. In addition to that, it is costly especially when handling large data sets, (Lai, 2023).

- e) **Knowledge graph-** This is a graphical representation capturing connections that exist between different entities. It is made up of nodes and edges. The nodes represent concepts or entities, while the edges represent the relationships between those entities, (Yasar, 2024). It functions by structuring and linking information in a formatted, graph-like arrangement. According to (Ojino & Ndolo, 2023), the scalability of the developed knowledge graph makes it well-suited for real-time or near-real-time fraud detection. It can manage large amounts of transactional data and quickly process queries to identify potential fraud as it happens. By analyzing the graph in real time, suspicious activities can be detected immediately, allowing for timely actions to mitigate potential losses.
- f) **Neural networks-** The functionality of this algorithm resembles that of a human brain. It utilizes cognitive computing that enables it to extract intricate patterns and relationships from transactional data hence capturing any indicators of fraudulent transactions. In addition, this algorithm contains multiple layers of interconnected nodes, known as neurons, that are solely responsible for passing information to the next computational layer. The final decision made by the algorithm is usually made by the final layer, (Mugoshi, 2021). Again, in detecting fraud in transactional data, the algorithm tends to learn from historical data from which it identifies patterns or anomalies that indicate fraudulent behavior. Some of the transactions flagged as fraudulent possess the following traits: a) they are unusually large, b) they tend to be out of the norm for the customer, and c) they suddenly occur more frequently. An additional method of identifying fraudulent transactions for this algorithm is through comparison with other sources of information such as location, and devices used to verify the authenticity of the parties involved, (Arya, 2024). One of the advantages of this algorithm is that it tends to adapt and learn from new data in real-time, enhancing its effectiveness in detecting evolving fraud patterns. Additionally, given that it utilizes cognitive computations, this algorithm is classified as one of the most accurate. Finally, the algorithm can identify complex patterns in any given data set. The drawback of the above however is that, like most of the other algorithms previously mentioned, it over-fits data especially when complexity is involved. It is also expensive to implement and time-consuming, (Negi, 2021). It is worth noting that this algorithm is a form of deep learning.

2.2. Proposed model

In this paper, the proposed model is an ensemble model, a machine-learning technique that combines different individual algorithms to improve predictive performance. Four different types of ensemble techniques exist, but in this case, the utilized type will be the stacking technique, which combines predictions from multiple models of different types. Specifically, the combined model will be made up of machine learning and deep learning. The algorithm for machine learning will be random forest, while on the other hand, the keras classifier will be used in the case of deep learning. Since random forest falls under supervised machine learning, the data used will contain two main classifications, “fraudulent and non-fraudulent” transactions. The goal of the model will be to predict and separate fraudulent transactions from legitimate ones through a thorough analysis of the transactional data. Some of the unique features of the data will include transaction ID, date, amount, frequency, account balance, location, etc. The model will afterward be evaluated by considering factors such as accuracy, precision, recall, etc. In addition, this model's performance will be compared to those of the standalone random forest algorithm.

2.3. Gap in literature

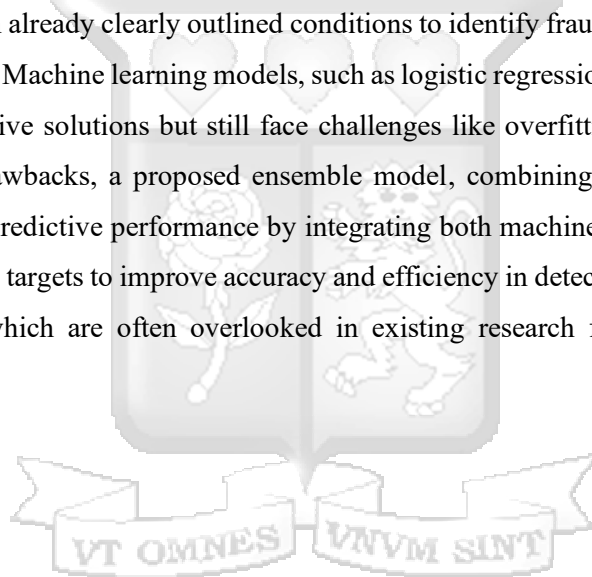
There is a notable absence of research channeled towards fraud detection models specifically for SACCOs. Existing literature mainly addresses fraud in larger financial institutions, such as banks, which have different operational frameworks and risk profiles compared to SACCOs. This discrepancy leaves a significant gap in the knowledge required to develop effective fraud detection and prevention mechanisms that cater to the unique needs of SACCOs, particularly in developing economies like Kenya.

Most existing studies focus on financial statement and credit card fraud, neglecting other critical forms of fraud that SACCOs frequently encounter, such as identity theft in transactions, online payments, etc. Transactional fraud poses a substantial risk to SACCO's operations and financial health, yet it remains underexplored in current research. This narrow focus on certain types of fraud in the literature fails to address the comprehensive fraud risk landscape SACCOs face, underscoring the need for more targeted studies in this area. Additionally, most fraud detection research is centered on European or other developed economies, with limited applicability to developing countries' economic and regulatory environments. This geographical bias further indicates the gap in the literature, as the fraud dynamics and institutional frameworks of SACCOs in countries like Kenya differ significantly from those in more developed regions.

2.4. Summary of the literature review

The literature review explores various theories explaining transactional fraud in financial institutions, focusing on human behavior and motivations behind fraudulent activities. The fraud triangle theory by Cressey identifies three key elements: pressure, opportunity, and rationalization, as essential for fraud occurrence. Wolfe and Hermanson's fraud diamond theory adds a fourth element, capability, emphasizing the need for skills to maximize the identified opportunities. Routine activity theory by Cohen and Felson on the other hand narrates the role of potential offenders, suitable targets, and the absence of capable guardians in crime occurrences. In summary, these theories provide a foundation for understanding and addressing transactional fraud which is through examining the behavioral patterns and environmental factors that facilitate fraudulent actions.

The review also analyzes existing fraud detection models, noting their limitations. The rule-based approach for instance, which relies on already clearly outlined conditions to identify fraud is time-consuming, costly, and prone to false positives. Machine learning models, such as logistic regression, decision trees, and neural networks, offer more adaptive solutions but still face challenges like overfitting and high computational costs. To address these drawbacks, a proposed ensemble model, combining random forest and a keras classifier aims to enhance predictive performance by integrating both machine learning and deep learning techniques. This new model targets to improve accuracy and efficiency in detecting fraudulent transactions, especially for SACCOs, which are often overlooked in existing research focused on larger financial institutions.



Chapter 3

3. Methodology

3.0. Introduction

This chapter outlines the methodology that is used in the study. In addition to that, the reasons that underpin the choice of the methodology are outlined. To clearly explain the methodology, the chapter is divided into seven sections. The first section explains the source and features of the data to be used, while the second section is about data processing and cleaning. The next section addresses the model type and its justification, after which another section explains how it will be created, trained, and tested. The architecture of the model is displayed in a different section, while another outlines how it will be optimized. Finally, the last section details how the proposed model will be evaluated using various valuation metrics.

3.1. Data collection

The dataset used to conduct this project is obtained from one of the leading and best-performing SACCOs in the Kenyan market. This is according to the regulatory authority SASRA. However, due to strict data privacy regulations, the name of the SACCO will not be disclosed at all in this paper. Given the high frequency of transactions that take place in regulated deposit-taking SACCOs, the data collected is for a duration of one month (roughly 30 days). The data used contains more than 900,000 rows and a total of 11 columns. The rows are mainly made up of different account holders, while the columns contain various features such as type of transaction, transaction amount, account balance before and after the transaction, date of the transaction as well as the time of the transaction, and whether the transaction has been classified as fraudulent or non-fraudulent.

3.2. Data preprocessing

This stage entails converting data into an understandable format to extract useful meaning and conclusions. There are different methods of data preprocessing. In this study, data cleaning will be employed. Data cleaning entails correcting or getting rid of inaccurate, corrupted, improperly formatted, duplicate, or incomplete data from a dataset. The cleaning process will involve the following:

- a) **Correcting missing rows in the data-** One way of addressing this error is by dropping the missing rows. An alternative method of handling this is imputation, which replaces missing values with substitute values usually obtained by calculating the mean, median, or mode using the remaining complete data values.
- b) **Checking for outliers-** Outliers are values that deviate from the normal range of the existing data set. These data values can either be too high or too low which may negatively affect the results of the data. In this case, calculating the z-score will be the method used to identify existing outliers in the collected data. This value identifies how many standard deviations a data point is from the mean

of the dataset, that is, either on the upper or lower side. For example, if the z-score value deviates further from zero, the higher the probability that it is an outlier. In cases where the outlier value deviation is very large, the data value will be dropped altogether.

$$Zscore = X - \frac{\mu}{\delta}$$

- c) **Data transformation-** The data will be converted into a format that is acceptable and applicable to the underlying system to be used. In this case, the format of the data should work well in both R and Python, which will be the main tools used in data analysis.
- d) **Data validation-** Finally, under data cleaning, the last thing to be considered is data validation. This will entail checking if the obtained data set meets the appropriate financial data requirements. Features such as quality, accuracy, and completeness of the data will be checked for.

3.3. Model type, technique, and justification

As previously highlighted, the model proposed in this study is an ensemble model. This model can be created using the same or a combination of different algorithms. In most cases, the random forest algorithm is usually used for single-algorithm ensemble models as it contains predictions from multiple decision trees. For combined algorithms, multiple different algorithms can be combined, for example, random forest can be used with K-Nearest neighbor, etc. The main advantage of the ensemble model is that the predictions made are usually better as they tend to address the shortcomings of individual algorithms. This is possible because the predictions from the multiple models are merged to make the final prediction, which is more precise and concrete, thus making it a better predictor. There are three main ensemble techniques. They include bagging, boosting, and stacking. Bagging is one of the ensemble techniques that utilizes only a single algorithm (Weak learner). This one tends to train multiple proportions of the training data using numerous instances of the same algorithm, (Soni, 2023). The results of the numerous instances of the base models are combined through aggregation/ voting to give the final predictions. There are numerous advantages of the bagging technique. Such include, (Cheruku, 2019):

- a) Reduces overfitting in the data through variance reduction.
- b) Improves accuracy in the predictions given that it combines different weak learners to form a strong one.
- c) Addresses the challenge of outliers.
- d) Helps reduce the challenges associated with an imbalanced data set.

Bagging can be diagrammatically represented as indicated below.

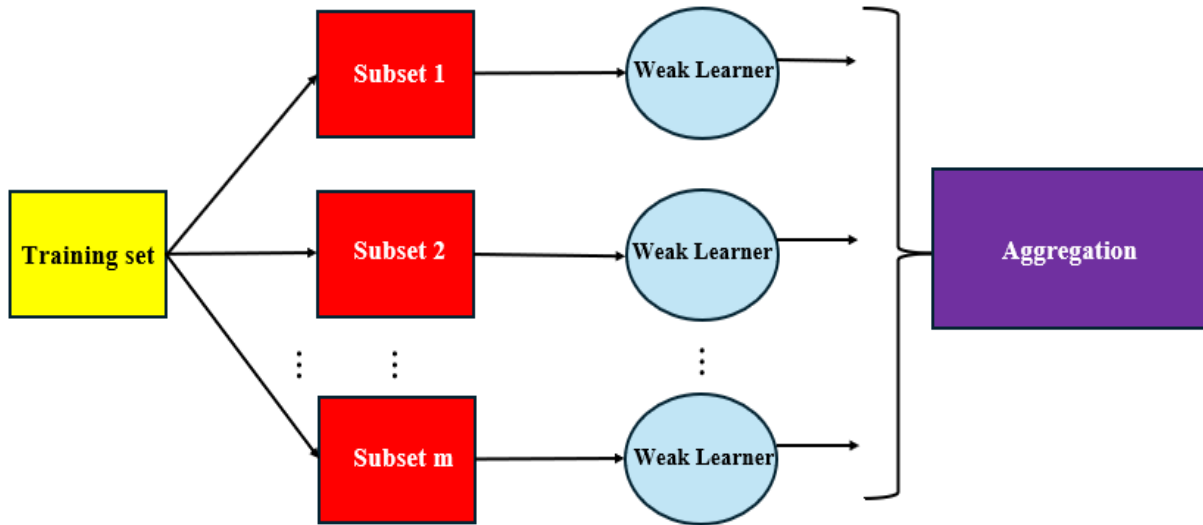


Figure 4: Bagging Technique
Source: (Cheruku, 2019)

Boosting technique on the other hand, unlike bagging which operates in parallel, tends to operate sequentially. The main idea behind the boosting technique is to train weak learners (multiple algorithms) sequentially so that each of the next learners addresses the drawbacks of the preceding learner, (Kalirane, 2024). Some of the advantages of the following model include the following:

- a) Reduces bias by thoroughly reweighting wrongly classified inputs.
- b) Produces strong predictors.
- c) Improves accuracy in the predictions by combining different accuracies from different models/ algorithms.

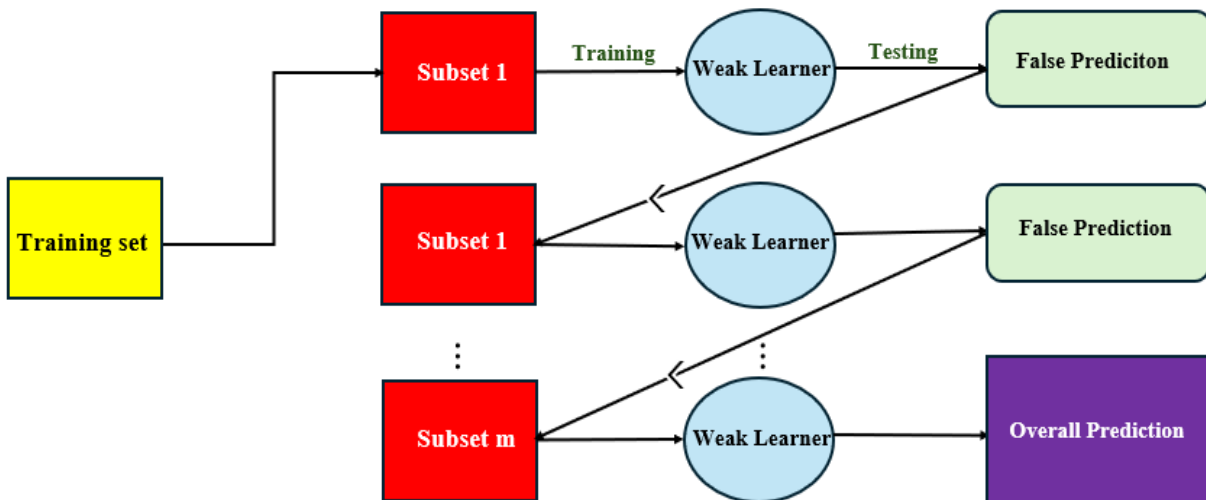


Figure 5: Boosting Technique
Source: (Kalirane, 2024)

The final ensemble technique is stacking. Unlike the previous two, this technique is more diversified as it combines multiple algorithms in the training data stage, whose predictions (Out of fold predictions- OOF-P) are then used as inputs for the meta-classifier/ Model. The predictions obtained from the meta-classifier are used as the final predictions. The algorithms used at the training stage are the base models. In this study, the stacking technique will be employed. The main reason why this technique is to be used in this study is because it allows for the combination of different algorithms to be employed at different stages, (Alexandropoulos et al., 2019) . Other advantages of this technique include the following:

- a) Allows for the identification of patterns in the data that cannot be easily identified by a single algorithm.
- b) Reduces the risk of overfitting.
- c) Improves accuracy level in the predictions.

The combined algorithms in this paper are random forest, a machine-learning algorithm, and keras classifier, from deep learning. Random forest is selected in the case of a machine-learning algorithm because of its ability to analyze large data sets with high dimensionality, though it takes time. In addition to that, given that it uses the predictions from different decision trees, it is less prone to overfitting compared to other machine-learning algorithms. Keras classifier on the other hand is selected and not any other classifier because it possesses deep learning algorithms that have high ability to identify complex patterns in data sets. These algorithms are also very flexible and, hence can be modified to suit different types of data sets. For some of the studies that have utilized random forest, it has been classified as one of the best machine-learning algorithms. According to (Islam et al., 2024), the table below summarizes the performance of the random forest algorithm against other algorithms.

No.	Model	Accuracy	Precision	Recall	F-Measure	AUC
1.	Random Forest	0.946	0.976	0.947	0.946	0.967
2.	Logistic Regression	0.934	0.975	0.936	0.935	0.957
3.	Decision Tree	0.936	0.967	0.938	0.938	0.919
4.	K-Nearest Neighbor	0.865	0.946	0.867	0.868	0.940

Table 1: Algorithms Performance Comparison 1

Source: (Islam et al., 2024)

The table indicates that the random forest algorithm performs better when predicting fraudulent transactions than other algorithms. Further to that, another study by, (Afriyie et al., 2023), on detecting and predicting credit card fraud for banks, while comparing the performance of different algorithms, again, random forest emerges as the best performer against other algorithms. This is depicted in the table below.

No.	Model	Accuracy	F1-Score	Recall	Precision	Specificity	AUC
1.	Decision tree	0.92	0.09	0.93	0.05	0.92	0.95
2.	Random Forest	0.96	0.17	0.97	0.09	0.96	0.99
3.	Logistics regression	0.92	0.08	0.76	0.04	0.92	0.88

Table 2: Algorithms Performance Comparison 2
Source: (Afriyie et al., 2023)

Additionally, as (Lokanan & Sharma, 2022) tried to predict financial fraud in financial markets, in the study, the conclusion was that random forest performed better than the rest of the algorithms used. The results of the performance of the different algorithms are depicted in the table below.

No.	Model	Sensitivity/ Recall	Specificity	Precision	F- Measure
1.	Logistics regression	0.90	0.80	0.80	0.84
2.	Decision tree	0.91	0.90	0.7	0.79
3.	Random Forest	0.98	0.90	0.99	0.99
4.	Cat-Boost	0.95	0.90	0.97	0.96

Table 3: Algorithms Performance Comparison 3
Source: (Lokanan & Sharma, 2022)

Combining these two algorithms offers several advantages, including the following:

- a) **Offers more stability-** Individual algorithms usually balance out variability in an ensemble model. Consequently, the predictions become more reliable and consistent.
- b) **Identifies more complex patterns-** Random Forest excels in capturing relationships between different features in the data set. Keras classifier- algorithms on the other hand thrive at identifying non-linear relationships. If the two are combined, wider ranges of data patterns are easily identified, which would not happen if an individual algorithm was used instead, (Islam et al., 2024).
- c) **Improves accuracy level -** The ensemble model combines the strengths of both random forest and keras classifier algorithms, leading to more accurate predictions than those of an individual algorithm. This will be well brought out when the results of the ensemble model are compared to those of an individual random forest algorithm, (Cheruku, 2019).
- d) **Reduces bias and overfitting-** Overfitting is reduced by combining the predictions from the different algorithms used as features for the new model, (Kalirane, 2024).

The following diagram represents an overview of the ensemble model used in this study.

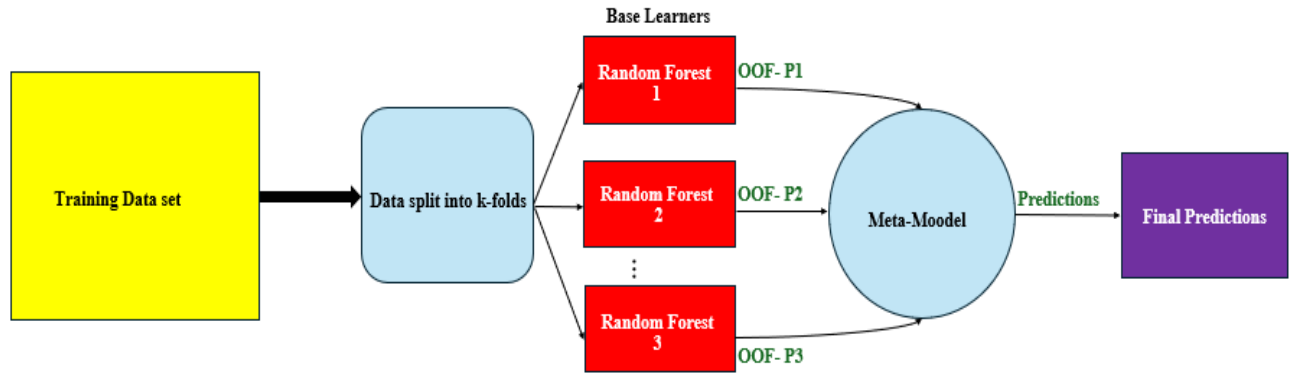


Figure 6: Ensemble Model
Source:(Alexandropoulos et al., 2019)

3.4. Model creation, training, and testing

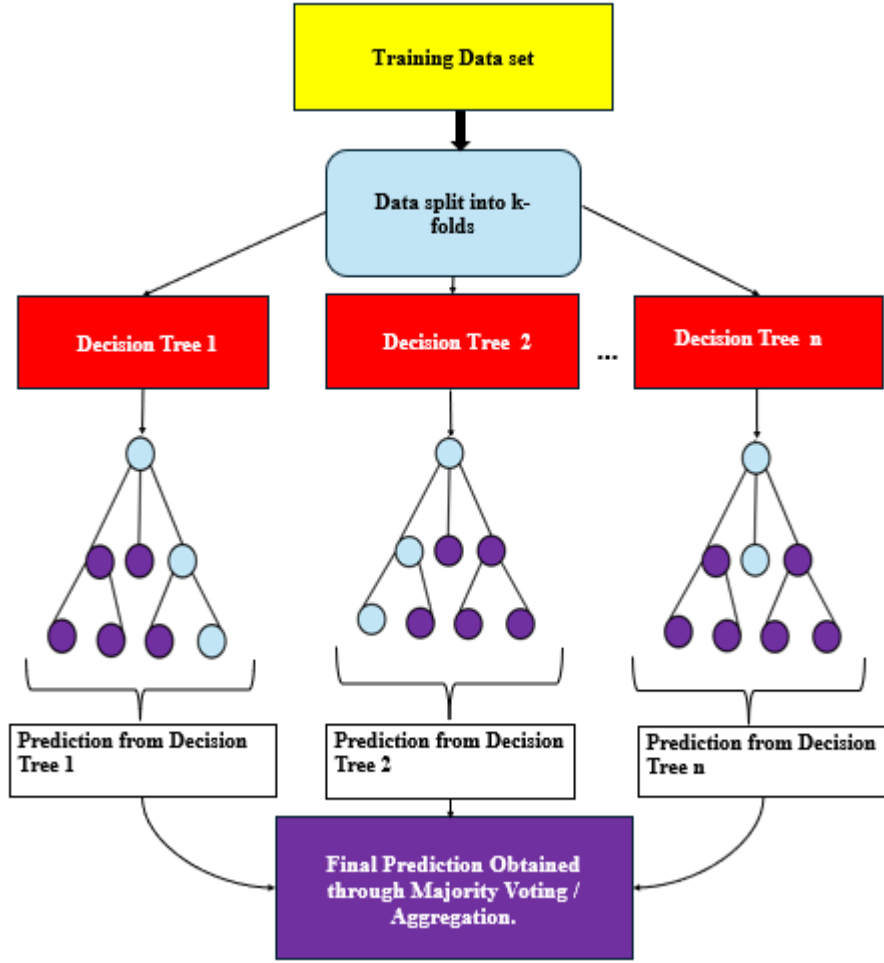
To create the model, the data set will be divided into three, training data, validation data, and testing data. Specifically, seventy percent (70%) of the total data will be used for training the model while thirty percent (30%) will be used as testing data. The training data will further be split into seen-training data (70%) and the remaining (30%) unseen-training data, also known as validation data. The purpose of the validation data is to assess the model's performance with unseen data during the training process. In other words, the validation data will help determine how well the model identifies patterns when fed with unseen data. The two algorithms, random forest and keras classifier will then be trained independently using both training data and validation data. The predictions from the random forest (base model/ learner) will be used as input features for the Keras classifier. Equally, predictions made by the deep learning algorithm will then be used as the final predictions using the testing data set which at this point will not be labeled.

3.5. Model architecture

3.5.1. Model architecture of random forest

A random forest algorithm is an ensemble model given that it consists of multiple decision trees that aid in making predictions or forecasts. The training data is split into numerous sub-sets through a process called bootstrapping, where random samples are drawn with replacement, and then fed to the decision trees. Each decision tree checks for a particular feature within that data set instead of all features, (Salman et al., 2024). The advantage of this is that overfitting is reduced, and the models' generalization capability is improved. Once the predictions are obtained from each decision tree, they are aggregated after which the final single output is obtained. It is important to note that, in the case of classification, majority voting is used to determine the final output, while in the case of regression tasks, aggregation is used instead. By averaging the results, the models' robustness is improved, and the variance is reduced, (Han et al., 2020). The same

rationale will be applied in the proposed ensemble model that is made up of three random forest algorithms as part of the base model. The random forest algorithm is explicitly represented as follows:



*Figure 7: Random Forest- Model Architecture
Source: (Wei, 2024)*

3.5.2. Model architecture of multilayer perceptron

Having discussed the machine learning algorithm, this section will focus on the deep learning aspect of the ensemble model. The Keras classifier is a high-level application programming interface (API) from the Keras library that is used to build and train neural networks for classification tests. Therefore, in this paper, the Keras library will be used to implement the model.

The specific classifier used under deep learning, and imported from Keras classifier is the multilayer perceptron (MLP). This is a type of artificial neural network consisting of multiple layers known as dense layers. Each layer comprises neurons that use non-linear activation functions that allow the network to learn easily any complex pattern within the data, (Rashedi et al., 2024). These dense layers are further divided into three main groups:

- a) **Input layer**- This is the first layer of the neural network that receives the input data, which is passed on to the next layer through the neurons that connect them. Each neuron represents a feature of the input data, (Jaiswal, 2024). In our model, the input data used are the predictions obtained from the base model which is made up of three random forest algorithms/ models. The input layer does not contain an activation function.
- b) **Hidden layers**- This is the middle layer that contains interconnected neurons which perform necessary computations on the data received from the input layer. The weighted sum of each input is computed first after which it is passed through an activation function. The equation of the weighted sum is given as follows, (Dagal, 2024):

$$\text{Weighted Sum} = \sum_{i=1}^n (W_i * X_i) + b$$

where:

W_i – denotes the corresponding weight of the i-th input.

X_i – denotes the i-th input.

b – denotes the bias associated with each hidden layer.

Different activation functions are usually used depending on the task to be performed by the network, as well as its properties. The purpose of the activation functions is to introduce nonlinearity in the network which would then allow it to learn complex patterns or relationships in the data. For this paper, the Rectified Linear Unit (ReLU) activation function is used in the case of all the hidden layers.

ReLU is one of the most widely used activation functions due to its high level of simplicity and effectiveness. Equally, it is preferred to other functions because it is easier to train and better still, outperforms them. This function tends to output the input directly, only if it is positive, and if not, the output becomes zero. In other words, if the input (X) is greater than zero, then the output from the function will be the same as the input. However, if the input (X) is less than zero, then the output from the function will be zero. ReLU function is usually expressed as follows, (Brownlee, 2019):

$$f(X_i) = \text{Max} (0, X_i)$$

where:

X_i – denotes the i-th input.

The equation can also be expressed explicitly as follows:

$$f(Xi) = \begin{cases} Xi, & \text{if } X > 0 \\ 0, & \text{if } X \leq 0 \end{cases}$$

where:

$f(Xi)$ – returns Xi if it is greater than 0.

$f(Xi)$ – returns 0 if Xi is less than or equal to 0.

The graph of the ReLU activation function further depicts how the function works.

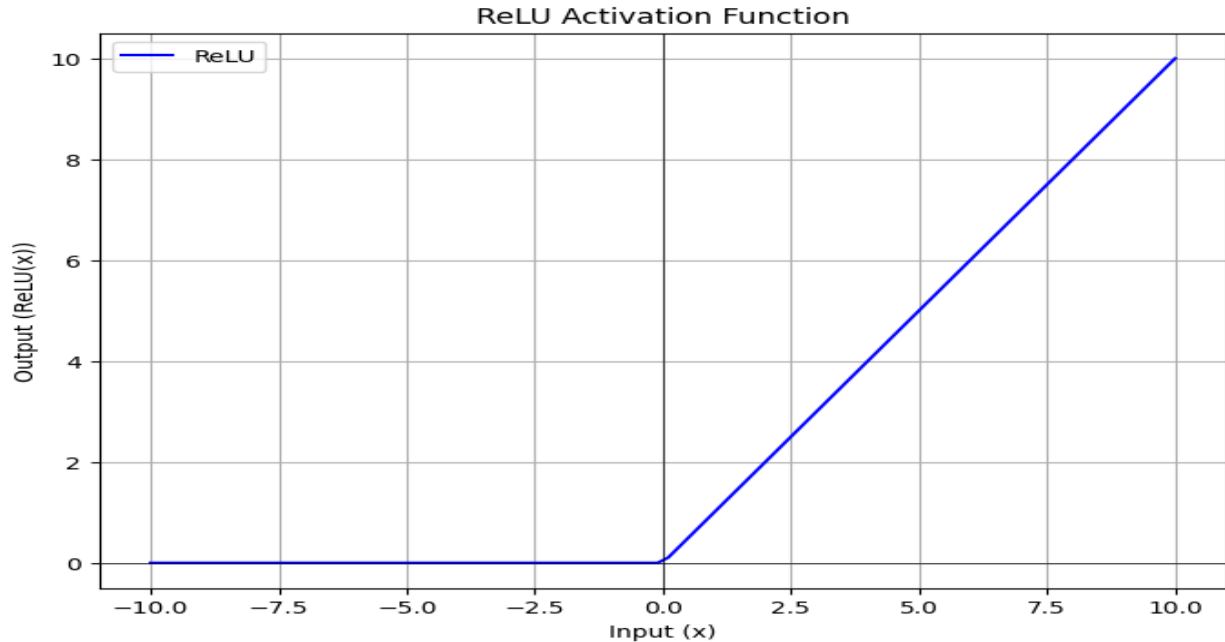


Figure 7: ReLU Activation Function Graph
Source: (Brownlee, 2019)

- c) **Output layer-** This is the last layer of the MLP. It produces the final output of the network. As is the case with the hidden layers, the number of neurons in the output layer depends on the task at hand. For example, in a binary classification problem, there would be a single output neuron that represents the probability of belonging to one class or the other. Given that in this paper we are dealing with a binary classification case, the output layer will also contain one neuron. Activation functions are sometimes found in the output layer. In this paper, the sigmoid function is employed. The output of a sigmoid activation function always lies between 0 and 1, which makes it suitable for binary classification tasks, (Patel, 2022). Now, with the output ranging between 0 and 1, the model can easily analyze it and determine if it leans more toward the positive (1) or negative (0) side, after which it can classify it as either fraudulent or non-fraudulent. The sigmoid function equation and graph are explicitly depicted below.

$$f(Xi) = \frac{1}{1 + e^{-Xi}}$$

where:

Xi – denotes the i-th input.

The graph is as follows:

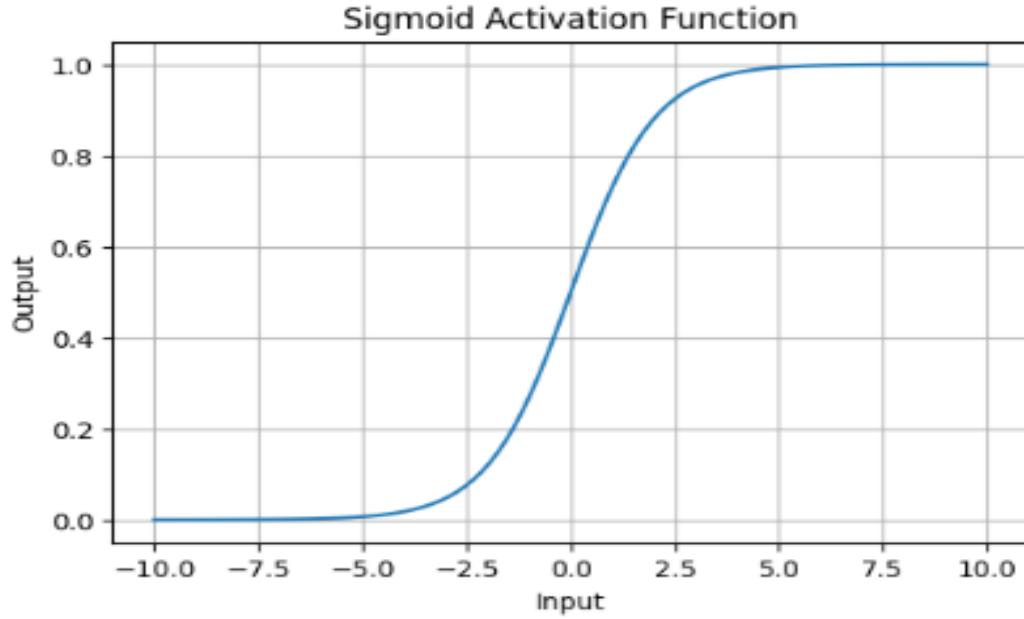


Figure 8: Sigmoid Activation Function Graph
Source: (Patel, 2022)

MLP is also considered to be a feed-forward neural network because, when making predictions, information flows in one direction, from input to output. Of the different types of neural networks, it is selected to be used in this paper because of its ability to learn non-linear relationships in data sets, making it suitable for classifications, regression, and pattern recognition tasks. MLP neural network is diagrammatically presented below.

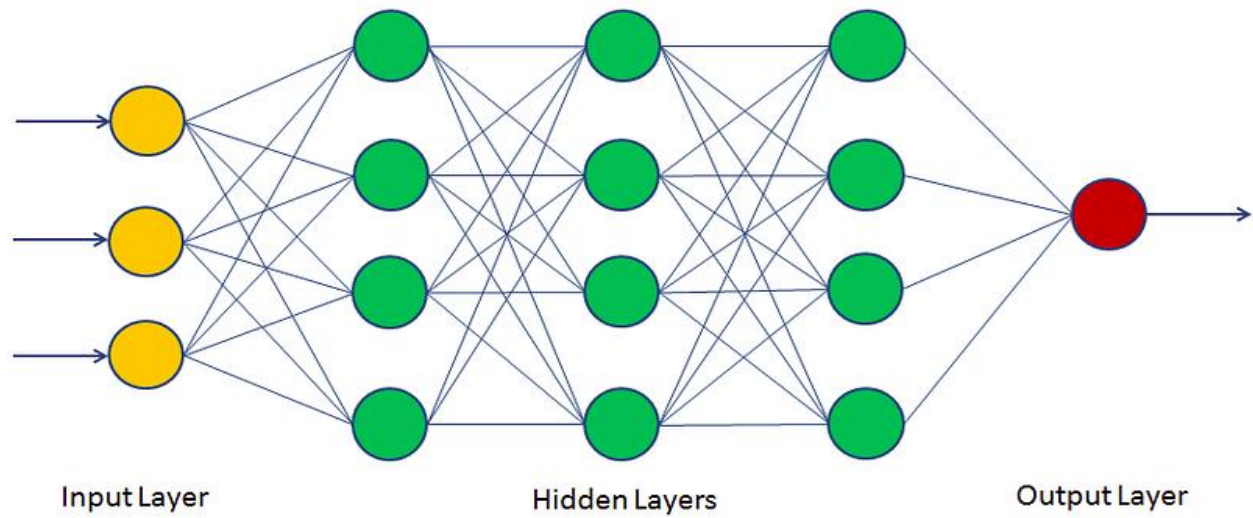
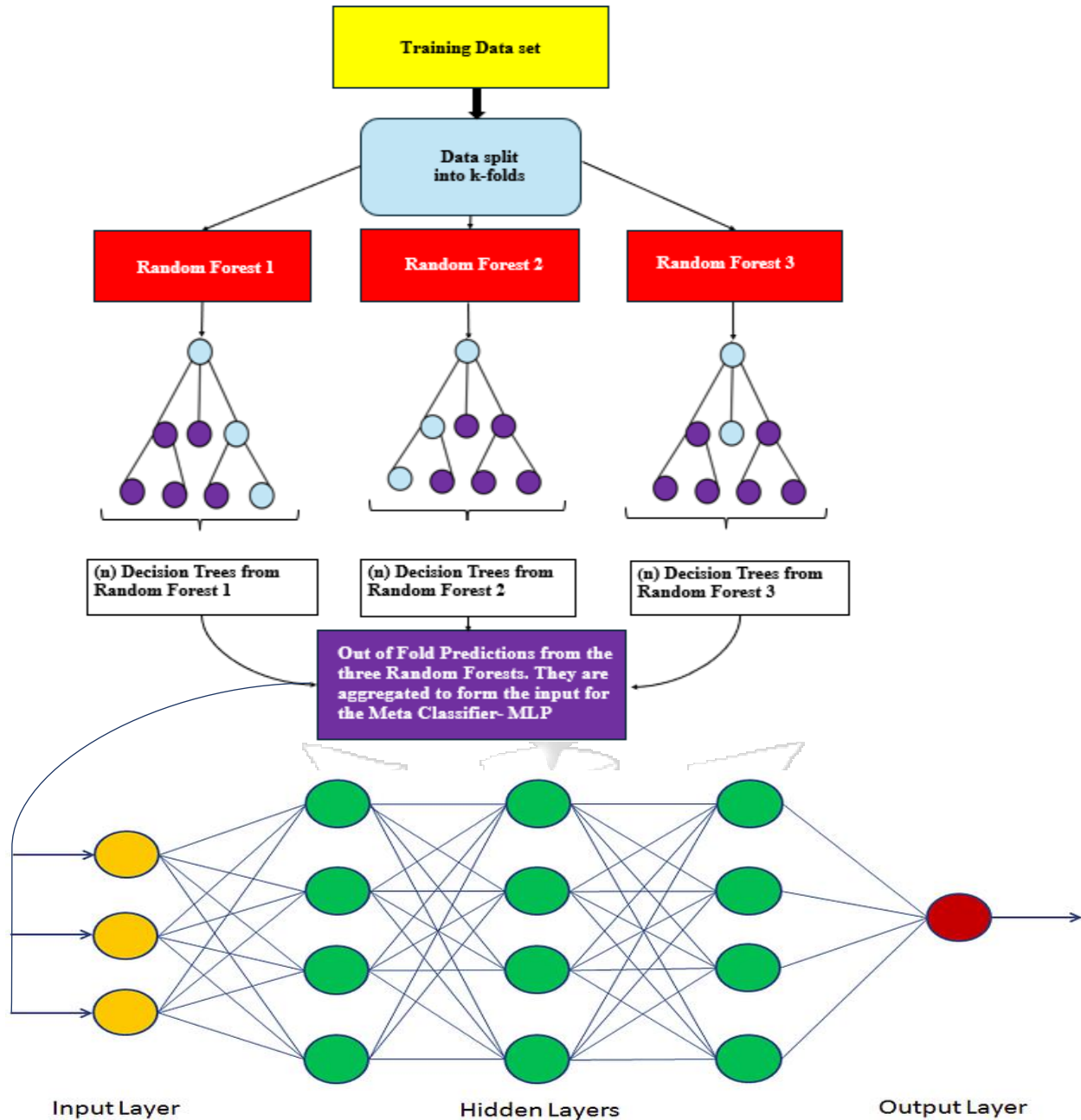


Figure 9: Multilayer Perceptron- Model Architecture
Source: (Multi-Layer Perceptron Neural Network Using Python – Machine Learning Geek, 2021)



3.5.3. Proposed model architecture

The ensemble model architecture is denoted below. It indicates that the predictions from the base model (3 random forest algorithms) will be used as inputs for the meta-model (MLP).



*Figure 10: Proposed Ensemble Model- Model Architecture
Source: Created Model*

3.5.4. Model compilation of multilayer perceptron

Model compilation outlines how the MLP is trained by specifying the essential components. In this paper, three components are employed, the optimizer, the loss function, and the metrics used to evaluate the model.

- a) **Optimizer-** An optimizer specifies how the models' weights are updated during the training stage. To adjust the weights, backward propagation always comes into play. The goal of backward propagation is to ensure that optimal parameters such as weight and bias are obtained, yielding more accurate results. In particular, the adaptive moment estimation (Adam) optimizer is selected because it is more robust across problems and converges faster compared to other optimizers. Furthermore, it is considered more efficient and adaptive, making it even more preferable. This is how the optimizer works, (Jamhuri, 2023).

Adam uses two main concepts:

1. First Moment Estimate (mt):

- This is the running average of the gradients (like momentum).

2. Second Moment Estimate (vt):

- This is the running average of the squared gradients (like RMSprop).

At each step:

- a) Compute the gradient of the loss function ($\nabla L(\theta_t)$).
- b) Update mt and vt using exponential decay:

$$mt = \beta_1 * mt - 1 + (1 - \beta_1) * gt$$

$$vt = \beta_2 * vt - 1 + (1 - \beta_2) * gt^2$$

Where:

β_1 : Decay rate for the moving average of gradients (e.g., 0.9).

β_2 : Decay rate for the moving average of squared gradients (e.g., 0.999).

η : Learning rate.

ϵ : Small constant to prevent division by zero.

- c) Apply Bias Correction:

$$mt^{\wedge} = \frac{mt}{(1 - \beta_1^t)}$$

$$vt^{\wedge} = \frac{vt}{(1 - \beta_2^t)}$$

d) Update Parameters:

$$\theta_t = \theta_{t-1} - \eta \frac{mt^{\wedge}}{(\sqrt{vt^{\wedge}} + \epsilon)}$$

b) **Loss function-** A loss function, sometimes referred to as an error function quantifies the difference between the predicted outputs and the actual target values. The loss obtained reflects the models' accuracy level of the predictions. A loss function is usually selected based on the task to be done. In this paper, the binary cross-entropy (BCE) loss function is used. As the name suggests, it is a binary classification loss function that works well with models that yield outputs with a probability between zero and one, making it ideal for this paper. As the gap between the actual and predicted value increases, the cross-entropy loss increases. The equation below depicts how the loss function works, (Sorokin, 2024).

$$Loss = -\frac{1}{N} \sum_{i=1}^N [y_{true} * \log(y_{pred}) + (1 - y_{true}) \log(1 - y_{pred})]$$

Where:

N: Total number of samples in the dataset.

y_{true} : This is the actual label for each sample. In this paper, it would be either 0 or 1.

y_{pred} : This is the predicted probability that the sample belongs to the positive class (class 1). It is the output of the model.

Log: This is the natural logarithm function.

The formula first calculates the loss/ error for each sample using the first section of the formula or equation above $[y_{true} * \log(y_{pred}) + (1 - y_{true}) * \log(1 - y_{pred})]$ and then averages these losses over all samples.

The term $y_{true} * \log(y_{pred})$ calculates the loss for the positive class (that is when the actual label y_{true} is 1). If the model predicts a high probability for the positive class (meaning y_{pred} is close to 1), this term will be close to 0. But if the model predicts a low probability (meaning y_{pred} is close to 0), this term will be a large positive number, contributing to a larger overall loss.

The term $(1 - y_{true}) * \log(1 - y_{pred})$ on the other hand calculates the loss for the negative class (that is when the actual label y_{true} is 0). If the model predicts a low probability for the positive class (meaning it predicts a high probability for the negative class, so y_{pred} is close to 0), this term will be close to 0. But if the model predicts a high probability for the positive class (meaning y_{pred} is close to 1), this term will be a large positive number, contributing to a larger overall loss.

- c) **Metrics-** These are used to measure the performance of the model during training and testing. They are similar to those highlighted in the previous chapter. They are accuracy, precision, recall, F-1 score, and AUC.

3.6. Model optimization

This is a process that aims to improve the performance and efficiency of a model either in machine learning or deep learning. Both the stand-alone random forest algorithm and the ensemble model are optimized. However, it is important to note that the optimized random forest algorithm is not the same as the one used in the proposed ensemble model. Instead, this is the random forest algorithm whose performance will be compared to the ensemble model's to determine which model performs better- the proposed model or the single random forest model. Given that the ensemble model is optimized, it would only be fair to compare the performance to that of an optimized random forest model.

3.6.1. Random forest optimization

A hyperparameter optimization using Bayesian optimization is carried out to improve the random forest classifier model performance. Bayesian Optimization (BO) is a probabilistic model-based optimization technique that is particularly useful when optimizing expensive or black-box functions by exploring the hyperparameter space more efficiently than traditional methods like grid or random search. Instead of evaluating the objective function exhaustively, BO builds a model of the objective function (usually a Gaussian Process, GP) and uses this model to select the most promising hyperparameters to evaluate next, (Popescu et al., 2009)

In other words, it uses prior knowledge about the function and the observations gathered during optimization to predict the regions of the hyperparameter space that are likely to yield the best results. The GPsampler in Optuna leverages a Gaussian Process model to perform an efficient hyperparameter search by balancing exploration and exploitation.

The optimization process in Optuna works as follows:

- a) **Define the Objective Function** - The first step is to define a function that represents the model's performance as a function of its hyperparameters. This function is usually a machine learning model, such as a random forest classifier.
- b) **Hyperparameter Search Space** - In this step, we define the range of hyperparameters to explore/tune. These are
 - **N-estimators** - The number of trees in the forest model.
 - **Max-depth** - The maximum depth of the trees
 - **Min-samples-split** - The minimum number of samples required to split an internal node
 - **Min-samples-leaf** - The minimum number of samples required to be at a leaf node
 - **Class weight** - Adjusts weights for the classes, with options including balanced and balanced subsample (useful for imbalanced datasets).
- c) **Optimization Process** - Optuna's GPSampler will perform optimization by modeling the objective function using a GP, and iteratively selecting new hyperparameter configurations to evaluate based on the current best observations.
- d) **Acquisition Function** - The optimizer selects the next set of hyperparameters to test based on the acquisition function, which is typically Expected Improvement (EI). The acquisition function determines how likely a new set of hyperparameters will improve the performance, considering the uncertainty of the GP.

3.6.2. Multilayer perceptron optimization

In addition to the previously highlighted hyperparameters, to optimize the model and prevent overfitting, dropout layers, and batch normalization techniques have been incorporated into the architecture. These techniques help to regularize the model and improve its generalization capabilities:

- a) **Batch normalization** helps by normalizing the activations of previous layers, ensuring that the network trains faster and is less sensitive to initializations.
- b) **Dropout layers** with varying rates (0.4, 0.3, and 0.2), randomly drop a fraction of the neurons during training, reducing the model's reliance on specific neurons and improving its ability to generalize to new data.

3.7. Model evaluation

To evaluate the model, the new model, whose features are those predictions obtained from the two algorithms, will be fed with the non-labeled testing data set, for it to make predictions. The predictions of the new model will then be compared to the actual labeled testing data set. To assess the performance of the ensemble model, different metrics will be used. However, before delving into the specific metrics, it's important to understand the following parameters as they play a role in the definition of the metrics.

- a) **True Positive (TP)**- This denotes the number of transactions that are fraudulent in the data set and are identified as fraudulent by the model. In this study, if a transaction is fraudulent, it is represented by a positive (1), and if it is not fraudulent, it is denoted by a negative (0).
- b) **True Negative (TN)**- This denotes the number of transactions that are non-fraudulent in the data set and are classified as non-fraudulent by the model.
- c) **False Positive (FP)**- This denotes the number of transactions that are non-fraudulent in the data set but are classified as fraudulent by the model.
- d) **False Negative (FN)**- This denotes the number of transactions that are fraudulent in the data set but are classified as non-fraudulent by the model.

Having defined the above parameters, the evaluation metrics will now be explored:

- a) **Accuracy**- Accuracy refers to the ability of the model to distinguish correctly between fraudulent and non-fraudulent transactions, (Riskiyadi, 2024). This is measured by dividing all the correctly predicted transactions by all the transactions. The higher the accuracy level, the better the model performance. The lower the accuracy value, the lower the model performance. This is represented by:

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

- b) **Precision**- It measures the proportion of correctly predicted positive cases (**TP**) out of all the cases the model predicted as positive (**TP and FP**) (Sai et al., 2023). It simply tells you how accurate the model is when it comes to identifying fraudulent transactions. A high value indicates that the model is doing well while a lower value indicates that the model is not performing well. The ratio is expressed as follows:

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

- c) **Recall/ Sensitivity**- This is also called the true positive rate. This rate is calculated as the fraction/ proportion of the true positive predictions over all the positive cases that were retrieved for testing. In other words, it assesses the models' ability to detect all instances of fraudulent transactions, (Khan et al., 2022). It is expressed as follows:

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

- d) **Specificity**- It is also known as the True Negative Rate. It simply depicts how good the model is at identifying non-fraudulent transactions. It is the proportion of the TN cases divided by the summation of TN and FP, (Khan et al., 2022). It calculates the true negative rate.

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

- e) **F-1 score or F-Measure**- This is the weighted average of the precision and recall results, (Mugoshi, 2021). It measures the accuracy of a model on the data set and is specifically useful in binary classification systems. It is usually expressed as follows:

$$F_Measure = \frac{2(Recall \times Precision)}{Recall + Precision}$$

- f) **Area Under the Receiver Operating Characteristic curve (AU-ROC-C)** - This demonstrates how well the model works, and how well it can distinguish between fraudulent and non-fraudulent transactions, (Sai et al., 2023). It is a more reliable performance metric because it considers both type I and type II errors. In addition to that, it accounts for imbalanced data sets and plots the trade-off between the FP rate and the TP rates for different choices of binary classification models, (Lokanan & Sharma, 2022). AU-ROC-C requires class probabilities to evaluate the model's performance across different decision thresholds. These probabilities indicate how confident the model is about its predictions. By using probabilities, the AU-ROC-C graph allows us to assess the model's ability to rank transactions from most to least likely to be fraudulent, rather than just making a hard decision at a fixed threshold. This helps in understanding the model's ability to distinguish between non-fraudulent and fraudulent transactions at various levels of confidence, leading to a more comprehensive evaluation. The higher the area under the curve, the better the performance of the model.

- g) **Confusion matrix**- This summarizes the performance of the model in a matrix form. In this study, the confusion matrix will be a 2 by 2 matrix given that it's a binary classification model, and will contain the TPs, TNs, FPs, and FNs. The confusion matrix is diagrammatically represented as follows:

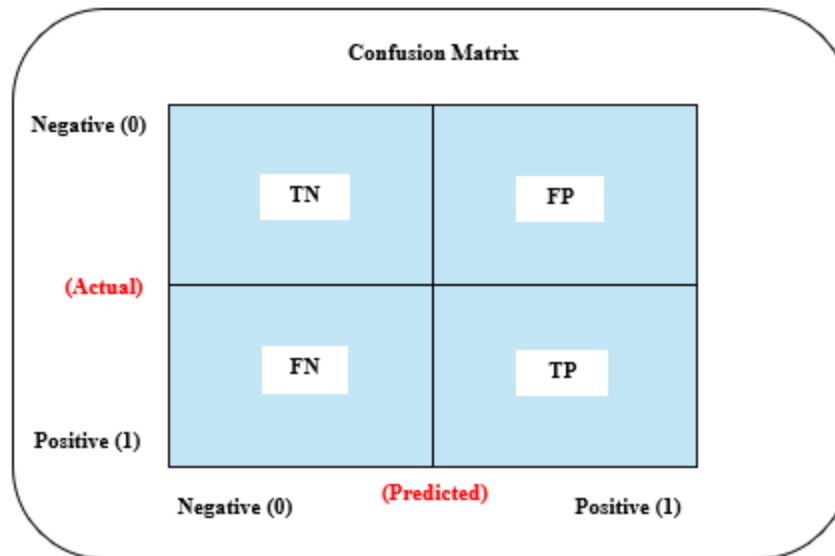
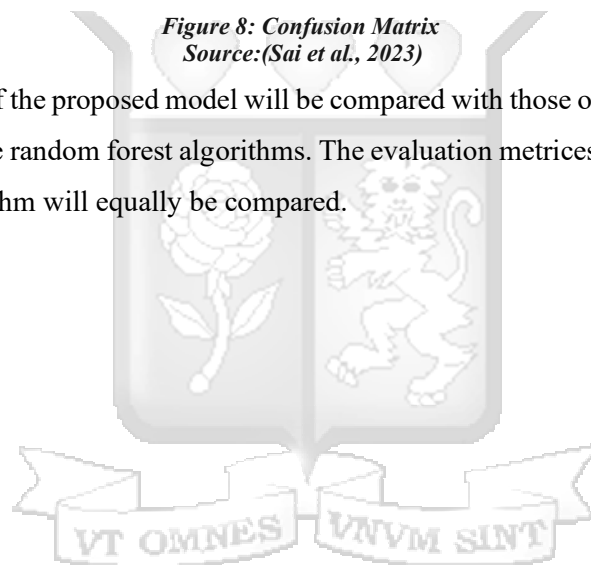


Figure 8: Confusion Matrix
Source: (Sai et al., 2023)

Further to that, the results of the proposed model will be compared with those obtained from both optimized and unoptimized standalone random forest algorithms. The evaluation metrics of the new model and those of the random forest algorithm will equally be compared.



Chapter 4

4. Data analysis and results

4.0. Introduction

This chapter presents the findings of the paper obtained after implementing the proposed model. It is categorized into five main sections. The first section covers the main steps involved in data pre-processing, followed by a section on model training, and another on how the model was optimized. The final section is on results and analysis where the previously mentioned valuation metrics were used on both models to evaluate their performances.

4.1. Data preprocessing

This stage entailed converting data into an understandable format to extract useful meaning and conclusions. Of the many preprocessing methods, data cleaning was employed. The cleaning process involved the following:

- a) **Correcting missing rows in the data**- There were no missing rows in the data.
- b) **Checking for outliers**- Equally, there were no outliers identified in the data set used.
- c) **Data validation**- Under this, the quality of the data was checked. A key problem identified was the data being highly imbalanced meaning that the non-fraudulent (**Majority class- 915,728 units**) transactions outnumbered the fraudulent transactions (**Minority class- 8,213 units**) by far.

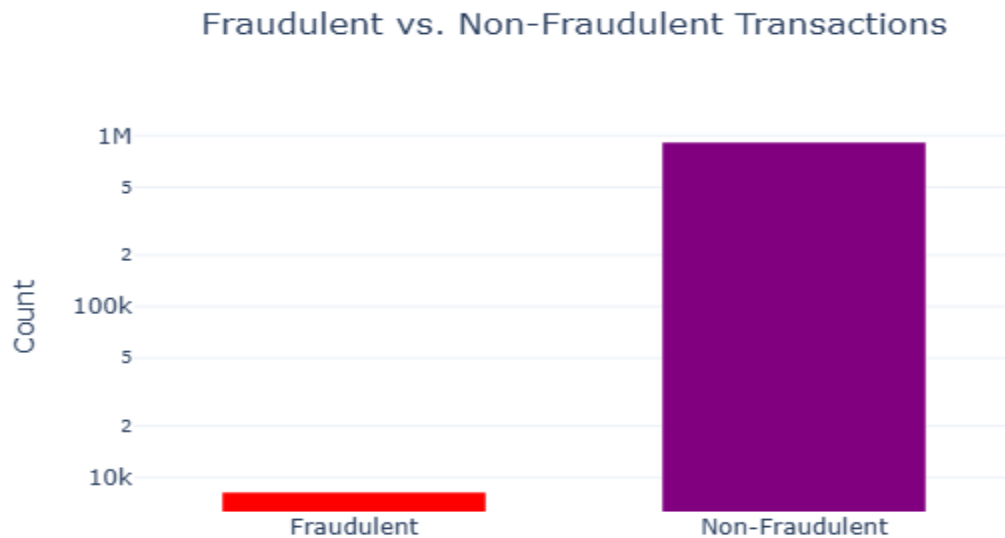


Figure 9: Fraudulent vs Non-Fraudulent Transactions
Source: Created Model

To address the problem, a resampling technique was used. In particular, the Synthetic minority oversampling technique (SMOTE) was implemented. However, it is important to note that the technique was deployed only after the data was split into training and testing datasets. This was done to prevent data leakage. Data leakage normally occurs when information from the testing dataset is used during training, which would then result in an overly optimistic model performance that would not be realistic. SMOTE works by generating new samples for the minority class through creating synthetic data points between existing cases that are similar in feature space. It starts by randomly selecting a sample from the minority class and finding a nearby neighbor. A new data point is then created by combining these two samples, positioning it somewhere between them in the feature space. SMOTE helps create a balance between the two classes without necessarily duplicating existing data sets, (Fernandez et al., 2018). After applying SMOTE, the balanced dataset was as follows:

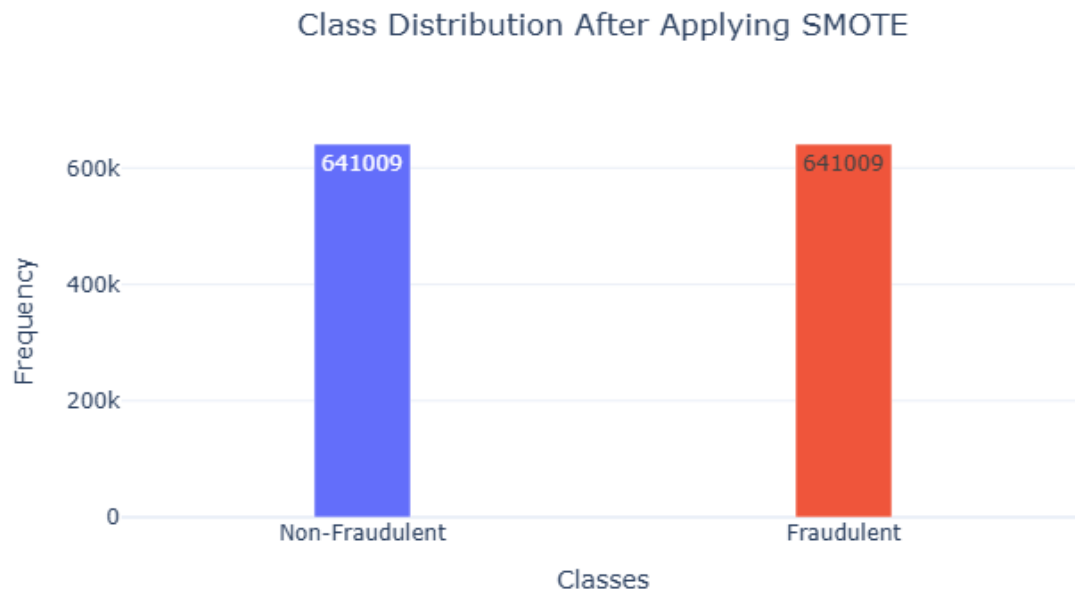


Figure 10: Fraudulent vs Non-Fraudulent Transactions after applying SMOTE
Source: Created Model

4.2. Model training- proposed model

At this stage, the proposed ensemble model was fed with the training data set from which it would learn any form of pattern. Given that the model is composed of different algorithms, they are discussed separately as follows.

4.2.1. Random forest

The base model, which forms the first part of the ensemble model is composed of three random forest algorithms, with each containing multiple decision trees. The training data set was split into k-folds (3), after which each unique fold was used as input data for the three random forest algorithms under the base model. The base model was run and the prediction from each random forest were obtained. These predictions are often referred to as out-of-fold predictions (OOF). The accuracy level obtained from each of the random forest algorithms was as follows:

Folds	Algorithm	Accuracy
1.	Random Forest	0.9710
2.	Random Forest	0.9703
3.	Random Forest	0.9697

*Table 4: Base- Model Accuracies
Source: Created Model*

4.2.2. Multilayer perceptron

MLP is the meta-model of the proposed ensemble model. Predictions obtained from the three random forest algorithms were aggregated and used as input data for the MLP. The meta-model was to further analyze the data and make its predictions. Like other neural networks, hyperparameters were tuned in to achieve optimal performance. These are parameters that control the training process of a model. In particular, epochs and batch size hyperparameters were used. Epoch defines the number of times the training data set is to be worked out through the learning algorithm. One epoch means the data is passed only once through the learning algorithm. For the epoch to be complete, it has to undergo both a front and backward pass. Batch size on the other hand represents the number of samples used in one front and backward pass per batch for each epoch. These two hyperparameters were used because they directly affect the efficiency and accuracy level of the training process. Additionally, the two would help in reducing the loss function error, thus improving the level of accuracy. Equally, they would help in determining the optimal weights as well as reducing bias. The number of batches used is usually computed as follows:

$$\text{Total number of batches} = \frac{\text{Total Training Samples}}{\text{Batch Size}}$$

In this paper, the number of epochs used was 5 and each epoch contained the following number of batches, with the last batch containing slightly fewer total samples (7), while the rest contained a total of (64) each.

$$\text{Total number of batches} = \frac{14,023}{64} = 219$$

MLP was also trained using feedforward propagation and backward propagation. For feedforward propagation, data flows in one direction, from the input layer to the output layer without forming any feedback loops. Backward propagation allows for feedback loops to be created. Usually, in the case of backward propagation, once the prediction is made, the error is calculated by the loss function (binary cross-entropy loss)- (the difference between the predicted output and the actual output). If it is too high, it is propagated back through the network, after which the models' weights are adjusted to minimize the error. The predictions obtained from the MLP were used as the final predictions from the ensemble model. To complete the training exercise, the model was run again, but this time using the validation data. The difference between the accuracy level of the model when using only training data and validation data is graphically represented below.

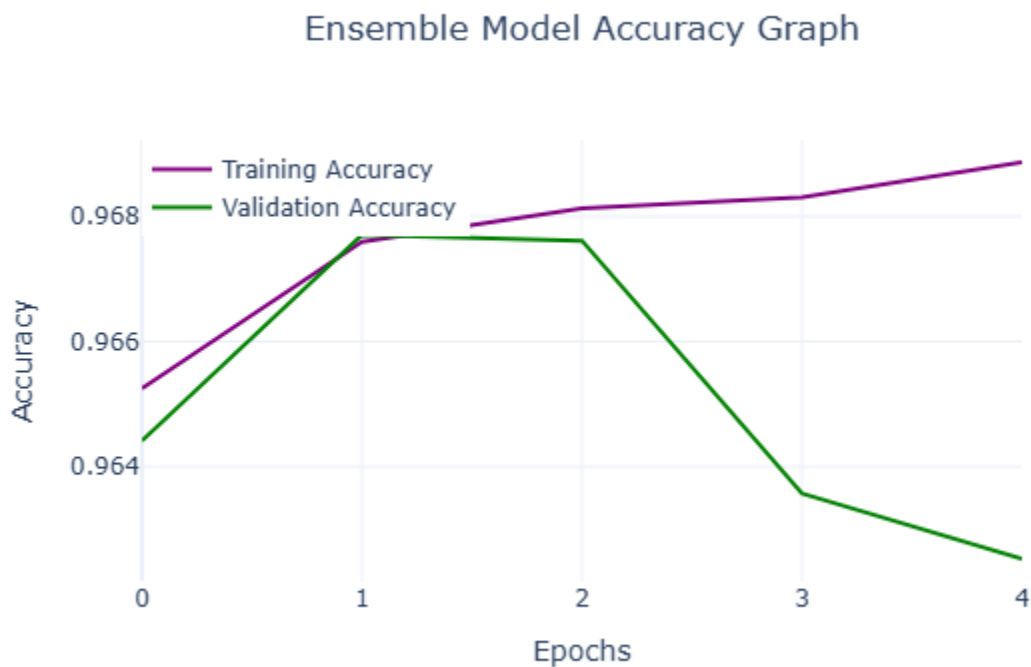


Figure 11: Ensemble Model Training Vs Validation Accuracies
Source: Created Model

Overall, the model performed consistently well on both the training and validation datasets at + 96.2%, showing no significant overfitting or underfitting issues. This was further supported by the loss function graph which indicated that overall, the training loss continued to improve, while the validation loss fluctuated slightly, suggesting the model was learning well but was still sensitive to variations in validation data. The model did not show significant signs of overfitting, as the validation loss did not drastically increase, and it stabilized by the end of the training, demonstrating that the model generalized effectively.

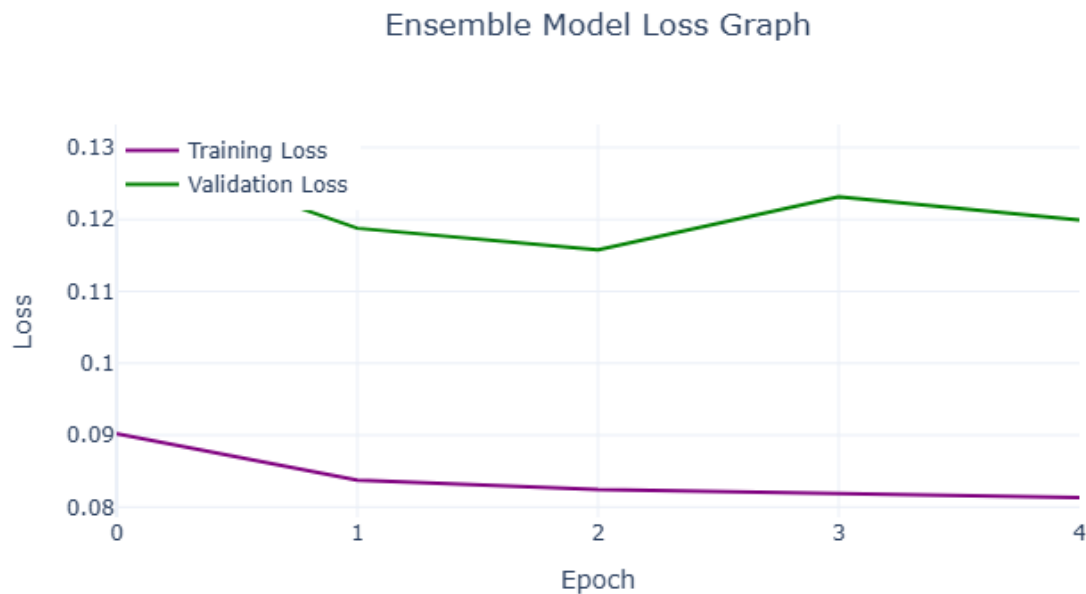


Figure 12: Ensemble Model Training Vs Validation Loss Function
Source: Created Model

4.3. Model optimization

4.3.1. Random forest optimization

Bayesian optimization was carried out to improve the random forest classifier model performance. Based on the level of importance, the optimization hyperparameters used can be represented in a graph as indicated below. From the graph, the number of estimators n depicted the highest level of importance, hence, it should be prioritized.

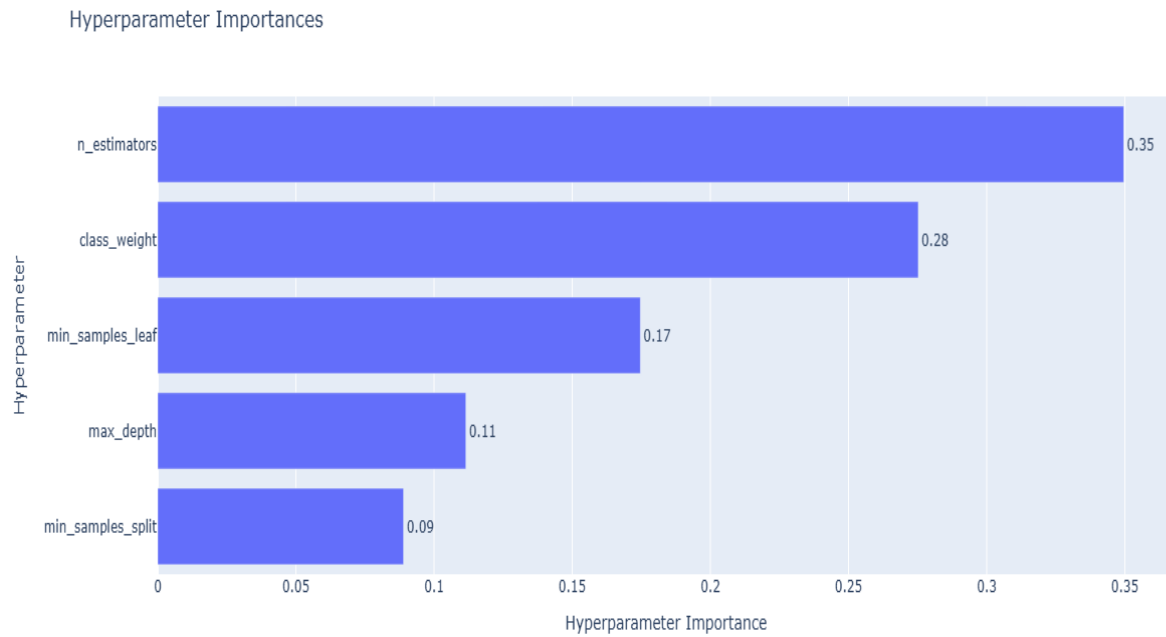


Figure 13: Hyperparameters Importance- Random Forest
Source: Created Model

4.3.2. Multilayer perceptron optimization

To prevent overfitting, dropout layers, and batch normalization techniques were incorporated into the architecture. These techniques helped to regularize the model and improve its generalization capabilities:

4.4. Model evaluation

For this section, the performance of the proposed model was evaluated based on the valuation metrics presented on the third chapter. In addition to that, the optimized random forest classifier model performance was compared to that of the ensemble model, to determine the best-performing model.

4.4.1. Valuation metrics

The valuation metrics for the ensemble model are represented by the following classification report.

Classification report for the ensemble model				
Class	Accuracy	Precision	Recall	F1-Score
0	0.9977	0.9997	0.9980	0.9988
1		0.8146	0.9627	0.8824

Table 5: Ensemble model- Evaluation Metrics
Source: Created Model

Next, the valuation metrics for the optimized random forest are presented below.

Classification report for the optimized random forest model				
Class	Accuracy	Precision	Recall	F1-Score
0	0.9437	0.9998	0.9434	0.9708
1		0.1339	0.9744	0.2354

Table 6: Optimized random forest model- Evaluation metrics
Source: Created Model

Lastly, the valuation metrics for the unoptimized random forest are presented below.

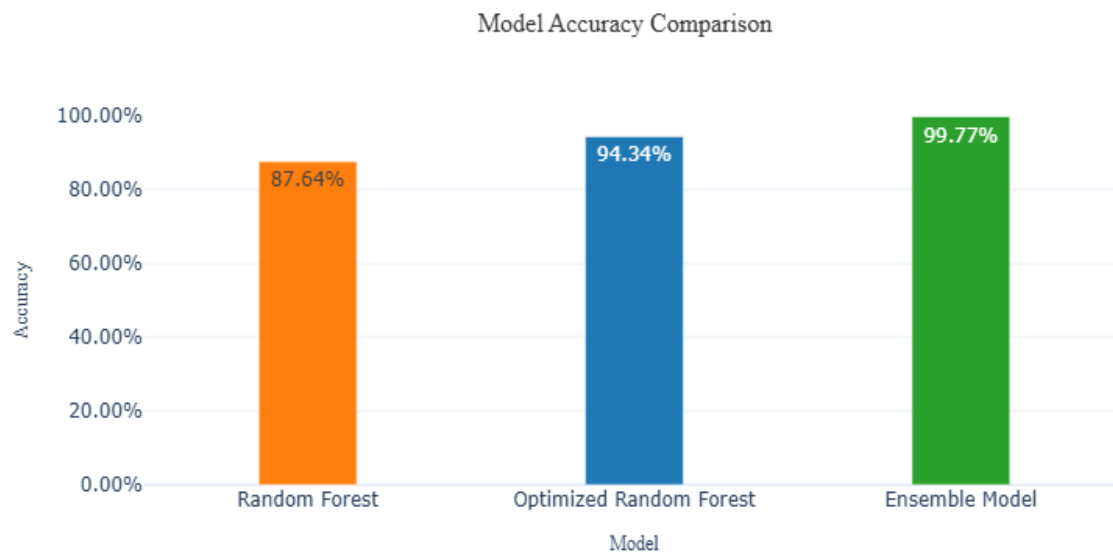
Classification report for unoptimized random forest model				
Class	Accuracy	Precision	Recall	F1-Score
0	0.8764	0.9987	0.8764	0.9336
1		0.0595	0.8713	0.1114

Table 7: Unoptimized random forest model- Evaluation metrics
Source: Created Model

From the tables above, it is evident that the ensemble model performed exceptionally well for non-fraudulent transactions, with near-perfect precision and recall. However, for fraudulent transactions, while the model had a high recall, its precision was lower, indicating that it sometimes misclassified non-fraudulent transactions as fraudulent. Nonetheless, the model struck a good balance, with an impressive overall performance in both classes compared to optimized random forest and random forest models.

Accuracy in particular surpassed the random forest models (Optimized & Unoptimized) by far. The random forest model achieved an accuracy of 87.64%, indicating solid performance, though there is room for improvement in distinguishing between the classes. Optimized random forest achieved an accuracy of 94.26%, which showed a significant improvement over the base random forest model, likely due to hyperparameter tuning and fine-tuning strategies.

The ensemble model on the other hand delivered a remarkable 99% accuracy, demonstrating near-perfect classification. This highlights the power of combining multiple models to enhance performance and reduce errors. The model accuracies are diagrammatically represented below.



*Figure 14: Model Accuracies Comparison
Source: Created Model*

Another important valuation metric used was the area under the curve (AUC). It quantifies how well the model distinguishes between the two classes, in this case, non-fraudulent and fraudulent transactions. A higher AUC score indicates a better-performing model that can accurately classify transactions as either fraudulent or non-fraudulent. The ROC curve is used to visualize the models' performance. It is a graph that evaluates the performance of a binary classification model. It plots the true positive rate (TPR, or sensitivity) against the false positive rate (FPR, or 1- specificity) at various threshold settings. In this paper, upon carrying out the AUC analysis:

- a) **All models outperform the random guess baseline:** The ROC curves for all three models (Random Forest, Optimized Random Forest, and Ensemble Model) lie significantly above the diagonal line (random guess), indicating that they are better than randomly predicting fraud.

Model-Specific Performance

- b) **Optimized random forest and ensemble model are top performers:** Both models had ROC curves that hug the top-left corner of the plot, suggesting excellent performance in terms of both sensitivity (True Positive Rate) and specificity (1 - False Positive Rate). Their AUC scores of 0.9859 & 0.9945 respectively, were very high, which indicated exceptional ability to

distinguish between fraudulent and non-fraudulent cases while minimizing both false positives and false negatives.

- c) **Random Forest Model:** While still better than random guessing, the random forest model showed slightly lower performance compared to the optimized and ensemble models. Its ROC curve was below the other two, and its AUC score of 0.9019 was lower.

Interpretation

- d) **High AUC Values:** The high AUC scores for all models, especially the optimized and ensemble models, suggested that they are very good at identifying fraudulent transactions while minimizing both false positives and false negatives. This is crucial in fraud detection, where accurate identification of fraudulent cases is essential to prevent losses and minimize disruption to legitimate transactions.

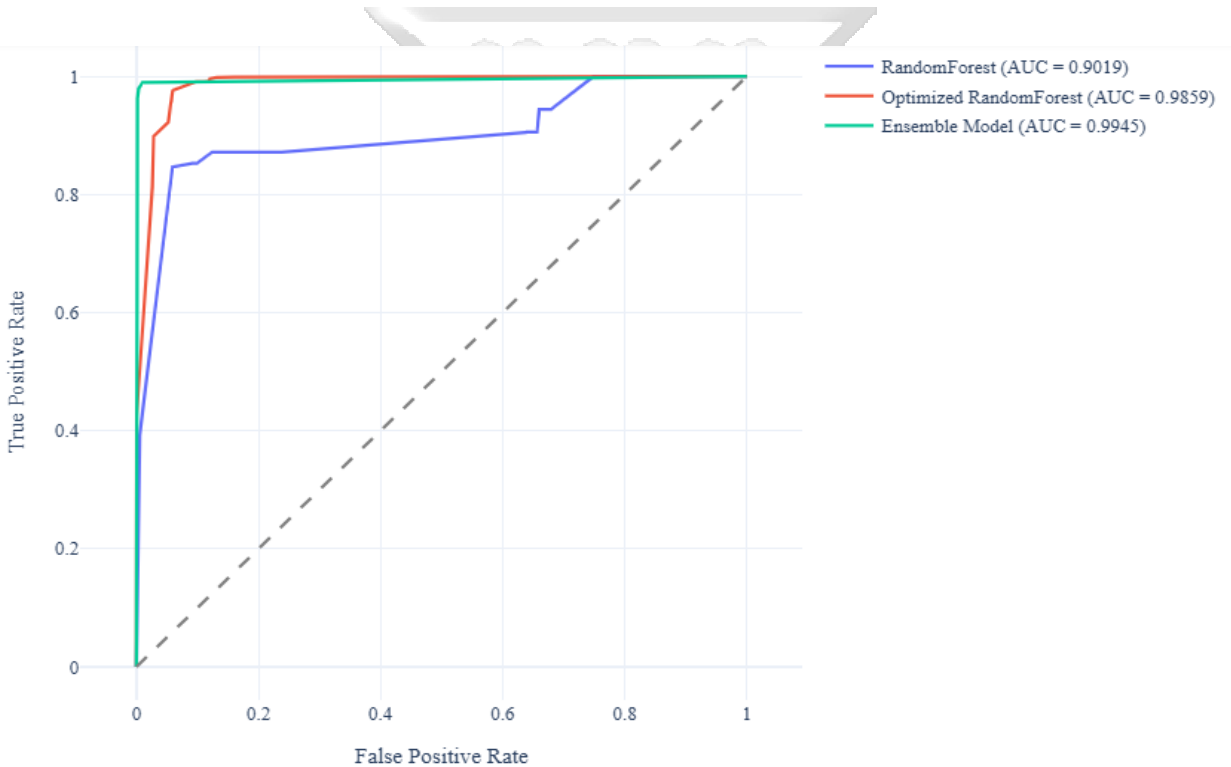


Figure 15: ROC Graph for the 3 Models
Source: Created Model

4.4.2. Confusion matrix

This one summarizes the model performance.

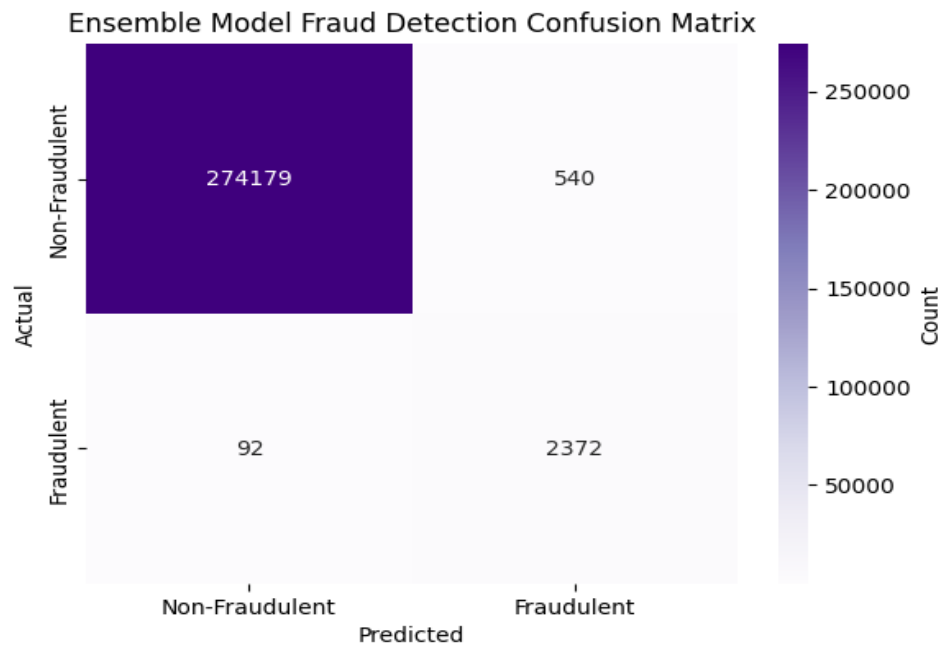


Figure 16: Ensemble Model- Confusion Matrix
Source: Created Model

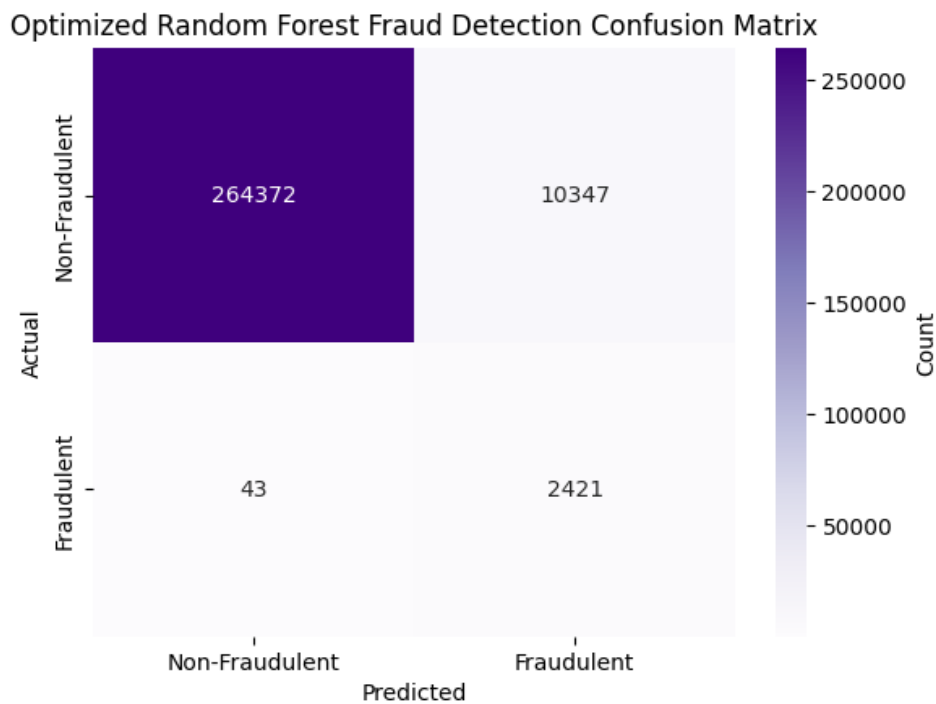


Figure 17: Optimized Random Forest Model- Confusion Matrix
Source: Created Model

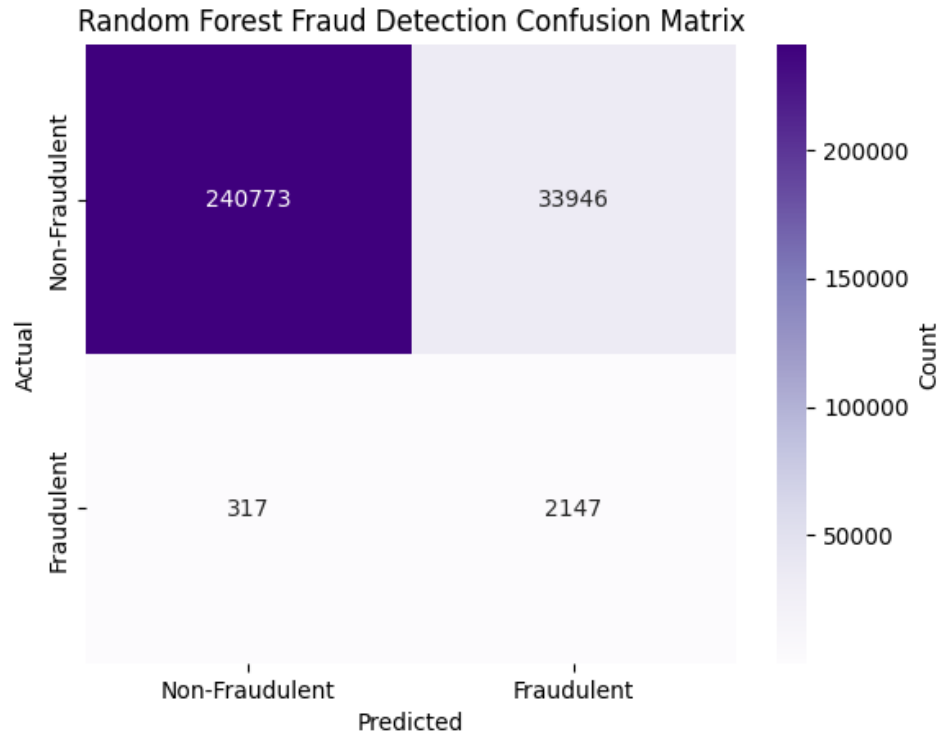


Figure 18: Random Forest Model- Confusion Matrix
Source: Created Model

By analyzing the confusion matrices above, it is evident that the ensemble model outperformed the other models. For example, if the actual non-fraudulent transactions are compared to the predicted non-fraudulent transactions, for all three models, it is easy to note that the ensemble model had the highest correctly predicted values (**274,179**), followed closely by the optimized random forest model (**264,372**), and finally, the random forest model (**240,773**). This further supports the other results that showcase the ensemble model as the best model.

4.4.3. Conclusion

In conclusion, the ROC curve analysis and confusion matrices demonstrate that the optimized random forest and ensemble model are the most promising models for fraud detection. They offer high accuracy and sensitivity in identifying fraudulent transactions while effectively minimizing both false positives and false negatives. However, of the two, the ensemble model does better.

Chapter 5

5. Discussion and conclusion

5.0.Introduction

This chapter presents the study's key findings, recommendations for future research to those interested in the topic, and limitations encountered during the study.

5.1.Summary of findings and conclusion

The paper's main objective was to develop a predictive model for detecting fraudulent transactions for regulated deposit-taking SACCOs in Kenya. To achieve this, an ensemble model that comprised two algorithms, one from machine learning and the other from deep learning was proposed. Random forest, a machine learning algorithm formed the base model, while multiple-layer perceptron, from deep learning, was used as the meta-model.

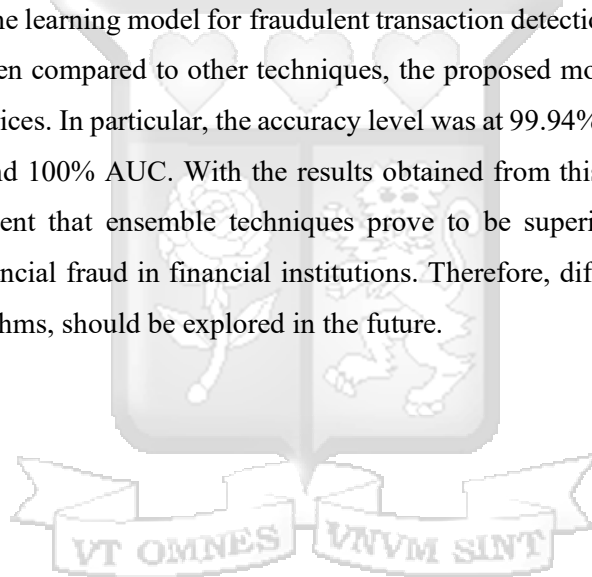
The proposed model's performance was evaluated against two other models: an optimized random forest and an unoptimized random forest. The purpose of doing so was to determine which model performed better. Once the three models were run, using the same training data set, and evaluated, the results indicated that the ensemble model, which utilized the stacking technique outperformed the other two models, hence making it a better model when it comes to predicting fraudulent transactions in the case of regulated deposit-taking SACCOs in Kenya.

In conclusion, overall, the ensemble model performs better if we consider the F1-score, precision, recall, and accuracy of both the optimized and unoptimized random forest on the positive class (Fraudulent). The results are summarized in the table below.

No.	Model Name	Accuracy	Precision	Recall	F1-Score
1.	Ensemble Model	0.9977	0.8146	0.9627	0.8824
2.	Optimized Random Forest	0.9437	0.1339	0.9744	0.2354
3.	Random Forest	0.8764	0.0595	0.8713	0.1114

*Table 8: Summary of Models' Performance
Source: Created Model*

It is important to note that the results are consistent with those obtained by other researchers. To mention a few, for example, (Baker et al., 2022) proposed the use of an ensemble learning technique- comprised of supervised machine learning models for credit card fraud detection. The results of the proposed model, when compared to those of other machine learning classifiers, proved to be superior, with an accuracy level of 100%, 97.3% precision, 83.7% F1-score, and 73.5% recall. Equally, (Khalid et al., 2024), in an attempt to address the shortcomings of previously proposed credit card fraud detection models- traditional machine learning methods and individual classifiers, the proposed ensemble model, which integrated Support Vector Machine (SVM), K-Nearest Neighbor (KNN), Random Forest (RF), Bagging, and Boosting classifiers, outperformed them across all the valuation metrics- accuracy, precision, recall, F1-Score. The conclusion and results obtained from that paper further emphasize the superiority of ensemble techniques in financial fraud detection. Finally, in a recent paper by (Talukder et al., 2024), which proposed an integrated multistage ensemble machine learning model for fraudulent transaction detection, the results were not short of impressive. Equally, when compared to other techniques, the proposed model attained the upper hand across all the valuation metrics. In particular, the accuracy level was at 99.94%, 99.91% precision, 99.14% recall, 99.52% F1-score, and 100% AUC. With the results obtained from this paper, as well as from the above examples, it is evident that ensemble techniques prove to be superior and, hence, suitable for addressing the issue of financial fraud in financial institutions. Therefore, different ensemble techniques, made up of different algorithms, should be explored in the future.



5.2.Recommendations

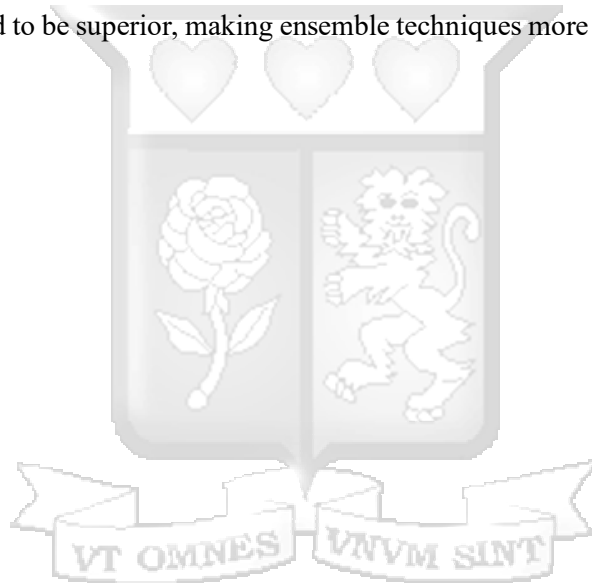
To future researchers who may be interested in this topic of research, and to SACCOs aiming to improve their security systems, a few things to consider include the following:

- a) **SACCOs should implement these models in their systems:** They should take advantage of the technological advancements taking place by integrating the proposed models into their transaction monitoring systems to improve fraud detection and operational efficiency. This would consequently increase customer confidence, as members feel assured that their funds are safe, considering their risk-averse nature.
- b) **Researchers should explore other models:** Future researchers should investigate different models that could help in addressing the issues of fraud. This paper presents a potential model but does not propose it as the only model that could work. Therefore, researchers should leverage both Artificial Intelligence (AI) and Machine Learning (ML) techniques to come up with more comprehensive predictive models.
- c) **SACCOs should work closely with researchers to ensure that better solutions are created:** Besides creating new models by the researchers, they should collaborate with the SACCOs and help them integrate this technology into their systems. Equally, the SACCOs should offer support to the researchers to ensure that the proposed models align with the specific needs of financial institutions. For example, they should avail real-world transactional data to the researchers and at the same time ensure that they adhere to the regulations that revolve around data protection.
- d) **Models created should be continuously improved:** With the rapid technological advances, everything is evolving including the tactics used by fraudsters. Therefore, the models proposed and implemented should be regularly updated to ensure that they consider the new patterns or fraud tactics adopted by the fraudsters.
- e) **Researchers should also pay attention to microfinance financial institutions:** Lastly, researchers should extend fraud detection research on microfinance institutions as well, given that they are also affected by financial fraud risk. Most papers have shifted their focus to banks, whereby they attempt to address the various types of fraud banks encounter. Very little has been done on SACCOs and other microfinance institutions. This, therefore, proves to be a potential research gap that needs to be addressed.

5.3. Limitations of the study

Some of the challenges encountered included the following:

- a) **Limited data access:** The main challenge encountered was on access to data. Financial data is considered to be highly confidential and thus not readily availed to the public. Also, with the now strict data policy regulations, it has become even more difficult to obtain the necessary data. To address the problem, financial institutions should consider depersonalizing the financial data, ensuring that sensitive information is still protected, and at the same time, available to the public.
- b) **Lack of standardized benchmarks’:** Another challenge was on performance benchmark. Given that there is little research done on addressing transactional fraud risk on SACCOs, it was hard to get the right comparable papers to assess the proposed model performance. To address the problem, two different models were created and their results were evaluated. The proposed model performance proved to be superior, making ensemble techniques more suitable, compared to single algorithms.



5.4. References

- ABSA. (2023). Absa Group Limited Integrated Report (IR) 2023. Nairobi: ABSA.
- Afriyie, J. K., Tawiah, K., Pels, W. A., Addai-Henne, S., Dwamena, H. A., Owiredo, E. O., Ayeh, S. A., & Eshun, J. (2023). A supervised machine learning algorithm for detecting and predicting fraud in credit card transactions. *Decision Analytics Journal*, 6, 100163. <https://doi.org/10.1016/j.dajour.2023.100163>
- Alammar, A., Yazeed Al Moayed, & Algeelani, N. A. (2022). Transaction fraud detector using KNN in deep learning. <https://doi.org/10.5281/ZENODO.7365019>
- Alexandropoulos, S.-A. N., Aridas, C. K., Kotsiantis, S. B., & Vrahatis, M. N. (2019). Stacking Strong Ensembles of Classifiers. In J. MacIntyre, I. Maglogiannis, L. Iliadis, & E. Pimenidis (Eds.), *Artificial Intelligence Applications and Innovations* (Vol. 559, pp. 545–556). Springer International Publishing. https://doi.org/10.1007/978-3-030-19823-7_46
- Ali, M. (2022, August). Supervised Machine Learning. Retrieved from Data Camp: <https://www.datacamp.com>
- Alushula, P. (2023, September 7). Retrieved from <https://www.businessdailyafrica.com>
- Amanja, D. M. (2015). Financial Inclusion, Regulation, and Stability: Kenyan Experience and Perspective.
- Andresen, M. A., & Ha, O. K. (2017). Routine activity theory. *CrimRxiv*. <https://doi.org/10.21428/cb6ab371.20b1f578>
- Arya. (2024, April 12). Fraud Detection Algorithms Using Machine Learning. Retrieved from <https://intellipaat.com>
- Baker, M., Mahmood, Z., & Shaker, E. (2022). Ensemble Learning with Supervised Machine Learning Models to Predict Credit Card Fraud Transactions. *Revue D Intelligence Artificielle*, 34, 509–518. <https://doi.org/10.18280/ria.360401>
- Bartsiotas, G. A., & Achamkulangare, G. (n.d.). Fraud prevention, detection, and response in United Nations system organizations.
- Beatrice, C. (n.d.). The effect of fraud on the financial performance of deposit-taking Savings and Credit Cooperative societies in Kenya.
- Bekele, W. D. (2023). Determinants of Financial Inclusion: A. *Journal of African Business*, 301-319.
- Brownlee, J. (2019, January 8). A Gentle Introduction to the Rectified Linear Unit (ReLU). *MachineLearningMastery.Com*. <https://www.machinelearningmastery.com/rectified-linear-activation-function-for-deep-learning-neural-networks/>
- CBK. (2023). Kenya Financial Sector Stability Report. Nairobi: Financial Sector Regulators.
- Chelangat, B. (2014). The effects of fraud on the financial performance of deposit-taking SACCOS in Kenya. 1-9.

- Cheruku, S. K. (2019, November 6). Retrieved from <https://www.linkedin.com/pulse/how-ensemble-learning-increases-model-performance-machine-cheruku/>
- Dagal, I. (2024). Multi-Layer Perceptrons-Based on Hill Function (Mlp-Hf) for Data Classification. <https://doi.org/10.2139/ssrn.4978377>
- Fernandez, A., Garcia, S., Herrera, F., & Chawla, N. V. (2018). SMOTE for Learning from Imbalanced Data: Progress and Challenges, Marking the 15th Anniversary. *Journal of Artificial Intelligence Research*, 61, 863–905. <https://doi.org/10.1613/jair.1.11192>
- FinAccess. (2021). Measuring Kenya's Financial Inclusion Journey.
- Gamlath Mohottige Mudith Sujeewa, M. S. (2018). The new fraud triangle theory - integrating ethical values of employees. *International Journal of Business, Economics, and Law.*, 52-53.
- Gopinathan, G. A. (2016). Fraud prevention, detection, and response in United Nations system organizations.
- Han, Y., Yao, S., Wen, T., Tian, Z., Wang, C., & Gu, Z. (2020). Detection and Analysis of Credit Card Application Fraud Using Machine Learning Algorithms. *Journal of Physics: Conference Series*, 1693(1), 012064. <https://doi.org/10.1088/1742-6596/1693/1/012064>
- Islam, S., Haque, Md. M., & Rezaul Karim, A. N. M. (2024). A rule-based machine learning model for financial fraud detection. *International Journal of Electrical and Computer Engineering (IJECE)*, 14(1), 759. <https://doi.org/10.11591/ijece.v14i1.pp759-771>
- Jaiswal, S. (2024, July 28). Multilayer Perceptrons in Machine Learning. Retrieved from <https://www.datacamp.com/tutorial/multilayer-perceptrons-in-machine-learning>
- Jamhuri, M. (2023, April 15). Understanding the Adam Optimization Algorithm: A Deep Dive into the Formulas. Medium. <https://jamhuri.medium.com/understanding-the-adam-optimization-algorithm-a-deep-dive-into-the-formulas-3ac5fc5b7cd3>
- Joshi, A. (2021). Decision Tree Algorithm for Credit Card Fraud Detection. Webology. <https://doi.org/10.29121/WEB/V18I4/103>
- Joshua, & Abosede, A. (2023). Fraud theories and global cybercrime during the COVID-19 pandemic: the extended S.C.C.O.R.E model.
- Kabaiku, P. M. (2018). Savings and Credit Cooperative Societies (SACCOs) in Kenya: A theo-social justification. 2-7.
- Kalirane, M. (2024, July 19). Retrieved from <https://www.analyticsvidhya.com/>
- Kamau, J. M. (2015). Strategies adopted by commercial banks in Kenya to combat fraud. *International Journal of Current Business and Social Sciences*, 1-18.

- Khalid, A. R., Owoh, N., Uthmani, O., Ashawa, M., Osamor, J., & Adejoh, J. (2024). Enhancing Credit Card Fraud Detection: An Ensemble Machine Learning Approach. *Big Data and Cognitive Computing*, 8(1), 6. <https://doi.org/10.3390/bdcc8010006>
- Khan, S., Alourani, A., Mishra, B., Ali, A., & Kamal, M. (2022). Developing a Credit Card Fraud Detection Model using Machine Learning Approaches. *International Journal of Advanced Computer Science and Applications*, 13(3). <https://doi.org/10.14569/IJACSA.2022.0130350>
- Khodabakhshi, M., & Fartash, M. (2016). Fraud Detection in Banking Using KNN (K-Nearest Neighbor) Algorithm.
- Kigerl, A. (2021). Routine activity theory and malware, fraud, and spam at the national level. *Crime, Law and Social Change*, 76(2), 109–130. <https://doi.org/10.1007/s10611-021-09957-y>
- Kumar, J., & Saxena, V. (2022). Rule-Based Credit Card Fraud Detection Using User's Keystroke Behavior. In R. Kumar, C. W. Ahn, T. K. Sharma, O. P. Verma, & A. Agarwal (Eds.), *Soft Computing: Theories and Applications* (Vol. 425, pp. 469–480). Springer Nature Singapore. https://doi.org/10.1007/978-981-19-0707-4_43
- Lai, G. (2023). Artificial Intelligence Techniques for Fraud Detection. <https://doi.org/10.20944/preprints202312.1115.v1>
- Lokanan, M. E., & Sharma, K. (2022). Fraud prediction using machine learning: The case of investment advisors in Canada. *Machine Learning with Applications*, 8, 100269. <https://doi.org/10.1016/j.mlwa.2022.100269>
- Malato, G. (2023, May 03). An Introduction to the Shapiro-Wilk Test for Normality. Retrieved from <https://builtin.com/data-science/shapiro-wilk-test>
- Marco Sánchez-Aguayo, L. U.-A.-J. (2021). Fraud Detection Using the Fraud Triangle Theory and Data Mining Techniques. 1-4.
- Motompa, D. K. (2016). Factors influencing the growth of saving and credit cooperative societies in Kenya: a case study of Kajiado East sub-county.
- Mugoshi, S. E. (2021). Use of artificial intelligence algorithms to enhance fraud detection in the Banking Industry.
- Multi-Layer Perceptron Neural Network using Python – Machine Learning Geek. (2021, April 23). <https://machinelearninggeek.com/multi-layer-perceptron-neural-network-using-python/>
- Mwithi, J. M. (2015). Strategies adopted by commercial banks in Kenya to combat fraud: a survey of selected commercial banks in Kenya. 3.
- Mytnyk, B., Tkachyk, O., Shakhovska, N., Fedushko, S., & Syerov, Y. (2023). Application of Artificial Intelligence for Fraudulent Banking Operations Recognition. *Big Data and Cognitive Computing*, 7(2), 93. <https://doi.org/10.3390/bdcc7020093>

- Ndege, A. (2024, April 17). Retrieved from Techcabal: <https://techcabal.com/>
- Ndolo, R. O. (2023). Knowledge Graph for Fraud Detection: Case of Fraudulent Transactions Detection in Kenyan SACCOs. Springer, 179-185.
- Negi, D. (2021). Automating Fraud Detection in Financial Services An AI-based Approach. 1315-1325.
- Ng'el, A. K. (1990). Evolution, Structure, and Performance of Kenya's Financial System. Savings and Development, 265-284.
- Njoki, C. (2023, March 13). Retrieved from Chamasoft: <https://blog.chamasoft.com>
- Ogola, D. O. (2016). The impact of financial development on economic growth in Kenya.
- Ojino, R., & Ndolo, R. (2023). Knowledge Graph for Fraud Detection: Case of Fraudulent Transactions Detection in Kenyan SACCOs. In Artificial Intelligence: Towards Sustainable Intelligence: First International Conference, AI4S 2023, Pune, India, September 4-5, 2023, Proceedings (Vol. 1907). Springer Nature Switzerland. <https://doi.org/10.1007/978-3-031-47997-7>
- Opany, N. (2022, October 31). Retrieved from <https://www.kanisa-sacco.org>
- Patel, J. (2022, December 28). Sigmoid Function: Derivative and Working Mechanism. Analytics Vidhya. <https://www.analyticsvidhya.com/blog/2022/12/sigmoid-function-derivative-and-working-mechanism/>
- Popescu, M.-C., Balas, V. E., Perescu-Popescu, L., & Mastorakis, N. (2009). Multilayer Perceptron and Neural Networks. 8(7).
- Predictive Analytics: Enhancing Risk Mitigation for SACCOs. (2023, August 29). <https://ronalds.co.ke/enhancing-risk-mitigation-in-kenyan-saccos-through-predictive-analytics/>
- Rashedi, K. A., Ismail, M. T., Al Wadi, S., Serroukh, A., Alshammari, T. S., & Jaber, J. J. (2024). Multi-Layer Perceptron-Based Classification with Application to Outlier Detection in Saudi Arabia Stock Returns. Journal of Risk and Financial Management, 17(2), 69. <https://doi.org/10.3390/jrfm17020069>
- Rawal, Y. (2021, July 2). Wakandi. Retrieved from <https://blog.wakandi.com>
- Riskiyadi, Moh. (2024). Detecting future financial statement fraud using a machine learning model in Indonesia: A comparative study. Asian Review of Accounting, 32(3), 394-422. <https://doi.org/10.1108/ARA-02-2023-0062>
- Ruankaew, T. (2013). The Fraud Factors. IJMAS, 2, 1-5.
- Sai, C. V., Das, D., Elmitwally, N., Elezaj, O., & Islam, M. B. (2023). Explainable AI-driven financial Transaction Fraud Detection Using Machine Learning and Deep Neural Networks. <https://doi.org/10.2139/ssrn.4439980>
- Saiful Islam, M. M. (2023). A rule-based machine learning model for financial fraud detection. International Journal of Electrical and Computer Engineering (IJECE), 759-769.

- Salman, H. A., Kalakech, A., & Steiti, A. (2024). Random Forest Algorithm Overview. *Babylonian Journal of Machine Learning*, 2024, 69–79. <https://doi.org/10.58496/BJML/2024/007>
- Sánchez-Aguayo, M., Urquiza-Aguiar, L., & Estrada-Jiménez, J. (2021). Fraud Detection Using the Fraud Triangle Theory and Data Mining Techniques: A Literature Review. *Computers*, 10(10), 121. <https://doi.org/10.3390/computers10100121>
- Sánchez-Aguayo, M., Urquiza-Aguiar, L., & Estrada-Jiménez, J. (2022). Predictive Fraud Analysis Applying the Fraud Triangle Theory through Data Mining Techniques. *Applied Sciences*, 12(7), 3382. <https://doi.org/10.3390/app12073382>
- SASRA. (2022). The SACCO Supervision Annual Report. Nairobi.
- Saxena, Jaydip Kumar & Vipin. (2022). Rule-Based Credit Card Fraud Detection Using User's Keystroke Behavior. *Researchgate*, 2-12.
- Sayal, K., & Singh, G. (2020). What Role Does Human Behaviour Play in Corporate Frauds?
- Skalak, S., Alas, M., & Sellitto, G. (2015). Fraud: An Introduction (pp. 1–23). <https://doi.org/10.1002/9781119200048.ch1>
- Shihembetsa, E. M. (2021, August). Use of Artificial Intelligence Algorithms to enhance fraud detection in the Banking Industry.
- Singh, K. S. (2020). What Role Does Human Behaviour Play in Corporate? ISSN.
- Soni, B. (2023, April 25). Medium. Retrieved from Medium: <https://medium.com/>
- Sorokin, M. (2024, January 17). Binary Cross-Entropy: Mathematical Insights and Python Implementation. Medium. <https://medium.com/@vergotten/binary-cross-entropy-mathematical-insights-and-python-implementation-31e5a4df78f3>
- Stephen N. Skalaka, M. A. (2015). Fraud introduction. 1-23.
- Steven. (2024, May 19). What is a Transaction Fraud Explained: Types, Impacts, and Transaction Fraud Detection. Retrieved from <https://www.idstrong.com>
- Sujeewa, G. M. M., Yajid, M. S. A., Khatibi, A., Azam, S. M. F., & Dharmaratne, I. (2018). The new fraud triangle theory - integrating ethical values of employees. 16(5).
- Talukder, Md. A., Khalid, M., & Uddin, M. A. (2024). An integrated multistage ensemble machine learning model for fraudulent transaction detection. *Journal of Big Data*, 11(1), 168. <https://doi.org/10.1186/s40537-024-00996-5>
- Vanini, P., Rossi, S., Zvizdic, E., & Domenig, T. (2023). Online payment fraud: From anomaly detection to risk management. *Financial Innovation*, 9(1), 66. <https://doi.org/10.1186/s40854-023-00470-w>
- Wanyama, D. N., & Reuben, D. J. M. (2022). Internal control measures on fraud detection and prevention in fintech companies in Kenya. 10(6).

Wei, D. (2024, June 8). Demystifying Machine Learning Models: Random Forest. Medium.
<https://medium.com/@weidagang/demystifying-machine-learning-models-random-forest-f992dc50b427>

Yasar, K. (2024, January). Knowledge graph in ML. Retrieved from TechTarget:
<https://www.techtarget.com/>

