

A Review of Machine Learning Algorithms for Fraud Detection in Credit Card Transaction

Kha Shing Lim, Lam Hong Lee and Yee-Wai Sim*

Quest International University, Ipoh, Perak, Malaysia

*Corresponding Author

Summary

The increasing number of credit card fraud cases has become a considerable problem since the past decades. This phenomenon is due to the expansion of new technologies, including the increased popularity and volume of online banking transactions and e-commerce. In order to address the problem of credit card fraud detection, a rule-based approach has been widely utilized to detect and guard against fraudulent activities. However, it requires huge computational power and high complexity in defining and building the rule base for pattern matching, in order to precisely identifying the fraud patterns. In addition, it does not come with intelligence and ability in predicting or analysing transaction data in looking for new fraud patterns and strategies. As such, Data Mining and Machine Learning algorithms are proposed to overcome the shortcomings in this paper. The aim of this paper is to highlight the important techniques and methodologies that are employed in fraud detection, while at the same time focusing on the existing literature. Methods such as Artificial Neural Networks (ANNs), Support Vector Machines (SVMs), naïve Bayesian, k-Nearest Neighbour (k-NN), Decision Tree and Frequent Pattern Mining algorithms are reviewed and evaluated for their performance in detecting fraudulent transaction.

Key words:

Data Mining; Machine Learning; Credit Card; Fraud Detection.

1. Introduction

Transactions between credit cards and electronic businesses are the primary areas that the fraudsters are abusing and exploiting. This is due to the fact that there are various loopholes in the existing detection methods, and a major factor – business negligence. There are many types of frauds that can be performed, such as deliberately underreporting or omitting income, overstating or claiming false amount of deductions, money laundering, swindle transactions [5] and more. To overcome fraud cases that are happening in the banking industry, rule-based systems, also known as Production Systems or Expert Systems, are used to store and manipulate fraud knowledge, in order to interpret the information in a meaningful way.

Rule-based approach is one of the simplest and most popular Artificial Intelligence (AI) techniques that uses rules as the representation of knowledge coded into the system, and provides reasoning for the context of pattern

matching. A conventional rule-based system consists of a list of rules (rule base) which represent the knowledge, an inference engine or semantic reasoner, which infers information or acts based on the interaction of input and the rule base, a temporary working memory, and a user interface which handles the incoming flow of data and the outgoing flow of the prediction results [2]. In short, a rule-based system infers knowledge that stored in its rule base and mimics the reasoning of a human expert in solving a knowledge-intensive problem.

Over the past few decades, rule-based systems have been widely implemented in fraud detection and prediction. Although rule-based systems can be accurate in detecting and predicting frauds, they require huge computational power for pattern matching and the rules have to be specially derived for the working domain [2]. Update and modification of the rule base are also complicated. For example, when introducing new knowledge to identify new fraud patterns, contradictions might be introduced between the new rules and existing ones. Moreover, fraudsters are highly adaptive-able and when given enough time, they will always be able to find ways to circumvent preventive measures [2]. For instance, they may bypass simple pattern matching or rule-based detection, and make the system regard the transaction made is a genuine one. This leads to the system failing to detect fraud occurrences accurately. Besides these shortcomings, rule-based systems are also not equipped with analysis and prediction capabilities. Having these capabilities to perform on received data would aid in better identification of future fraud incidences. Therefore, recent advancements which implement data mining and machine learning techniques for fraud prediction have widely been found. These new technologies aim to improve the detection and prediction of frauds, and offer correlation analysis in fraud data.

Being the subfields of AI, data mining and machine learning techniques can analyse and discover patterns in large datasets and subsequently produce hidden insights, through learning from historical relationships and the trends in the data [3]. There are multiple data mining and machine learning techniques that can be used in not only identifying fraud transactions, but also predicting the suspicious rate of

transaction data that may transact over time. In this paper, methods such as Artificial Neural Networks (ANNs), Support Vector Machines (SVMs), Bayesian Classification, k-Nearest Neighbour (k-NN), Decision Tree (DT) and Frequent Pattern (FP) Mining algorithms are reviewed for their capability and performance in discovering frauds. These methods involve distinguishing fraudulent financial data from authentic data, thereby disclosing fraudulent behaviour or activities, and enabling decision makers to develop appropriate strategies to decrease the negative impact of fraud [15].

In a nutshell, recent research works that use data mining and machine learning techniques to overcome the difficulty in detecting fraud activities are investigated and analysed. These techniques carry out the work of scrutinizing the behaviour of user's transaction history and analysing the transaction data in-depth to determine the legitimacy of a transaction. The ultimate goal of this paper is to provide a comprehensive review of data mining and machine learning techniques in analysing fraudulent behaviours, identifying the major sources and characteristics of the data, based on the fraud detection and prediction works that had been investigated.

2. Literature Review

In the context of AI, fraud detection is viewed as a classification problem [2,69], in which the objective is to correctly classify credit card transactions as legitimate or fraudulent. The detection methods were usually embedded in a Fraud Detection System, whereby the model contains the rules or knowledge that is used to identify and prevent fraud. A Fraud Detection System, as depicted in Fig. 1, takes the accumulated data in the bank's database, performs training and learning, then output a model that best represents the characteristics of the transaction data.

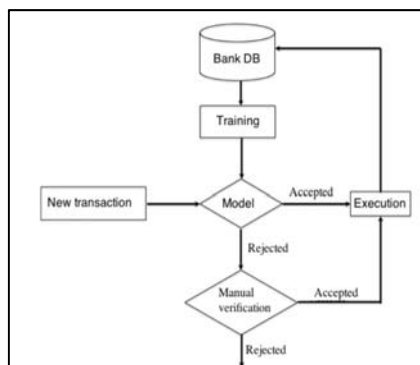


Fig. 1: Flow of Fraud Detection System [2].

The model is then used to decide new transactions, whether it can be accepted as a genuine transaction, or reject it as a fraudulent transaction. A transaction that is accepted by the model will be executed and then added to the database to improve the model. Whereas, a rejected transaction will pass to manual check. If the rejected transaction is regarded as normal after checking, the transactions are then executed and the information will be added into the bank's database, otherwise, the transaction is rejected.

A big part of the fraud detection and prevention process is related to training and learning from the transaction data to identify new frauds [2] [14] [15]. As such, it is important to design a model with the best Data Mining and Machine Learning algorithm that quickly detects fraud and takes immediate preventive action. A well-designed model would not only identify frauds accurately but also determine the probability of fraudulent behaviour.

2.1 Artificial Neural Networks

Artificial Neural Networks, or short for ANNs, are a computing system that is inspired by the way of how biological neural networks works. ANNs are composed of a large number of highly interconnected processing elements called an artificial neuron, which receives signals, processes it and transmits the processed signal to the next artificial neurons. A neural network when used for fraud detection, is typically a collection of neuron-like processing units with weighted connections between the units. With the ANNs' remarkable ability to derive meaning from complicated or imprecise data, it is increasingly proposed as a state-of-the-art way to identify frauds [10,70-71].

Montague (2012) [53] proposed Auto Encoder (AE) as one of the types of neural networks for fraud detection. An AE is divided into two parts: an encoder and a decoder. AE follows a feed-forward ANN architecture, except that the output layer has the same number of neurons as the input layer. The transition between the first and second layer represents the encoder and the transition between the second layer and the third layer represents the decoder. AE will use the input data itself as its target value and learn the common patterns that shared by the majority of the training data during the construction period. Fraud cases will have a different distribution as compared to a normal transaction. For data points that exhibit high errors and show anomaly compared to those majority transaction patterns, these characteristics are reflective of fraudulent transactions.

On the other hand, Bansal and Suman (2014) [64] proposed Self-Organizing Map (SOM) as one of the types of neural networks and trained in an unsupervised learning environment for fraud detection. Like ANNs, the SOM operates in two modes, training and mapping. The training phase will build the map by using inputs, while mapping

automatically classify a new input vector. Getting the best matching node is done by running through all weight vectors and calculating the distance from each weight to the sample vector. The most commonly used method to determine the distance is the Euclidean Distance. The mathematical formula of Euclidean Distance is:

$$Dist = \sqrt{\sum_{i=0}^n (V_i - W_i)^2} \quad (1)$$

where $\mathbf{V}$ is the current input vector and $\mathbf{W}$ is the node's weight vector.

Based on the Bansal and Suman's research (2014) [12], which successfully applied SOM for fraud detection, it describes that during the process of reducing dimensions and clustering the transaction data into fraudulent and genuine sets. It takes into account of such values and finally outputs the patterns and cluster as an output vector:

- Account related: Account number, currency of account, account opening date, last date of credit or debit available balance.
- Customer related: Customer ID, customer type like high profile or low profile.
- Transaction related: Transaction number, location, currency and timestamp.

Serrano, Costa, Cardonha, Fernandes, and Júnior (2012) [54] proposed the use of ANN as a predictor transaction. In the case of ANN predictor, the ANN returns the prediction of certain input data and classify those data into a group from a set of predefined groups (fraudulent and non-fraudulent). Based on their research work, they make use of Feed-Forward Networks, which are the most widely used for time-series prediction [12] [54] [58]. In Feed-Forward Networks, the information travels only in one direction, from the input to the output, without any feedback nor between neurons of the same layer. The most commonly Feed-Forward Network used is the Multilayer Perceptron (MLP), where all the neurons of a layer are connected to all the neurons of the following layer. By using this type of ANN, they able to gives a hard output, which can be '1' for the fraudulent transaction and '0' for non-fraudulent transaction.

2.2 Naïve Bayesian Classifier

Naïve Bayesian Classifier or naïve Bayesian is a supervised Machine Learning method that uses a training dataset with known target classes to predict future or any incoming class value. It can predict the class membership probabilities such as the probability that a given tuple belongs to a particular class. Naïve Bayesian is based on Bayes' theorem of the posterior probability. It assumes class-conditional independence, which means the effect of an attribute value

on a given class is independent of the values of the other attributes. It is made to simplify the computations involved and, in this sense, is considered as "naïve". Bayes theorem provides a way of calculating the posterior probability. More description can be found below:

$$p(C_k | x) = \frac{p(C_k)p(x|C_k)}{p(x)} \quad (2)$$

where $\mathbf{p}(C_k|x)$ is the posterior probability of class (target) given predictor (attribute), $\mathbf{p}(C_k)$ is the prior probability of class, $\mathbf{p}(x|C_k)$ is the likelihood which is the probability of predictor given class, and $\mathbf{p}(x)$ is the prior probability of predictor.

In plain English, using Bayesian theorem terminology, the above equation can be written as:

$$Posterior = \frac{Prior \times Likelihood}{Evidence} \quad (3)$$

Milgo & Carolyne (2016) [13] proposed a Bayesian approach as the main classification method for fraud detection targeting onto ATMs, as shown below:

$$Fraud \setminus Evidences = \frac{P(Evidences \setminus Fraud) \times P(Fraud)}{P(Evidences)} \quad (4)$$

Their research work suggests that by using probability-based model, banks would be able to identify the main security issues encountered with the use of ATM cards and establish internal control mechanisms, and in the same time sought to deter the possibility of card fraud.

Milgo & Carolyne also focused heavily onto data pre-processing stage as Naïve Bayesian is subjected to missing data and number of attributes in a dataset. They also detail the implementation of the solution from data pre-processing to fraud classification fraud by feeding Bayesian model reasonable training data and finally provide transaction probability of the transaction. With the Bayesian model, it can ease many of the theoretical and computational difficulties of rule-based systems [2] [5] by offer posterior probability of fraud and managing probabilistic knowledge. Viaene, Derrig, and Dedene (2004) [55] proposed Naive Bayesian by combining with AdaBoost algorithm which developed by Freund and Shapire (1995) [63] in classifying valid or non-valid Personal Injury Protection (PIP) automobile insurance claims. In their research work, they described the mechanics of boosting rest on the construction of a sequence of classifiers, where each classifier is trained on a resampled (or reweighted) training set where those training data instances that were poorly predicted in the previous runs receive a higher weight in the next run. After a fixed number of iterations, the constructed models are then combined by weighted or simple voting schemes. The idea underlying the sequential perturbation of the training

data is that the base learner – specifically, Naïve Bayes gets to focus incrementally on those regions of the data that are harder to learn. With Boosting Naïve Bayes as fraud detection model, they able to deter fraudulent claims and also identified the characteristics of such claims that distinguish them from valid claims.

Kiran, Guru, Kumar, Kumar, Katariya, and Sharma (2008) [56] described the implementation of Naïve Bayes and k-NN algorithm on same credit card dataset and calculate the precision of algorithms to identify the fraudulent transactions. For k-NN algorithm it will be described in the next section C. Their research work suggests that by implementing Naïve Bayesian as Machine Learning classifier for fraud detection, it not only provides probabilities of fraud but also able to learn in an efficient, fast and high in accuracy for real-world scenarios.

2.3 k-Nearest Neighbour

k-Nearest Neighbour (k-NN) falls into the supervised Machine Learning family, and is an example of instance-based learning where it categorised objects based on closest feature space in the training set. k-NN is considered as a “lazy” learning algorithm as it does not use the training data points to do any generalization. There is no explicit training or learning phase in k-NN, and the training happens at the time when prediction is requested. New instance is compared with existing ones in the feature space by using Euclidean Distance, and the closest existing instance is used to assign the class to the new one. As k-NN algorithm is based voting scheme, which the most many neighbour is consider as winner and is used to label the query.

Sudha and Nirmal Raj (2017) [19] presented k-NN as a Machine Learning model in helping detecting fraudulent transaction. In the process of k-NN during detecting fraudulent transaction, the model classifies any incoming transaction by calculating most many nearest neighbours to new transaction. If the most neighbours are fraudulent, then the transaction is classified as fraudulent, else, it is classified as legal transaction. They also described that various factors had to be considered before classifying the new transaction. For instance, the financial characteristics of a transaction such as card number, transaction amount, or time since last purchase needs to be collected first before perform classification. With those information, k-NN can then properly classify the new transaction as fraudulence or legitimate.

Heta (2018) [65] proposed k-Nearest Neighbour, Random Forest, AdaBoost and Logistic Regression for fraud detection and a comparative study is shown between them. In his research work, he described that k-NN algorithm gives better accuracy and efficient result as it often recognizes as “outlier detection” [16] [17] for classification

between fraudulent and non-fraudulent transactions. It also shown that k-NN uses less memory compared to other algorithms in the author’s experiment result. Like other research works, Heta also focused on data pre-processing and exploratory data analysis to identify how many variables that needed for best training in Machine Learning models.

Malini and Pushpa (2016) [66] proposed k-NN algorithm and Hidden Markov Model (HMM) as classification Machine Learning model for credit card fraud detection. Their research work aims to minimize the false alarm and increase the detection rate by implement these two models in fraud detection. They described that by using the two models together, HMM able to analyse user’s behaviour and helps in minimizing the fraud rates occurred by k-NN, thus retaliate further fraudulent activities more efficiently [67] [68]. In other words, k-NN algorithm will be responsible for detect fraudulent transaction using distance calculation while the HMM will process based on credit card user’s behaviour pattern and check over the upcoming transaction.

2.4 Support Vector Machines

Support Vector Machines (SVMs) is a supervised Machine Learning algorithm that given a set of training samples, each marked as one or other categories, it will assign the best category to the new data. In classification, new data are mapped into the feature space and predicted to belong to which category based on which side of the gap they fall into. The algorithm is based on finding the Hyperplane in the feature space to divide the data points according to their categories or classes. Hyperplane, in human language, is a decision boundary that separates between a set of different classes.

The margin of the Hyperplane is the one that gives the greatest separation between the classes and is known as maximum-margin Hyperplane. The instances or data points that are nearest to the maximum-margin of the Hyperplane are called Support Vectors, and there is always at least one Support Vector for each class, and often there are more.

Abdelhamid, Khaoula and Atika (2014) [26] proposed SVMs for credit card fraud detection, which in each test instances of the transaction data, SVMs classifier will use and determine the instance falls into which category. For example, if the instance falls into fraudulent class, it is declared as fraud, else it is declared as a legitimate transaction. According to Abdelhamid, Khaoula and Atika, they describe that multiple attributes that can be used to improve the detection rate. However, adding more or irreverent attributes can make the classifier inefficient. The important attributes that listed out in the paper and can be used for detecting fraud in SVM classifier are:

- Customer ID

- Transaction amount
- Frequency of card usage including Date and Time
- Average amount of transactions
- Place

They also utilized a trick which called as Radial Basis Function Kernel, or known as RBF Kernel in helping implicitly map the input vector into the feature space and find the non-linear decision boundary. RBF Kernel able map the and highly imbalance data points from non-linearity to linearity. The RBF Kernel is introduced as:

$$K(x|x') = \exp\left(-\frac{\|x-x'\|^2}{2\sigma^2}\right) \quad (4)$$

where $\|x-x'\|^2$ is the squared Euclidean Distance between two data points x and x' , σ is a standard deviation, and x and x' are a vector containing data about a single observation.

The next step – cross-validation is performed after running the SVMs classification algorithm with RBF Kernel onto transaction datasets. Cross-validation is executed for assessing how accurately a model will perform in practice. After validation is done, a learned model is then outputted. Sallehuddin, Ibrahim, Zain and Elmi (2014) [57] proposed SVMs and Artificial Neural Network for detecting SIM box bypass fraud, whereby SIM cards are used to channel national and multinational calls away from mobile operators and deliver as local calls. According to their research work and experiment result, they focused heavily in data pre-processing stage that comprises of feature extraction, handling missing data, removing outliers and data normalization before feeding the data into training the Machine Learning model. After the data pre-processing stage, the development of ANN and SVM model is then started.

Training, Testing and hyper-parameter tweaking are done in both model development and a comparison of accuracy is shown based on their experiment result. From Fig. 2, it illustrates that by having SVM train on more data, it outperforms ANN and able to have a consistent accuracy on SIM Box fraud detection.

Banerjee, Bourla, Chen, Kashyap and Purohit (2018) [58] proposed and experimented several algorithms in their research works for fraud detection as presented in Fig. 3. Algorithms include Random Forest Classifier, k-NN, Naive Bayes Classifier, Logistic Regression, ANN with Multilayer Perceptron and SVMs had been experimented and reviewed based on a dataset provided by University of California, San Diego and the Fair Isaac Corporation. In their research work, they described that SVM achieved significantly higher accuracy compared other five algorithms. However, in contrast to other algorithms, the SVM algorithm took much more time and computing power to complete the fitting of the model. Compared to the

Random Forest Classifier, the model takes much more time to execute. As a result, when processing real-time data such as credit card transactions on datasets that would be much larger than the current dataset that they trained with, the SVM model would need be made more efficient in order to process and classify fraud data in a reasonable amount of time.

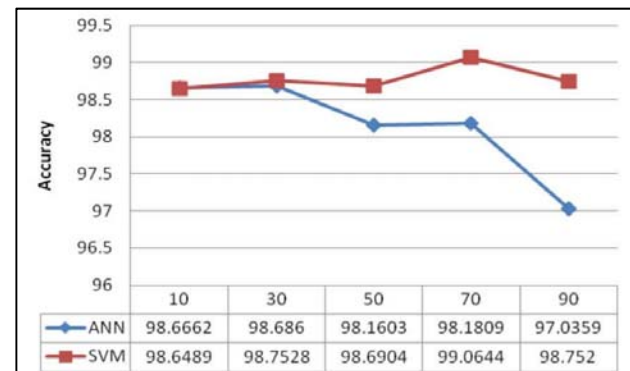


Fig. 2: Comparison of SVM and ANN classification accuracy [57].

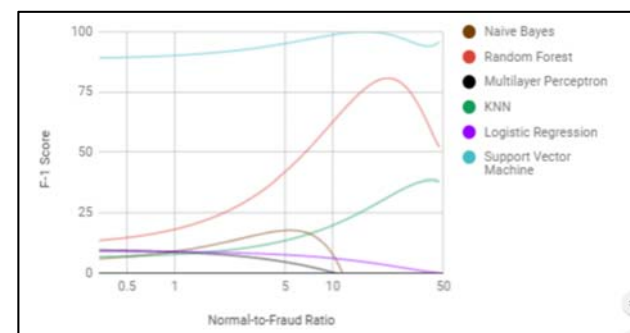


Fig. 3: Accuracy of classification approaches for fraud detection [58].

2.5 Decision Tree Induction

Decision Trees Induction falls into supervised Machine Learning family whereby it learns the class-labelled training tuples and try to predicts a class for a given input vector. A Decision Tree is a flowchart-like tree structure, where it typically starts with a single topmost node, which branches into possible nodes. Each of those branched nodes leads to additional internal nodes, which test on an attribute and branch off into other nodes. By continuing branching into more nodes, this eventually gives it a treelike shape. Decision Tree will continue to expand until every branch reaches an endpoint. Meaning that there are no more conditions to consider. The endpoint, or called as lead node will represents the final labels or choices in the Decision Tree.

There are three main algorithms in Decision Tree, namely ID3, C4.5 and CART, which are devised to not only best classify a new data, but also best building the Decision Tree in an optimal way and generates rules out from it. The growing of Decision Tree algorithms is often related to Decision Tree Learnings which it deals with the construction of an optimal Decision Tree from class-labelled training dataset.

Jayasree and Balan (2017) [41] presented fraud detection by using Decision Tree specifically for money laundering in credit card transaction. According to the research conducted by them, they take ID3 algorithm as a base and further improved it with their own algorithm, which named as Advanced Iterative Dichotomiser 3 (AID3). They integrated another method called Bitmap Index in helping getting the best attributes to split, instead of using Information Gain.

Bitmap Index works by provide pointers to the rows and columns in the transaction tables, then stores the row ID, column ID and the key values. The attributes are subsequently analysed using the given distinct key value on each transaction for a particular account. The indexing procedure is applied on a dummy banking database that Jayasree and Balan obtained. Bitmap Index in AID3 is carried out by using the "SELECT" query and logical AND operator. The query results are then use to construct the Decision Tree and apply it onto fraud detection application. Gaikwad, Deshmane, Somavanshi, Patil, and Badgujar (2014) [59] suggested Decision Tree Induction (ID3) algorithm as primary Machine Learning model for credit card fraud detection. Their approach started by investigating the dataset to understand each of the features and verify is there any anomalies or missing data. Then, training and testing are performed to test the performance of ID3 classification. After a series of testing and k-fold cross-validation, the final ID3 model is output and implemented on a bank's web server. Combined with the One-Time Password (OTP), ID3 is used to verify the transaction first, if the transaction is classified as legitimate, OTP will be sent out and the transaction will proceed.

Sahin and Duman (2011) [60] proposed and applied Decision Trees and Support Vector Machines as classification models on credit card fraud detection problem. Their study focused on experimenting and comparing the two approaches with a real credit card dataset provided by the Hong Kong National Bank with the required permissions. The past data in the dataset are used to form a data mart representing the card usage profiles of the customers. In their research works, they used multiple tree algorithms and different kernel function of SVMs to test out the model's performance in classifying fraudulent transaction. The chosen methods to build classifier models are C5.0, Classification and Regression Tree (CART) and

Chi-square Automatic Interaction Detector (CHAID) from decision tree methods and SVMs with kernels of polynomial, sigmoid, linear and radial basis functions. All these methods are used to develop classification models using three data sets that have a different ratio of fraudulent and legitimate transactions. In the first set, there is one normal transaction for each fraudulent one. In the second set, there are four normal transactions for each fraudulent one and there are nine normal ones for each fraudulent one in the third set.

Sahin and Duman also mentioned that due to the imbalanced of the dataset (978 fraudulent records and 22 million normal ones with a ratio of about 1:22500), they utilized stratified sampling in order to undersample the normal records so that the models have chance to learn the characteristics of both the normal and the fraudulent transaction.

2.6 Frequent Pattern Mining

Frequent Pattern Mining is one of the subfields in Data Mining in which it mines frequent sets of items and the interesting patterns in a given item set. Here, the interesting patterns can be a set of items that appears frequently in a database, or want to discover rare and negative patterns that have not seen in other existing itemset. Application of Frequent Pattern Mining are largely inter-related with Association Rule Learning, whereby it finds all the itemset patterns and then post-process them into rules in helping solving problems.

In this credit card fraud detection context, Apriori and FP-growth algorithm will be focused on as these two algorithms are often used to detect new fraud patterns [3] [4] and generate those patterns as new rules in order to prevent same method of fraud happening in future.

Seeja and Zareapoor (2014) [15] utilized Apriori as the main Frequent Pattern Mining technique in mining the patterns from a database. According to them, they described that instead of finding patterns for fraudster behaviour, Apriori algorithm is used in identify buying patterns for fraud and legal transaction. They had set the MinSup as 0.9 and constructed two patterns for each customer – legal pattern and fraud pattern.

After finding the legal and fraud patterns and stored in the database, the Fraud Detection System then traverses these pattern databases to detect frauds. Seeja and Zareapoor developed a matching algorithm, which traverses the pattern databases for matching the incoming transaction in detecting fraud. If a closer match is found with a legal pattern of the corresponding customer, then the matching algorithm returns "0" giving a green signal to the bank for allowing the transaction, else it returns as "1", giving an alarm to the bank that this might be a fraudulent transaction.

Tripathia, Nigamb and Edlaa (2017) [61] introduce Apriori algorithm as Association Rule Mining for fraud detection that caused by fraudulent websites. Their approach focused on Phishing activity whereby it intends to mislead the users to fraudulent websites, steal their sensitive information and use that information to perform the fraud transaction. To resolve the mentioned problem, they proposed Apriori algorithm to analyses the web access log which addresses the activities performed by the end clients and detecting a fraud sequence of repeated web URLs since fraudulent websites will have the same and frequent URLs [15] [47] [62]. After applying data mining techniques like pre-processing, transformation and frequent pattern generation, frequent patterns are generated. These frequent patterns in their model will be used to detect frauds on incoming request of the end clients.

In their experimental results and discussion, they described that the accuracy of the proposed implementation is significantly affected based on the size and the noise in the dataset. As the amount of dataset increased, the performance of the Apriori algorithm is dropped and this reflects the complexity of the algorithm itself. During the experiment, it is also observed that the number of items set and the amount of input transactions can have a negative effect onto the quality of Apriori algorithm in terms of error rate, memory consumption and search time.

Choudhary and Divya (2017) [62] proposed both Apriori and FP Tree algorithm to mine frequent itemset and discover interesting patterns that lead to credit card fraud. They described that by using frequent pattern matching schemes, they able to extract the frequent patterns from a dataset which they transformed it into a log file. From their experiment and discussion, they state although two algorithms generated the same patterns in the end, however, the time and memory consumed by Apriori algorithm is worse than FP Tree as it performed repeated dataset scans. Whereas FP Tree suffers from the complexity of its implementation since it performs recursive operations to generate a new tree every time, making the process complicated.

3. Discussion and Challenges

Fraud detection is an important part of the modern banking industry. This paper studied intelligent approaches of fraud detection, both statistical and computational. Though review of each algorithm and their performance in different credit card dataset, each technique was shown to be reasonably capable at detecting various forms of credit card fraud. In particular, the heavy-computational methods such as ANNs [55] and SVMs [26] are much more capable to

learn and adapt to new fraud patterns, which is highly effective to the evolving tactics of fraudsters.

Although other supervised Machine Learning algorithms like k-NN, Naïve Bayesian and Decision Tree are not effective in detecting new frauds and they required a more comprehensive re-train process onto the new data [30], but the nature of their algorithms is easy to be understood by non-AI experts. Particularly, Naïve Bayesian and Decision Tree provide probability and decision rules, which is helpful to the banking industry in understanding the underlying approaches of why and they can make use of the probability and rules to make better business decisions.

The same goes for Frequent Pattern Mining methods. Both Apriori and FP-growth are great at finding existing patterns from the credit card dataset but lacks the capability in detecting rare patterns from the new credit card transaction [29] [30]. In most of the fraud detection context, only mining frequent patterns is not sufficient to detect frauds and often it needs to be incorporate with classification approach [29] [30] [38]. Associative Classification approaches such as Classification Based on Associations (CBA) and Classification based on Multiple Association Rules (CMAR) are frequently used to enhance the detection of fraud and improve fraud rules accuracy that generated from Decision Tree Induction [47] [53] [38].

One of the main challenges while studying Literature Review and various types of Machine Learning algorithms is there is no standard, comprehensive or benchmark credit card dataset published publicly for comparative study. As credit card transaction dataset is inherently private property in the banking sector, therefore, to get a proper dataset and have a benchmark is very difficult. Lack of the standard transaction dataset and not having full information onto its fields or dimensions also makes the comparison of various Machine Learning techniques harder and more difficult to justify why certain approaches are better one another.

Moreover, there are also no specific rules for knowing how much data that is needed to produce a good Machine Learning model. This is often associated with bias and variance trade-offs [28] [30] [32]. For example, more data does not always mean a better model and it may lead to overfitting, which the model just “memorizing” the training data. On the other hand, if there is too little data, the model may become underfitting, which pays little attention and over-simplifying the training data. In other words, it just does the same thing over and over again regardless of what the data might be trying to tell it to do. Hence, it is important to assess the bias and variance trade-offs by adding more data or features that can help offset bias and variance problem.

4. Conclusion

In summary, the investigation of current practices in credit card fraud detection, as well as intelligent methods using Data Mining and Machine Learning algorithms are reviewed in this paper. Instead of using Rule-based system with simple pattern matchmaking for fraud detection, utilizing Supervised Machine Learning techniques provides a more novel way in catching fraudulent credit card transaction as it offers high accuracy and high confidence of risk score.

Besides that, it is important to repeatedly to examine the performance of Machine Learning model and fine-tune it to achieve the best scores in detecting fraudulent transactions. If done properly, the model would provide high confidence in distinguishing legitimate and fraudulent transaction while adapting over time to new or previously unseen fraud tactics [50] [52].

Through the studies of Data Mining and Machine Learning algorithms, it can be concluded that no one model works best for detecting fraudulent transaction. There is no single powerful algorithm in credit card fraud detection that outperforms all others algorithm. Each Data Mining and Machine Learning techniques have their strengths and weaknesses, whether it has high accuracy but slow training speed, or vice-versa. By using different algorithms onto the provided credit card transaction dataset, the models will have different classification performances in deciding whether a new transaction is fraudulent [31] [33] [38].

As a conclusion, fraud detection can become very complex when there is a need to consider numerous patterns in the data that can lead to fraudulent transactions. This also may require huge computational power to train and to accurately detect fraud. Although there is some limitation when using Data Mining and Machine Learning techniques to detect fraud, it still gives high dependability and adaptability as compared to Rule-based system, which only statically matches the patterns based on the transaction data.

References

- [1] R. Marmo, Data Mining for Fraud Detection System. Encyclopedia of Data Warehousing and Mining, 2nd ed, 2013, pp. 411-416.
- [2] C. Tyagi, P. Parwekar, P. Singh, and K. Natla, "Analysis of Credit Card Fraud Detection Techniques," Solid State Technology, vol. 63, no. 6, 2020, pp. 18057-18069.
- [3] C. Chee, J. Jaafar, I. Aziz, M. Hassan, and W. Yeoh, "Algorithms for frequent itemset mining: a literature review," Artificial Intelligence Review, vol. 52, 2019, pp. 2603-2621.
- [4] C. Ordonez, and K. Zhao, "Evaluating association rules and decision trees to predict multiple target attributes," Intelligent Data Analysis, vol. 15, no. 2, 2011, pp. 173-192.
- [5] D. Excell, Bayesian Inference – the Future of Online Fraud Protection. Computer Fraud & Security, 2nd ed., 2012, pp. 8-11.
- [6] J. Xu, A. Sung and Q. Liu, "Behaviour Mining for Fraud Detection," Journal of Research and Practice in Information Technology, vol. 39, no. 1, 2007, pp. 3-18.
- [7] C. Paasch, Credit Card Fraud Detection using Artificial Neural Networks tuned by Genetic Algorithms, 2014, doi:10.14711/thesis-b1023238
- [8] R. Porkess, and S. Mason, "Looking at Debit and Credit Card Fraud," Teaching Statistics, vol. 34, no. 3, 2011, pp. 87-91.
- [9] T. Sweer, Autoencoding Credit Card Fraud, Radboud University, 2018, Retrieved from https://www.cs.ru.nl/bachelors-theses/2018/Tom_Sweers___4584325___Autoencoding_credit_card_fraude.pdf
- [10] S. Yusuf, and D. Ekrem, "Detecting Credit Card Fraud by ANN and Logistic Regression," in Proceedings of the International Symposium on Innovations in Intelligent SysTems and Applications, 2011.
- [11] L. Mukhanov, "Using Bayesian Belief Networks for credit card fraud detection," In Proceedings of the Conference: Proceedings of the 26th International Conference on Artificial Intelligence and Applications, 2008, pp. 221-225.
- [12] S. Maes, K. Tuyls, B. Vanschoenwinkel, and B. Manderick, "Credit Card Fraud Detection Using Bayesian and Neural Networks," In Proceedings of the First International NAISO Congress on NEURO FUZZY THECNOLOGIES, 2002.
- [13] C. Milgo, "A Bayesian Classification Model for Fraud Detection over ATM Platforms," Journal of Computer Engineering, vol. 18, no. 4, pp. 26-32, 2016.
- [14] A. Desai, and D. Deshmukh, "Data mining techniques for Fraud Detection," International Journal of Computer Science and Information Technologies, vol. 3, pp. 1-4, 2013.
- [15] K. Seeja, and M. Zareapoor, "FraudMiner: A Novel Credit Card Fraud Detection Model Based on Frequent Itemset Mining," The Scientific World Journal, 2014, pp. 1-10.
- [16] Kevin Zakka. (n.d.), A Complete Guide to K-Nearest-Neighbours with Applications in Python and R, Retrieved from <https://kevinzakka.github.io/2016/07/13/k-nearest-neighbor>
- [17] I. Sutedja, Y. Heryadi, L. Wulandhari, and B. Abbas, "Recognizing debit card fraud transaction using CHAID and K-nearest neighbour: Indonesian Bank case," in Proceedings of the 11th International Conference on Knowledge, Information and Creativity Support Systems, 2016, pp. 1-5.
- [18] C. Sudha, and T. Raj, "Credit Card Fraud Detection in Internet Using K-nearest Neighbor Algorithm," International Journal of Computer Science, vol. 5, issue 11, pp. 22-30, 2017.
- [19] I. Rajak and K. Mathai, "Intelligent Fraudulent Detection System based SVM and Optimized by Danger Theory," in Proceedings of International Conference on Computer, Communication and Control, 2015, pp. 1-4.
- [20] j2kun, "Formulating the Support Vector Machine Optimization Problem," 2017 Retrieved from <https://jeremykun.com/2017/06/05/formulating-the-support-vector-machine-optimization-problem/>
- [21] C. Burges, "A tutorial on support vector machines for pattern recognition," Data Mining and Knowledge Discovery, vol. 2, no. 2, 1998, pp. 121-167.
- [22] Y. Sahin, and E. Duman, "Detecting credit card fraud by decision trees & support vector machines," in Proceeding of the International Multi Conference of Engineers & Computer Scientist, vol. I, 2011.
- [23] V. Dheepa, and R. Dhanapal, "Behavior Based Credit Card Fraud Detection Using Support Vector Machines," Journal on Soft Computing, vol. 2, no. 4, 2012, pp. 391-397.

- [24] Q. Lu, and C. Ju, "Research on Credit Card Fraud Detection Model Based on Class Weighted Support Vector Machine," *Journal of Convergence Information Technology*, vol. 6, no. 1, 2011, pp. 62-68.
- [25] D. Abdelhamid, S. Khaoula, and O. Atika, "Automatic Bank Fraud Detection Using Support Vector Machines," in *Proceedings of the International conference on Computing Technology and Information Management*, pp. 10-17, 2014.
- [26] R. Porkess, and S. Mason, "Looking at debit and credit card fraud," *Teaching Statistics*, vol. 34, no. 3, 2011, pp. 87-91.
- [27] L. Oghenekaro, and C. Ugwu, "A Novel Machine Learning Approach to Credit Card Fraud Detection," *International Journal of Computer Applications*, vol. 140, no. 5, 2016, pp.45-50.
- [28] N. Carneiro, G. Figueira, and M. Costa, "A data mining-based system for credit-card fraud detection in e-tail," *Decision Support Systems* 95, Elsevier B.V, 2017, pp. 91–101.
- [29] S. Ong, S. Sagadevan, and N. Malim, "Credit Card Fraud Detection Using Machine Learning As Data Mining Technique," *Journal of Telecommunication, Electronic and Computer Engineering*, vol. 10, no. 1-4, pp. 23-27, 2014.
- [30] CyberSource, "Annual Fraud Benchmark Report: A Balancing Act," North America Edition, 2016.
- [31] G. James, D. Witten, T. Hastie, and R. Tibshirani, "An Introduction to Statistical Learning," Springer, 2013, pp. 204.
- [32] A. Bănărescu, "Detecting and Preventing Fraud with Data Analytics," *Procedia Economics and Finance*, vol. 32, 2015, pp. 1827-1836.
- [33] B. Zolfaghari, K. Bibak, T. Koshiba, H. Nemati, and P. Mitra, "Statistical trend analysis of physically unclonable functions: An approach via text mining," *CRC Press*, 2021, pp. 55-74.
- [34] W. Loh, "Classification and regression trees," *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, vol. 1, no. 1, 2011, pp. 14-23.
- [35] J. Quinlan, "Improved Use of Continuous Attributes in C4.5," *Journal of Artificial Intelligence Research*, vol. 4, 1996, pp. 77-90.
- [36] B. Hssina, A. Merbouha, H. Ezzikouri, and M. Erritali, "A comparative study of decision tree ID3 and C4.5," *International Journal of Advanced Computer Science and Applications*, 4(2), 2014, pp. 13-19.
- [37] B. Gupta, A. Rawat, A. Jain, and A. Arora, "Analysis of Various Decision Tree Algorithms for Classification in Data Mining," *International Journal of Computer Applications*, Vol. 163, no. 8, 2017, pp. 0975 – 8887.
- [38] S. Priyanka, "Comparative Study ID3, CART and C4.5 Decision Tree Algorithm: A Survey," *International Journal of Advanced Information Science and Technology*, vol. 27, no. 27, 2014, pp. 97-103.
- [39] Y. Sahin, S. Bulkan, and E. Duman, "A cost-sensitive decision tree approach for fraud detection," *Expert Systems with Applications*, vol. 40, no. 15, 2013, pp. 5916-5923.
- [40] V. Jayasree, and R. Balan, "Money laundering regulatory risk evaluation using Bitmap Index-based Decision Tree," *Journal of the Association of Arab Universities for Basic and Applied Sciences*, vol. 23, no. 1, 2017, pp. 96-102.
- [41] T. Minegishi, and A. Niimi, "Proposal of Credit Card Fraudulent Use Detection by Online-type Decision Tree Construction and Verification of Generality," *International Journal for Information Security Research*, vol. 3, no. 1, 2013, pp. 229-235.
- [42] O. Aodha, and G. Brostow, "Revisiting Example Dependent Cost-Sensitive Learning with Decision Trees," in *Proceedings of the 2013 IEEE International Conference on Computer Vision*, 2013, pp. 193-200.
- [43] I. Monedero, F. Biscarri, C. León, J. Guerrero, J. Biscarri, and R. Millán, "Detection of frauds and other non-technical losses in a power utility using Pearson coefficient, Bayesian networks and decision trees," *International Journal of Electrical Power & Energy Systems*, vol. 34, no. 1, 2012, pp. 90-98.
- [44] S. Agarwal, "Data Mining: Data Mining Concepts and Techniques," in *Proceedings of the 2013 International Conference on Machine Intelligence and Research Advancement*, 2013, pp. 203-207.
- [45] R. Agrawal, and R. Srikant, "Fast algorithms for mining association rules in large databases," *Research Report RJ 9839*, 1994, IBM Almaden Research Center, San Jose, California.
- [46] C. Ordóñez, and K. Zhao, "Evaluating association rules and decision trees to predict multiple target attributes," *Intelligent Data Analysis*, vol. 15, no. 2, 2011, pp. 173–192.
- [47] S. Nasreen, M. Azam, K. Shehzad, U. Naem, and M. Ghazanfar, "Frequent Pattern Mining Algorithms for Finding Associated Frequent Patterns for Data Streams: A Survey," *Procedia Computer Science*, vol. 37, 2014, pp. 109-116.
- [48] D. Excell, "Bayesian inference – the future of online fraud protection," *Computer Fraud & Security*, vol. 2, 2012, pp. 8-11.
- [49] J. Xu, A. Sung, and Q. Liu, "Behaviour Mining for Fraud Detection," *Journal of Research and Practice in Information Technology*, vol. 39, no. 1, 2007, pp. 3-18.
- [50] K. Hu, Y. Lu, L. Zhou, and C. Shi, "Integrating Classification and Association Rule Mining: A Concept Lattice Framework," *Lecture Notes in Computer Science New Directions in Rough Sets, Data Mining, and Granular-Soft Computing*, 2001, pp. 443-447.
- [51] B. Liu, Y. Ma, and C. Wong, "Classification Using Association Rules: Weaknesses and Enhancements," *Data Mining for Scientific and Engineering Applications Massive Computing*, 2001, pp. 591-605.
- [52] F. Thabtah, "A review of associative classification mining," *The Knowledge Engineering Review*, vol. 22, no. 1, 2007, pp. 37-65.
- [53] D. Montague, *Essentials of online payment security and fraud prevention*, Wiley, 2011, pp. 183-189.
- [54] A. Serrano, J. Costa, C. Cardonha, A. Fernandes, and R. Júnior, "Neural Network Predictor for Fraud Detection: A Study Case for the Federal Patrimony Department," in *Proceedings of the Seventh International Conference on Forensic Computer Science*, 2012, pp. 61-66.
- [55] S. Viaene, R. Derrig, and G. Dedene, "A case study of applying boosting naive bayes to claim fraud diagnosis," *IEEE Transactions on Knowledge and Data Engineering*, vol. 16, no. 5, 2004, pp. 612-620.
- [56] S. Kiran, J. Guru, R. Kumar, N. Kumar, D. Katariya, and M. Sharma, "Credit card fraud detection using Naïve Bayes model based and KNN classifier," *International Journal of Advance Research, Ideas and Innovations in Technology*, vol. 4, 2018, pp. 44-47.
- [57] R. Sallehuddin, S. Ibrahim, A. Zain, and A. Elmi, "Detecting SIM Box Fraud by Using Support Vector Machine and Artificial Neural Network," *Jurnal Teknologi*, vol. 74, no. 1, 2015, pp. 137-149.
- [58] R. Banerjee, G. Bourla, S. Chen, M. Kashyap, S. Purohit, and J. Battipaglia, "Comparative Analysis of Machine Learning Algorithms through Credit Card Fraud Detection," in *Proceedings of the 2018 IEEE MIT Undergraduate Research Technology Conference*, 2018, pp. 1-4.
- [59] J. Gaikwad, A. Deshmane, H. Somavanshi, S. Patil, and R. Badgujar, "Credit Card Fraud Detection using Decision Tree

Induction Algorithm,” International Journal of Innovative Technology and Exploring Engineering (IJITEE), vol. 4, no. 6, 2014, pp. 66-69.

- [60] Y. Sahin., and E. Duman, “Detecting Credit Card Fraud by Decision Trees and Support Vector Machines,” in Proceedings of the International of MultiConference of Enginners and Computer Scientists, 2011, pp. 442-447.
- [61] D. Tripathi, B. Nigam, and D. Edla, “A Novel Web Fraud Detection Technique using Association Rule Mining,” Procedia Computer Science, vol. 115, 2017, pp. 274-281.
- [62] V. Choudhary, and E. Divya, “Credit Card Fraud Detection using Frequent Pattern Mining using FP- Tree And Apriori Growth,” International Journal of Advance Technology and Innovation Research, vol. 09, no. 13, 2017, pp. 2370-2373.
- [63] R. Schapire, “Explaining AdaBoost,” Empirical Inference, 2013, pp. 37-52.
- [64] M. Bansal, and Suman, “Credit Card Fraud Detection Using Self Organised Map,” International Journal of Information & Computation Technology, vol. 4, No. 13, 2014, pp. 1343-1348.
- [65] H. Naik, “Credit Card Fraud Detection for Online Banking Transactions,” International Journal for Research in Applied Science and Engineering Technology, vol. 6, no. 4, 2018, pp. 4573-4577.
- [66] N. Malini, and M. Pushpa, “Investigation of Credit Card Fraud Recognition Techniques based on KNN and HMM,” in Proceedings of the International Conference on Communication, Computing and Information Technology, 2018, pp. 9-13.
- [67] M. Franzese, and A. Iuliano, “Hidden Markov Models,” Encyclopedia of Bioinformatics and Computational Biology, vol. 1, 2019, pp. 753-762.
- [68] M. Pietrzykowski, and W. Sałabun, “Applications of Hidden Markov Model: state-of-the-art,” International Journal of Computer Technology & Applications, vol. 5, no. 4, 2014, pp. 1384-1391.
- [69] B. Baesens, S. Höppner, and T. Verdonck, “Data engineering for fraud detection,” Decision Support Systems, 2021, article 113492,
- [70] X. Zhang, Y. Han, W. Xu, and Q. Wang, “HOB A: A novel feature engineering methodology for credit card fraud detection with a deep learning architecture,” Information Sciences, vol. 557, no. 10, 2021, pp. 302-316.
- [71] P. Craja, A. Kim, and S. Lessmann, “Deep learning for detecting financial statement fraud,” Decision Support Systems, vol. 139, 2020, article 113421.



Lam Hong Lee serves as an associate professor at Quest International University, Malaysia. He received a Bachelor of Computer Science from Universiti Putra Malaysia in 2004, and a Ph.D. in Computer Science from the University of Nottingham in 2009. He joined Universiti Tunku Abdul Rahman, Malaysia in 2009 as an assistant professor, before joining Quest International University in 2013. His research interest lies in improving intelligent systems using enhanced machine learning techniques. Besides this, he is also investigating the implementation of data mining, pattern recognition and machine learning techniques in various kinds of IR4.0 applications.



Yee-Wai Sim serves as a professor at Quest International University, Malaysia. He is also the Dean of Faculty of Computing and Engineering at his current university. He received a Bachelor of Engineering in Electromechanical from University of Southampton (UK) in 2000, and a Ph.D. from the University of Southampton (UK) in 2004. He has also founded a company in providing child locator solution. His research interests are in the areas of agent technology, machine learning, and drone technology.



Kha Shing Lim is currently working as an Android Software Engineer at Fave Asia, Malaysia. He completed his Bachelor's Degree of Information Technology from Quest International University, Malaysia. During his internship, he served as a Data Integrator at an oil and gas company named PETRONAS at Malaysia.