

EVALUATING THE EFFICIENCY OF THE K-NEAREST NEIGHBORS (KNN) ALGORITHM FOR CREDIT CARD FRAUD DETECTION

Rosary Nancy¹

¹Dept. of CS, Fatima College, India.

ABSTRACT

Credit card fraud is a major and rising threat, particularly with the advent of e-commerce and online transactions in the modern day. The implications of identity theft and financial losses caused by such malicious activities might affect millions of individuals worldwide, posing a huge danger to the banking industry. The efficiency of fraud detection in credit card transactions is significantly influenced by the data set measurement method, variable selection, and detection algorithms employed. Information extraction is crucial in detecting online payment fraud, as are the strategies employed to address this issue. This study assesses the performance of the K-Nearest Neighbor algorithm on heavily distorted credit card fraud data. These procedures are evaluated using measures including accuracy, sensitivity, precision, and specificity. The results show K-Nearest Neighbor have optimum accuracy percentages of 96.91%, respectively.

1. INTRODUCTION

Financial fraud is a continuous and growing issue that has far-reaching consequences for the financial services, industry, and government sectors. Fraud is described as dishonest techniques used to obtain monetary benefit, and it is a rising threat to organizations, corporate entities, and governments, with serious implications. The increased dependence on internet technologies for both physical and digital credit card transactions has resulted in a significant increase in credit card fraud.

Detecting credit card fraud involves analyzing spending behavior, and various methods, such as support vector machines, genetic algorithms, decision trees, artificial neural networks, and naive Bayes, have been employed for fraud detection. Credit card companies are striving to predict the authenticity of purchases by evaluating disparities in purchasing locations, transaction amounts, and user purchase history. However, the rise in credit card fraud cases emphasizes the need for optimized algorithmic solutions for credit card companies.

Credit card fraud detection faces challenges such as dynamic fraudulent behavior patterns that resemble authentic operations, scarce and imbalanced credit card transaction datasets, and the influence of testing approaches, variable selection, and identification techniques on fraud detection efficiency. Additionally, the constantly changing nature of data over time poses a challenge in distinguishing between past genuine operations and potential current or future fraudulent activities. This paper aims to conduct a comparative analysis of fraudulent activity identification on credit cards using k-nearest neighbor techniques.

2. LITERATURE REVIEW

Fraud detection is traditionally seen as a problem of data mining classification, with the aim of correctly classifying credit card transactions as legitimate or fraudulent. Categorizing of card operations is primarily a concern of Boolean categorization. Credit card transactions in this regard is either legit transactions or a fraudulent transaction.

A. Data Mining

Data mining refers to digging into or mining the data in different ways to identify patterns and get more insights into them. It involves analysing the discovered patterns to see how they can be used effectively. Classification in data mining is a common technique that separates data points into different classes. It allows you to organize data sets of all sorts, including complex and large datasets as well as small and simple ones. Binary classification is the simplest classification issue. In binary classification, the target attribute might be either high risk or typically safe for fraud. Binary classifiers are the most suited approach for detecting credit card fraud.

B. Scam on Credit Card

Card fraud can be categorized into two types: internal theft for card identity and external fraud. A broader classification has also been undertaken, dividing it into several categories such as Traditional card-related scams, vendor-related, and cyber fraud. The total losses from fraudulent transactions involving financial institutions and organizations globally are reported in. The study concluded that, with an increase of approximately \$2.6 billion from the previous year's reported losses, the total exceeded \$16 billion in 2015. This implies that every \$100 involved 6.1 cents attributed to fraud.

C. Machine Learning

Machine Learning, a branch of Artificial Intelligence, involves training computers to recognize patterns within extensive datasets and enhance these patterns autonomously, without requiring human intervention. The training process commences with a basic machine-learning algorithm that utilizes training data to analyse the relationships among different elements and an objective value. Once trained, the model can then be employed to predict unknown target values for other data instances. Depending on whether the training data is labelled, machine learning can be categorized as supervised or unsupervised. Supervised learning focuses on establishing a relationship between an input value and an output value, enabling the prediction of additional output values when presented with more input. Classification issues categorize output (e.g., fraud and non-fraud), while regression issues output a specific value (e.g., height). Machine learning algorithms that analyse the input-output relationship without producing an output are termed unsupervised, as the training data are neither labelled nor classified.

D. Fraud on Credit Card

As card payments become the most prevalent form of transactions, both online and offline, the incidence of credit card fraud is rapidly increasing. Traditional manual methods for identifying fraudulent transactions are time-consuming and prone to inaccuracies, making them less effective as the volume of data grows. To address this, corporate organizations are turning to sophisticated techniques, particularly those based on artificial learning. Methods for predicting and detecting fraud generally fall into two broad categories: supervised and unsupervised. Supervised fraud detection methods project designs based on features of both deceptive and legitimate operations, allowing the classification of new transactions as fraudulent or legitimate. In contrast, unsupervised fraud detection identifies potential instances of fraudulent transactions by detecting outliers.

E. Feature Selection

The primary basis for fraud detection on credit cards involves the analysis of the cardholder's spending behaviour. This spending behaviour is evaluated through the identification of relevant features that capture unique patterns in credit card usage, as genuine and fraudulent transactions tend to exhibit distinct profiles. These features are derived from a combination of historical transaction data, with all features categorized into fundamental input types, including operational data, geographic data, merchant type data, and time-based data on quantities and transaction periods. Variables falling under the category of transaction statistics depict an overall pattern of card utilization, those under geographic data illustrate the card's spending habits considering different geographical regions, and variables under merchant data showcase the card's usage across various merchant categories.

3. METHODOLOGY

This section explains the procedure and methods used in the experiment the dataset used to train the model and the techniques used in the study for machine learning, K-Nearest Neighbor Algorithm. The learning phase is where the classifier's system are created and supplied with the extracted information, the test is evaluated utilizing confusion matrix measurement rates.

A. Dataset

Dataset emerges from Kaggle Machine Learning platform. This dataset presents 3075 transactions with 12 features of transactions in CSV file. Due to confidentiality issues, the features details and background information cannot be presented. The features contains the average amount of transaction per day, transaction amount, if declined or not, foreign transaction or not, if it's of high risk, and six month average balance in the dataset. This fraudulent feature is perhaps the label for the Boolean evaluation and it contains precedence Y in case of illegal operations (fraud) and N for legal operations (not fraud).

B. Data Preprocessing

Data preprocessing involves handling the dataset, with a significant portion of the feature columns being categorical data. To facilitate processing, conversion into integers was deemed necessary. In instances where the class of the feature is labelled as 'Y,' a binary conversion to 1 is performed, while for a feature class labelled 'N,' the conversion is to 0. The dataset is then divided, with 80% allocated for the training set and 20% for the test set. Given the imbalanced ratio between non-fraudulent and fraudulent cases observed in the dataset predictions from the initial training set, which is not evenly distributed, a specific range for the fraudulent class is selected for training data (from 1 to 400) and testing data (from 401 to 450). For the non-fraudulent class, training data is chosen from 451 to 2800, while testing data is selected from 2801 to 3075.

C. K-Nearest Neighbor (KNN)

K-nearest neighbor algorithm is a classification algorithm that predicts the attributes of an informational point to other points based on its relative position. Its classification is based on similarity measures such as Euclidean distance, Manhattan distance measure. It is assumed that the data point in the training set that has the shortest Euclidean distance to the test point has the same unknown attribute as the test point. The Euclidean distance measure for the KNN classifier is used in this study. The range between Euclidean (EC)'s two-point vectors (x_1, x_2) is determined by:

$$EC = \sqrt{\sum (x_1 - x_2)^2} \quad k=1, 2, \dots, n$$

Manhattan distance measure between two points (x_i, y_i) and (x_n, y_n) is a metric in which the distance between two points is the absolute difference of their Cartesian coordinate.

$$M = (x_i - x_n) + (y_i - y_n)$$

D. Evaluation and Result

To evaluate this machine learning models we considered two different method namely;

(1) Classification accuracy: which is the ratio of number of correct predictions to the number of input sample. But this is very effective only if there are equal number of samples in each class.

$$\text{Accuracy} = \frac{\text{number of correct predictions}}{\text{total number of predicted made}}$$

(2) Confusion Matrix: this gives a matrix as output and describe the complete performance of the model. Four essential measurements are utilized in evaluating the analyses, to be specific True Positive Ratio (TPR), True Negative Ratio (TNR), False Positive Ratio (FPR) and False Negative Ratio (FNR) rates metric individually. In which true positive, true negative, false positive and false negative are the quantity characterized by true positive, false positive, true negative, and false negative experiments, thus p and n are the absolute values of positive and negative class cases being tested. True positives are classes predicted to be positive and are, true negative classes are predicted as negative but are negative. False positive are classes predicted to be positive and are negative, false negatives are classes predicted to be negative but that are positive.

		Actual Values	
		Positive (1)	Negative (0)
Predicted Values	Positive (1)	TP	FP
	Negative (0)	FN	TN

Accuracy is the ratio of the sum of true positive and true negative to the sum of all the predicted samples

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

$$\text{Sensitivity (recall)} = \frac{TP}{TP+FN}$$

Sensitivity which is also called recall is the measure of the ratio of true positive predictions to the sum of true positive and false negative.

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

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

Specificity is the measure of the ration of true negative to the sum of true negative and false positive

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

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

Precision is the ratio of the number of true positives to the sum of true positive and false positive. It can be said to be the measure of the quality of the positive feedback data.

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

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

To evaluate the implementation of the classifiers, specificity, precision, accuracy, and sensitivity are used.

Metrics	K-nearest neighbor
Accuracy	96.91
Sensitivity	89.36
Specificity	98.19
Precision	89.36

4. CONCLUSION

classifiers models are being developed in this study based on K-Nearest Neighbor 80% of the dataset is used for validation and testing. Precision, Sensitive, Specificity, Accuracy are used to assess performance. However, an unrealistic expectation is the presence of a balanced training and testing dataset of the same distribution. Based on this exploration, a credit card organization ought to consider executing a KNN algorithm that investigates the buy time to distinguish whether a credit card transaction is fraud.

5. REFERENCES

- [1] Maes, S., Tuyls, K., Vanschoenwinkel, B. and Manderick, B., (2002). Credit card fraud detection using Bayesian and neural networks Proceeding International NAISO Congress on Neuro Fuzzy Technologies.
- [2] Ogwueleka, F. N., (2011). Data Mining Application in Credit Card Fraud Detection System, Journal of Engineering Science and Technology.
- [3] Bhattacharyya, S., Jha, S., Tharakunnel, K., & Westland, J. C. (2011). Data mining for credit card fraud: A comparative study. Decision Support Systems.
- [4] Shen, A., Tong, R., & Deng, Y. (2007). Application of classification models on credit card fraud detection. In Service Systems and Service Management, 2007 .
- [5] Data Analytics vs Data Science: Two Separate, but Interconnected Disciplines, Data Scientist Insights, 28-Apr-2018.
- [6] Bhattacharyya, S., Jha, S., Tharakunnel, K., & Westland, J. C. (2011). Data mining for credit card fraud: A comparative study. Decision Support Systems, 50(3),