



FRAUD DETECTION USING DATA ANALYTICS IN THE CUSTOMS CONTEXT: A CASE STUDY OF UNDER INVOICING IMPORTATION FRAUD IN Indonesia

Siti Arifa'atus Sa'adah ^a, Arief Hartanto ^b

^a Directorate General Customs and Excise, Indonesia. Email: St.arifa92@gmail.com (penulis berkorespondensi)

^b Directorate General Customs and Excise, Indonesia. Email: ariefhartn@gmail.com

ARTICLE INFORMATION

ARTICLE HISTORY

Received

11 September 2023

Accepted to be published

24 November 2023

KEYWORDS:

Under-invoicing

Fraud Detection

Machine Learning

Customs Value

Importation

WCO

ABSTRACT

Fraud detection menjadi perhatian besar bagi semua instansi pemerintah. Di ranah kepabeanan, *fraud detection* sangat dibutuhkan untuk memastikan tidak adanya kebocoran penerimaan negara yang salah satunya disebabkan oleh *fraud* pada importasi *under invoicing*. Implementasi data analitik telah digunakan dalam banyak penelitian untuk menangani masalah pada *big data* dan memberikan solusi. Penelitian ini bertujuan untuk menjelaskan bagaimana data analitik dapat diimplementasikan untuk mendeteksi *fraud* pada importasi *under invoicing*. Beberapa variabel yang disertakan dalam penelitian ini, termasuk variabel yang menunjukkan tingkat risiko importir, komoditas, pemasok, dan negara eksportir. Penelitian ini membandingkan berbagai model *machine learning* termasuk *Logistic Regression*, *Decision Trees*, *Random Forest*, *Extreme Gradient Boost*, *Artificial Neural Networks*, *Gaussian NB*, dan *K-nearest Neighbors*. Untuk mengevaluasi model-model tersebut, penelitian ini mengukur kinerja model dengan membandingkan nilai akurasi, nilai presisi, dan nilai *log loss*. Hasilnya menunjukkan bahwa *Xtreme Gradient Boost* memiliki performa terbaik dalam mendeteksi *fraud under invoicing* dengan nilai akurasi sebesar 63%, nilai presisi sebesar 63%, dan nilai *log loss* sebesar 62%. Sejauh yang kami ketahui, ini adalah penelitian pertama yang membandingkan sejumlah model *machine learning* untuk mendeteksi *fraud under invoicing*. Hasil dari penelitian ini akan membantu para petugas bea cukai dalam proses pemeriksaan impor dengan memberikan peringatan dini terhadap transaksi *under-invoicing*. Hal ini dapat menghasilkan pemeriksaan yang lebih efektif dan efisien, sehingga bea cukai dapat menjalankan fungsi pelayanan dan pemeriksaan dengan baik, meskipun dengan sumber daya yang terbatas.

Fraud detection is a big concern for all the government agencies. In customs areas, fraud detection is needed to ensure that there is no leakage in state revenue, one of which is caused by the under invoicing importation fraud. The data analytic implementations have been used in many studies to handle problems in big data and give solutions. This study aims to explain how data analytics can be implemented to detect the under invoicing importation fraud. Several variables were included in this study, including the variables that show the risk level of importers, commodities, suppliers, and the exporter countries. This study compared various machine learning models including Logistic Regression, Decision Trees, Random Forest, Extreme Gradient Boost, Artificial Neural Networks, Gaussian NB, and K-nearest Neighbors. To evaluate the models, this study measures the performance of the models by comparing accuracy score, precision score and log loss score. The result shows that the Xtreme Gradient Boost performs best in detecting under invoicing fraud with accuracy score at 63%, precision score at 63% and log loss score at 62%. As far as we know, this has been the first work to compare a number of machine learning models to create under invoicing fraud detection. The results of this study will assist examiners in the import clearance process by providing an early warning of the under-invoicing transaction. It can lead to more effective and efficient examination, so that customs agencies can perform well in their service and inspection functions, despite the limited resources.

1. INTRODUCTION

1.1. Background of Study

The industrial revolution 4.0 brings many new opportunities for commercial trade, including international trade. Currently, the market is

increasingly borderless, which has led to an increase in the intensity of international trade. As illustration based on import data at the Directorate General of Customs and Excise, the general importation using

document BC 2.0 in 2022 reached more than 4.5 million documents with a foreign exchange value of more than USD 196 billion. This importation value does not include importations using other documents such as BC 2.3 (a document for bounded zone), Consignment Notes (CN), or other importation documents. This importation value is projected to increase in the coming years, given the historical upward importation trend in previous years. This fact is certainly a challenge for the Directorate General of Customs and Excise as the unit in charge of customs activities. The Directorate General of Customs and Excise is required to perform well in providing both services and inspections of export and import activities.

In conducting service and inspection activities, the Directorate General of Customs and Excise carries out risk management for the efficiency and effectiveness of task implementation, one of which is by implementing a channeling system for the clearance process of export and import activities. Currently, the channeling system consists of red, yellow, and green lanes. The difference between the three lanes is in the type of inspection carried out in the clearance process. The red lane applies physical checks and document checks before the release of goods. In the yellow lane, no physical inspection is carried out, but document inspection is carried out before the release of imported goods. Meanwhile, in the green lane, no physical inspection is carried out and document inspection is carried out after the release of imported goods.

In 2022, green channel imports dominated, with more than 92% on average, while red channel imports were around 7%. This is certainly in line with the government's efforts to reduce dwelling time and improve the ease of movement of goods throughout the port by carrying out the inspection after releasing the imported goods. This condition is certainly a challenge for the Directorate General of Customs and Excise in conducting supervision or inspection. The Directorate General of Customs and Excise is required to be able to carry out inspections more effectively and efficiently, one of which is related to examining the reported customs value.

1.2. Under-invoicing importation

Indonesia applies the self-assessment system in the import document, which means that importers or parties acting on behalf of importers declare the import document, pay import taxes and report to the customs authorities. Otherwise, the customs agencies will carry out the inspection of goods and import documents, including the amount of import taxes that should be paid. In conducting the inspection, the customs agencies applied risk management using channeling procedures to ensure that there is no congestion in port due to the inspection process.

Relating to the customs value that should be declared by importers, WCO Released Customs Valuation and Transfer Pricing Guidelines [13]. WCO

stated that the primary basis for customs value is transaction value, that's the price actually paid or should be paid. In addition, WCO also stated that there are other methods that can be applied to determine the customs value if the transaction value can't be applied.

During the inspection, the customs officer will examine and compare the customs value declared by importers in the import document with the customs value based on the WCO Guideline. In Indonesia, the customs officers are limited to conduct the inspection until 30 days after releasing goods. In so many cases, *under invoicing* fraud could be found through deep examination. *Under invoicing* fraud is the condition when importers declare the customs value of imported goods at a lower amount than it should be, in order to cut off the import duty or other taxes. In the current Indonesian trading volume, it's impossible to perform a deep examination for each transaction. As a solution, customs officers can assign a priority list of transactions that should be examined in detail, which leads to the transaction that indicates under-invoicing fraud.

At present, there are no tools that can provide alerts or warnings on *under invoicing* fraud. Therefore, the author argues that there is a need to build tools or models that can assist the customs valuation and provide alerts or warnings on certain transactions indicated under invoicing fraud. The purpose of this study is to explain how data analytics can be implemented to provide an early warning system that can detect under-invoicing fraud at an early stage.

2. LITERATURE REVIEW

The previous study stated that conventional techniques to detect fraud will be imprecise, costly, and time consuming [1]. So, they conclude in the previous study that the researchers usually use support vector machine (SVM) and Artificial Neural Network (ANN) algorithms to detect fraud.

The study that developed by [11] shows that good results in fraud detection were performed by machine learning methods like classification trees and logistic regression. Furthermore, in advancing standardized analytic models, the researchers usually include the feature engineering and instance engineering in the model.

Previous research [10] explains how to apply the ANN (Artificial Neural Networks) model to detect fraud in income tax. The results of the study successfully showed what factors affect the occurrence of fraud in income tax. The results also show a high level of accuracy with an accuracy value of 92% and a precision level of 85%. Based on this research, the author is interested in making this project, namely the application of data analytics, especially ANN (Artificial Neural Networks) to detect fraud in under invoicing importation.

This research is a development of previous research [10], with the novelty of applying and comparing analytical data modelling for fraud

detection in under invoicing importation, which may have been developed, but has not yet produced sufficient level of accuracy. In addition, the novelty of the model to be used is that the author will combine several risk factors, not only commodity risk or imported goods, but also supplier risk, importer risk, and country risk, that greatly affect the behavior patterns of importers in reporting the importation documents.

3. RESEARCH METHODOLOGY

The authors used descriptive statistical analysis to describe the general information conditions of the data. The use of descriptive statistics provides a more detailed picture of the dataset used. In addition, descriptive statistics also generate more understanding that is needed to conduct further analysis. To make an initial identification of the relationship between the variables, the authors use Pearson correlation analysis to determine whether there is a correlation between these variables and how significant the correlation is.

In this research, we considered four more steps after conducting business understanding: (i) data understanding, (ii) data preparation, (iii) modelling, and (iv) evaluation.

3.1. Data Understanding

The authors perform the data understanding to identify the variables that can be involved in predicting the under invoicing fraud. Based on the business understanding, the authors conclude that the variables affecting the under invoicing fraud are under invoicing flag, facility types, incoterm, import tax rate, risk level of importers, risk level of supplier, risk level of exporter country, and risk level of commodities. To run the model, the under-invoicing flag is a variable target.

3.2. Data Preparation

Data preparation aim to prepare the data to be ready used in modelling and reduce the biases. In this paper, data preparation conducted in several steps:

- (i) Cleansing data to ensure that the data is clean from empty fields or errors;
- (ii) Labeling of under invoicing transactions by providing the description of the documents that showed under invoicing fraud.
- (iii) Clustering the risk level of the supplier country, supplier, commodity, and importer into three classes: low risk, medium risk, and high-risk category. This clustering uses the difference of customs value shortage between the customs value declared and customs value based on examination by customs officers. Furthermore, clustering proceeds according to rule-based criteria as follows:

Table 1 Clustering ruled based

Customs value shortage	Risk Category
Less than IDR 100 million	Low
IDR 100 million – IDR 1 billion	Medium
More than IDR 1 billion	High

- (iv) Balancing the data to avoid bias. The total data in this study contains 921,426 rows that consist of 159,359 with under invoicing fraud and 762,067 rows that shows the correct customs value. From this data, the author performs balancing by down-sampling the majority of data. So, the data is ready to be used in modelling to predict the pattern of under invoicing fraud.

3.3. Modelling

In this paper, the author comparing several machine learning models to detect under invoicing fraud. It's intended to generate better understanding and comprehensive exploration of the models that led to finding the best fit model. For the execution, the author used python and google collab. The following machine learning models that used in this paper:

- (i) Artificial Neural Network (ANN)
Artificial Neural Networks works techniques similar to how the human brain works to identify hidden patterns or relationships in data. The model used neurons to pass the information or insight during the training data process and networks take data and train themselves to recognize the hidden patterns in this data [7].
- (ii) Logistic Regression
Logistic regression is a multivariable method for modeling the relationship between multiple independent variables and a categorical dependent variable. Logistic regression didn't use a linear relationship between dependent and independent variables, but this model performs and generates the relationship between the logit of the outcome and the predictor values [3].
- (iii) Decision tree
Decision tree is a successive model that combines a series of the basic test where a numeric feature is compared to a threshold value in each test.

In data mining, classification algorithms are useful to handle a big volume of data and information. It can be used to make assumptions, to classify the insight on the training datasets, and then classify the result based on the pattern found in the training dataset. [4].

The nodes and branches are composed of each tree. Each node shows features in a category to be classified and generates a value that can be taken by the node [6].

(iv) Random Forest

A random forest classifier is an machine learning algorithm that perform multiple decision trees, using a randomly selected subset of training dataset and variables [2].

The random forest model uses many individual trees to predict the classification. The individual trees perform the model through bootstrap data rather than the original data. This technique is proven to decrease overfitting problems [12].

(v) Xtreme gradient boost

Xtreme gradient boost is one of the implementations of gradient boosting machines (GBM) that perform one of the best performing supervised machine learning algorithms [8].

In the prediction process, the model will deduct the last residual and then build a new tree model in the negative gradient direction. then the previous decision tree will affect the sample input when constructing the next tree model.

Xtreme gradient boost also minimizes the volatility of model prediction and improves the overfitting phenomenon of the model by gaining a regular term to objective function [9].

(vi) Gaussian NB

The Naive Bayes algorithm, as a probabilistic classification model, computes a set of probabilities by counting up the frequency and combination of values in a particular data set. The probability of certain features in the data arises as members in a probability sequence, and it is gained by computing the frequency of each feature value in the class from the training data set. The training data set is a subset of the data that used to train the model and find the pattern of the classification. The training process uses known values to predict unknown values [14].

(vii) K-nearest Neighbor

K-nearest Neighbor techniques are a class of multivariate, nonparametric approaches to continuous or categorical prediction. The auxiliary variables are appointed feature variables and the space explained by the feature variables is appointed the feature space. To determine the reference set, the model used a set of sample population units for which observations of both response and feature variables are available. And the set of population units for which predictions of response variables are desired is appointed the target set. Having

complete sets of observations for all feature variables is required to ensure all population units have both the reference and target sets [5].

3.4. Evaluation

To evaluate the modelling results, the author compares the results (output) of the modelling using parameters such as precision score, accuracy score and log loss.

Precision score is the ratio of positive correct predictions compared to the overall positive predicted results using the following formula:

$$\text{Precision Score} = \frac{\text{True Positive}}{(\text{True Positive} + \text{False Positive})}$$

while accuracy score is the ratio of correct predictions (positive and negative) of all data, by using the formula:

$$\text{Accuracy Score} = \frac{\text{Number of Correct Predictions}}{\text{Total Number of Predictions}}$$

whereas, Log loss measures the performance of the modelling, where the output is a probability of a value between 0 and 1. As the predicted probability gets further away from the actual label, the Log loss increases. A perfect model will have a Log loss of 0, so the smaller the Log loss value, the better the result. The formula for calculating Log loss is as follows:

$$H(P, Q) = -\frac{1}{N} \sum_{i=1}^N y_i \cdot \log(P(y_i)) + (1 - y_i) \cdot \log(1 - P(y_i))$$

4. RESULTS AND FINDINGS

4.1. Findings

In this section we look at the dataset that we used, present and discuss the results, and describe what these results mean in terms of under-invoicing fraud detection.

4.1.1. Data Description and Preparation

In this paper, we used Indonesian importation data and examination result data during 2022 to build the pattern and predict the under invoicing fraud. The examination data used to verify the under invoicing flag through the examination conducted by customs officers. The total data is 921,426 rows, consisting of 159,359 rows showing under-invoicing fraud and 762,067 rows not showing under-invoicing fraud. We balance the data by down-sampling most of the data because of the imbalance issue, so the final data consist of 159,359 with under invoicing fraud flags and 159,359 rows that show no under invoicing fraud.

The variables used consist of under invoicing flag, facilities types, incoterm, import tax rate, risk level of importers, risk level of supplier, risk level of exporter country, and risk level of commodities. Under invoicing flag is the target variable. Under invoicing flag was the dummy variable which is 0 value showing the transaction that's free from under invoicing fraud and 1 value used for the transaction

that shows the under invoicing fraud. While facility types indicate whether the importation is exempted from import duty or otherwise. Incoterm is a dummy variable for the type of the incoterm used in the importation transaction, of which 0 value was used for the non EXW, FCA, FOB, and CFR transaction and 1 value was used for the EXW, FCA, FOB, and CFR transaction. Import tax rates calculated as import duty paid divided by the total of customs value. Risk level of importers, risk level of supplier, risk level of exporter country, and risk level of commodities defined based on the risk mapping using clustering techniques mentioned in the methodology section.

To obtain general information on the data owned, the authors use descriptive statistical analysis of the import duty rate and other variables. For the import duty rate, the maximum rate is 78%, the minimum rate is 0%, and the average rate is 4%. The table below describes about the level of risk in our dataset.

Tabel 2 level of risk of the dataset

Level of risk	High	Moderate	low
Importer	232.267	47.753	38.698
Commodity	242.346	48.118	28.254
Supplier	201.861	61.221	55.636
Exporter country	316.860	1.421	437

For the incoterm variable, our dataset consists of 242,612 data that use the EXW, FCA, FOB, and CFR incoterm in the importation transaction and 76,104 data that used incoterm except EXW, FCA, FOB, and CFR incoterm in the importation transaction.

4.1.2. Pearson correlation

The author uses Pearson correlation analysis as an initial identification about the relationship between the variables. The results of the correlation between the variables obtained are shown below:

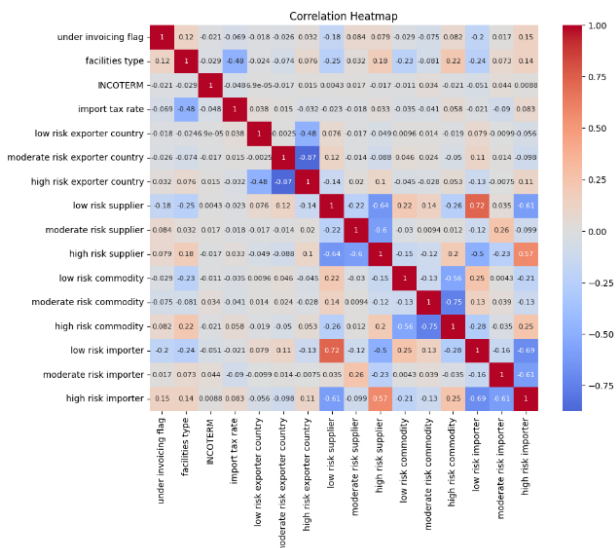


Figure 1 Pearson Correlation Matrix

Based on the figure above, a value of 1 indicates the strongest correlation between variables, so the smaller the number, the least influence between variables. In this paper, the strongest relationship between the target and independent variable is the relationship between under-invoicing flags and low risk importers (-0.22). It means that low risk importers will lead to other than under-invoicing transactions. It can be explained that the importer profiling can affect the behavior of the importer in importation declaration, the higher the risk level of importers, the bigger the risk of under invoicing fraud happening.

4.1.3. Comparing analysis

After conducting the data preparation and initial analysis using Pearson correlation, we ran the models and compared the result through evaluation of the precise, accuracy, and log loss score. In running our models, we split our datasets into training and testing dataset using ratio 70:30. After experimenting with some of the above modelling, the author measured the accuracy score, precise score and log loss values of some of the above modelling, with the following results:

Table 3 Accuracy Score, Precise Score and Log Loss

Models	Accuracy Score	Precise Score	Log Loss
<i>Logistic Regression</i>	0.5940	0,6047	0,6631
<i>Decision Trees</i>	0.6313	0,63193	0,6460
<i>Random Forest</i>	0.6308	0,6313	0,6295
<i>Xtreme Gradient Boost</i>	0.6314	0,63198	0,6258
<i>Gaussian NB</i>	0.5767	0,6238	1,3472
<i>K-nearest</i>	0.5784	0,5835	2,2349

Models	Accuracy Score	Precise Score	Log Loss
<i>Neighbor</i>			
ANN	0.6292	0,6204	0,6287

From the modelling results above, the accuracy score obtained with the lowest score is Gaussian NB modelling with an accuracy score of 0.5767, and the model with the highest accuracy score is Xtreme Gradient Boost with a value of 0.6314. Meanwhile, for the precise score, the lowest score was obtained using the K-nearest Neighbor Classifier modelling with a score of 0.5835, and the highest score using Xtreme Gradient Boost with a score of 0.63198. As for the Log Loss score, the modelling with the highest score is obtained using K-nearest Neighbor Classifier with a score of 2.2349, and with the lowest and ideal score using Xtreme Gradient Boost with a score of 0.6258. Based on the data above, it can be concluded that the Xtreme Gradient Boost Classifier modelling provides better results than other modelling in the under-invoicing fraud detection, although the slight difference from the accuracy score, precise score and log loss values with other modelling.

In addition, the confusion matrix used to evaluate the performance of the Xtreme Gradient Boost models. As it is seen in the figure below, the model predicted most of the under invoicing fraud transactions in both training dan testing data. In the training data, the models predicted 66% the under invoicing fraud transaction while in the testing data, the models predicted 66.9% the under invoicing fraud transaction.

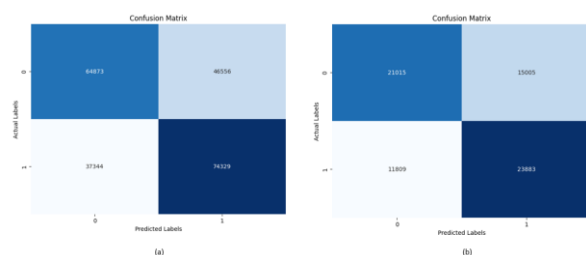


Figure 2 Training and Testing Confusion Matrices.
(a) training confusion matrix; (b) testing confusion matrix

4.2. Discussion

The result of this study shows that the Xtreme Gradient Boost algorithm performs best in terms of accuracy, precision and log loss, while the Artificial Neural Network performs less accuracy and precise score and bigger log loss. This study shows the different result of previous research by [9] which shows that the Artificial Neural Network generates the best performance. It could happen because there could be different characteristics between the income tax and under-invoicing importation fraud. Even though, it may generate a different result or may be

the better result if we explore more about the variables or apply the hyper-parameter tuning. Nevertheless, we conclude that the result of this study is still acceptable since they successfully detected more cases of under-invoicing fraud transactions than those that went undetected.

5. CONCLUSIONS

Based on the results of Pearson correlation analysis, the facility types, the risk level of importer and risk level of supplier variables have a higher correlation rate than other variables. This shows that the facility code, risk importer, and risk supplier can be factors that may influence fraud in under invoicing importation, followed by other factors.

In addition, based on the results of the research above, it can be concluded that the Xtreme Gradient Boost Classifier modelling provides the best results for detecting fraud in under invoicing importation. This can be seen from the higher accuracy score and precise score values while the log loss score is the lowest compared to other modelling.

Therefore, the authors recommend that Xtreme Gradient Boost Classifier modelling to be used as a data analytic tool to detect fraud in the under invoicing importation. With the use of these data analytic tools, hopefully it can provide an early warning/alert for customs and excise officials on transactions that are indicated to be fraud in under invoicing importation so that they can be followed up with an in-depth examination. It is also expected that document examination and physical examination of import transactions with a risk of fraud are more targeted, effective and efficient.

However, the result of this modelling isn't very accurate, with an average accuracy of less than 63%. It means that there's still room to improve the variables and models in the future research to get better results and recommendations for under-invoicing fraud detection. The limitation of this paper is that the examination result of customs officer used is only the examination result of clearance examination, which may affect the under invoicing fraud flag. The other limitation is the variable level of risk used was determined by the amount of under invoicing fraud found in the clearance examination, which may not represent the level of fraud risk comprehensively.

6. RECOMMENDATION

For future research, we suggest that the study should include better risk mapping by conducting the collaboration with the related unit. Thus, the mapping of the fraud risk level can represent the true condition and perform better modelling. Moreover, the additional examination result such as post clearance audit result and audit desk result can be included to perform under invoicing fraud flag. Thus, the under-invoicing fraud detection can generate a better alert system with higher both accuracy and precise score, and with lower log loss.

REFERENCES

Ali, A., Abd Razak, S., Othman, S. H., Eisa, T. A. E., Al-Dhaqm, A., Nasser, M., Elhassan, T., Elshafie, H., & Saif, A. (2022). Financial Fraud Detection Based on Machine Learning: A Systematic Literature Review. In *Applied Sciences (Switzerland)* (Vol. 12, Issue 19). MDPI. <https://doi.org/10.3390/app12199637>.

In-text reference: (Ali et al., 2022)

Belgiu, M., & Drăguț, L. (2016). Random forest in remote sensing: A review of applications and future directions. *ISPRS Journal of Photogrammetry and Remote Sensing*, 114, 24–31. <https://doi.org/10.1016/j.isprsjprs.2016.01.011>.

In-text reference: (Belgiu & Drăguț, 2016)

Boateng, E.Y. & Abaye, D.A. (2019). A Review of the Logistic Regression Model with Emphasis on Medical Research. *Journal of Data Analysis and Information Processing*, 7, 190-207. <https://doi.org/10.4236/jdaip.2019.74012>.

In-text reference: (Boateng & Abaye, 2019)

Charbuty, B., & Abdulazeez, A. (2021). Classification Based on Decision Tree Algorithm for Machine Learning. *Journal of Applied Science and Technology Trends*, 2(01), 20–28. <https://doi.org/10.38094/jastt20165>.

In-text reference: (Charbuty & Abdulazeez, 2021)

Chirici, G., Mura, M., McInerney, D., Py, N., Tomppo, E. O., Waser, L. T., Travaglini, D., & McRoberts, R. E. (2016). A meta-analysis and review of the literature on the k-Nearest Neighbors technique for forestry applications that use remotely sensed data. In *Remote Sensing of Environment* (Vol. 176, pp. 282–294). Elsevier Inc. <https://doi.org/10.1016/j.rse.2016.02.001>.

In-text reference: (Chirici, Mura, McInerney, Py, Tomppo, Waser, Travaglini, & McRoberts, 2016)

Dey, Ayon. (2016). "Machine learning algorithms: a review," *International Journal of Computer Science and Information Technologies*, vol. 7, no. 3, pp. 1174–1179, 2016).

In-text reference: (Dey, 2016)

Géron, A. (2019). *Hands-on Machine Learning with Scikit-Learn, Keras, and TensorFlow: Concepts, Tools, and Techniques to Build Intelligent Systems*; O'Reilly Media.

In-text reference: (Geron, 2019)

Ibrahim Ahmed Osman, A., Najah Ahmed, A., Chow, M. F., Feng Huang, Y., & El-Shafie, A. (2021). Extreme gradient boosting (Xgboost) model to predict the groundwater levels in Selangor Malaysia. *Ain Shams Engineering Journal*, 12(2), 1545–1556. <https://doi.org/10.1016/j.asej.2020.11.011>.

In-text reference: (Ibrahim Ahmed Osman, Najah Ahmed, Chow, Feng Huang, & El-Shafie, 2021)

Li, Z., Lu, T., He, X., Montillet, J. P., & Tao, R. (2023). An improved cyclic multi model-eXtreme gradient boosting (CMM-XGBoost) forecasting algorithm on the GNSS vertical time series. *Advances in Space Research*, 71(1), 912–935. <https://doi.org/10.1016/j.asr.2022.08.038>.

In-text reference: (Li, Lu, He, Montillet, & Tao, 2023)

Murorunkwere, B.F.; Tuyishimire, O.; Haughton, D.; Nzabanita, J. (2022). Fraud Detection Using Neural Networks: A Case Study of Income Tax. *Future Internet* 2022, 14, 168. <https://doi.org/10.3390/fi14060168>.

In-text reference: (Murorunkwere, Tuyishimire, Haughton, Nzabanita, 2022)

Rangineni, S., & Marupaka, D. (2023). Analysis Of Data Engineering for Fraud Detection Using Machine Learning and Artificial Intelligence Technologies. *International Research Journal of Modernization in Engineering Technology and Science*. <https://doi.org/10.56726/IRJMETS43408>.

In-text reference: (Rangineni & Marupaka, 2023)

Schonlau, M., & Zou, R. Y. (2020). The random forest algorithm for statistical learning. *Stata Journal*, 20(1), 3–29. <https://doi.org/10.1177/1536867X20909688>.

In-text reference: (Schonlau & Zou, 2020)

WCO Guide to Customs Valuation and Transfer Pricing. (2018). Brussel, Belgium: WCO.

In-text reference: (Brussel, 2018)

Wibawa, A. P., Kurniawan, A. C., Murti, D. M. P., Adiperkasa, R. P., Putra, S. M., Kurniawan, S. A., & Nugraha, Y. R. (2019). Naïve Bayes

Classifier for Journal Quartile Classification. *International Journal of Recent Contributions from Engineering, Science & IT (IJES)*, 7(2), 91. <https://doi.org/10.3991/ijes.v7i2.10659>.

In-text reference: (Wibawa, Kurniawan, Murti, Adiperkasa, Putra, Kurniawan, & Nugraha, 2019)