

Predicting the Risk of Online Sales Fraud with the Naïve Bayes Approach on Facebook Social Media

Leony Ayu Diah Pasha

Department of Information System, Universitas Muhammadiyah Sumatera Utara, Indonesia

ABSTRACT

The rapid development of digital shopping media is accompanied by increasing cases of online fraud, especially through social media platforms such as Facebook. This study aims to develop a prediction model for the risk of online sales fraud using the Naïve Bayes algorithm. The data used is the data of buying and selling transactions that occur through the Facebook marketplace. The data has been collected on the Kaggle platform so that it can be used directly. Data in the form of extracted features include seller characteristics, products sold, number of transactions, device usage and other fraud indicators. Important features that affect the potential for fraud are identified and used in the machine learning process. The results of the study show that the Naïve Bayes model is able to provide accurate predictions in identifying the risk of online sales fraud, with a satisfactory accuracy rate of 95%. The results of the study are expected to contribute to the development of a more effective fraud detection system and increase user confidence in making online transactions.

Keyword : Naïve Bayes, Facebook, Fraud, Transactions



This work is licensed under a Creative Commons Attribution-ShareAlike 4.0 International License.

Corresponding Author:

Leony Ayu Diah Pasha,
Department of Information System,
Universitas Muhammadiyah Sumatera Utara,
Jalan Kapten Muktar Basri No 3 Medan 20238, Indonesia.
Email: leonyayudiah@gmail.com

1. INTRODUCTION

In this digital era, online buying and selling through platforms such as social media is increasing rapidly around the world, including Indonesia. Online buying and selling is now increasingly popular due to the easy and instant access to products. This is an advantage for consumers and customers as consumers do not need to go anywhere to shop, and sellers do not need to sell their products everywhere for others to see. Social media is another platform option for online buying and selling besides shopping apps. Unlike ordinary e-commerce, in social commerce buyers and sellers can interact more freely, and can transact directly on social media without having to enter another digital site or application. The rise of social media users is a separate target market for sellers to sell their wares to social media users, even some social media also create special platforms for online buying and selling through social media such as Tik-Top Shop and Facebook Marketplace.

According to Databoks, Facebook is the most widely used social media for e-commerce transactions in 2017. Based on a survey by the Indonesian E-Commerce Association (idEA), online transactions through social media such as Facebook and Instagram reached 66%. In the top position, Facebook takes a market share of up to 43%. Only 16% of sellers and buyers use the marketplace platform and there are 7% who choose to use their own website. This survey shows the phenomenon that buyers and sellers, who are mostly micro-entrepreneurs, use social media as a place for e-commerce transactions rather than the widely available marketplace platform or through their own website. The survey was conducted among around 2,000 MSMEs in 10 cities in Indonesia in 2017.

According to the Populix survey, out of 1,020 Indonesian respondents, only 86% have ever shopped via social media. Of this group, the majority shopped through Tiktok Shop. Meanwhile, fewer respondents have shopped via WhatsApp, Facebook, Instagram, and other social media applications, as shown in the graph. Populix also found that the most common products purchased by respondents through social media are clothing (61%), beauty products (43%), food and beverages (38%), and cellphones and accessories (31%). The survey was conducted on July 28-August 9, 2022 to 1,020

respondents spread across urban areas in Indonesia, the majority of them from Jabodetabek (35%), Bandung (7%), and Surabaya (7%).

However, behind the convenience and practicality of online shopping, there is also a growing problem of fraud in online sales. Online sales fraud is a serious threat, harming consumers and reducing public trust in e-commerce. Online selling through social media, particularly Facebook, is a popular trend for both small businesses and large corporations. With millions of active Facebook users, the platform has become one of the most popular places to sell online. However, this also poses a high risk of online selling scams that can harm consumers.

Online selling scams on Facebook have become a serious threat to consumers. In the form of fake accounts, scammers offer products at low prices, but do not deliver the goods after payment is made. Scams also occur through fake investments and the lure of partnerships with large rewards that are not realized. To avoid scams, it is important to verify sellers, use secure payment methods, and be wary of overly tempting offers. Facebook Data collected from online sales activities on social media, from comments, reviews, and interactions between sellers and buyers, provides valuable information when analyzing fraud risks.

Early prediction of fraud risk is essential to reduce the losses that consumers and merchants may experience when transacting online. In modeling predictions, of course, the existing data must support so that the model created will produce maximum results. Therefore, to predict the risk of fraud in online sales transactions, data sourced from Kaggle is used. The data is in the form of data on cases of online buying and selling fraud in the Facebook marketplace. In this case, the fraud in online buying and selling in question is the transaction part. Data taken from the Kaggle site has parameters in the form of columns that support analysis and prediction, using the right approach and good model building is expected to get maximum prediction results.

The Naïve Bayes approach, which is one of the classification methods in machine learning, has proven effective in predicting risk or data classification in various fields, including sentiment analysis and fraud detection. By applying the Naïve Bayes approach to online sales data on Facebook social media, it is expected to provide accurate predictions related to the possibility of fraud research by (Sunardi et al., 2022). For profiling online fraud victims in Indonesia, the accuracy value of the Naïve Bayes algorithm is 75.28%.

In the previous literature review, there are several studies that have used the Naïve Bayes method in classifying Facebook user posts to find out the phishing links of Facebook users and get an accuracy value of 99.01% (Fahmi et al., 2023). Research on prediction using Naïve Bayes has also been conducted by (Rifai et al., 2019) on predicting the graduation rate of web application-based Microsoft Office trainees and getting good results, namely an accuracy rate of 99%. These results prove that the Naive Bayes method can be used to assist in prediction. However, to predict online sales fraud, it is necessary to have transaction data that describes the characteristics of online sales fraud so that predictions can be made from this. Therefore, this research aims to fill the knowledge gap and contribute new understanding related to online sales fraud risk detection using a machine learning approach.

2. RESEARCH METHOD

This research uses hardware and software. The needs of the research determined the research tools selected. Based on the capacity and capability of each device, the software and hardware were optimized to enable the research to run properly. This research uses a quantitative approach with concrete data from current phenomena. In addition to developing a machine learning model using the Naive Bayes algorithm to predict the outcome of the current data, statistical methods will be used to extract information from the data. The research process starts with selecting the topic and background of the problem. Next, a solution strategy is designed and the results are evaluated. Then the research report To achieve the research objectives, planning the stages of activities is very important. To collect data for this research, literature was used to search for scientific articles, journals, or other sources related to the research topic. The foundation, background, and concepts of this research are supported by several references related to this research. Undoubtedly, the supporting literature will relate to data collection, Naive Bayes algorithm, machine learning modeling, and other sources related to data collection and variables such as the Facebook market.

3. RESULTS AND DISCUSSION

A. Data Collection

1) Related Research

Data is collected from various relevant sources, both from the source to be processed and supporting sources, such as relevant literature reviews, which are used to help formulate problems and support current theories. The process of collecting data by searching for related research is a stage to find material related to the research problem discussed. Research related to machine learning to predict the risk of fraud from an e-commerce, especially the Facebook marketplace using the Naïve Bayes method.

2) Kaggle Data

The data that will be processed is taken from the Kaggle website. The dataset used is called Fraudulent E-Commerce Transactions. The data set is sourced from the Kaggle website. The dataset is a collection of buying and selling transaction data on the Facebook marketplace which is collected as material for the development of machine learning prediction of fraudulent e-commerce transactions. The dataset is divided into 2 parts, the first dataset contains 1,472,952 data and the second dataset contains 23,634 data. In this research, the second dataset is used. The data is a Facebook marketplace purchase transaction data designed to simulate transaction data from e-commerce platforms with a focus on fraud detection. It contains various features commonly found in transaction data, with 16 additional attributes specifically designed to support the development and testing of fraud detection algorithms.

Fraudulent_E-Commerce_Transaction_Data_2.csv (6.3 MB)

Detail Compact Column 10 of 16 columns

Transaction ID	Customer ID	Transaction Amount	Transaction Date	Payment Method	Product Category
23634 unique values	23634 unique values	10 9.72k	2024-01-01 2024-04-07	debit card 25% credit card 25% Other (11759) 50%	home & garden electronics Other (14100)
c12e87a8-8a86-4c8d-b5cc-84f3a768578	8ca9f182-82a4-4287-ab63-484e83a1bd18	42.32	2024-03-24 23:42:43	PayPal	electronics
7d187683-7961-4fce-9827-9698e2b6a201	4d158416-caae-4b09-bd5b-15235deb9129	301.34	2024-01-22 00:53:31	credit card	electronics
f2c14f9d-92df-4aaf-9931-ceaf4e63ed72	ccae47b8-75c7-4f5a-aa9e-957deced2137	348.32	2024-01-22 08:05:03	debit card	toys & games
e99496fa-194d-486b-b4da-9565fca9e5ce	b94960c8-eeee-4987-b1cd-4819816adce1	95.77	2024-01-16 20:34:53	credit card	electronics
7362837c-7538-434e-8731-bdf713f8f26d	de9d6391-b3a7-4bc7-9a55-8f013ab66928	77.45	2024-01-16 15:47:23	credit card	clothing
5da506fe-d4df-474a-b773-14633f86cfe	83033baf-2b0c-4688-b5b8-9c86976f4948	345.27	2024-02-22 13:49:27	PayPal	toys & games
47b35c5d-d4c9-4a7d-a354-cd4158abf67	6a5395a3-b47c-4bdc-91d7-3bf126530e91	53.69	2024-03-21 13:42:10	debit card	toys & games

The dataset contains 23,634 data which is a collection of transaction data on buying and selling in the Facebook marketplace. Download data in the form of csv files and named the file "Data.csv". Furthermore, to process the data, data importing is carried out into the Google Colab application for further processing and analysis. Importing data using the python programming language with the help of the pandas library.

```
import pandas as pd
df = pd.read_csv('Data.csv')
df.head()
```

	Transaction ID	Customer ID	Transaction Amount	Transaction Date	Payment Method	Product Category	Quantity	Customer Age	Customer Location	Device Used	IP Address	Shipping Address	Billing
0	c12a07a0-8a0b-4c0d-b5cc-04f3a768570	8ca9f102-02a4-4207-ab63-484e83a1bd10	42.32	2024-03-24 23:42:43	PayPal	electronics	1	40	East Jamestown	desktop	110.87.246.85	5308 Rachel Stravenue Suite 718nNorth Blakebu...	5308 Rachel Stravenue Suite 718nNorth Blakebu...
1	7d187683-7961-4fce-9827-9698e2b6a201	4d158416-caae-4b09-bd5b-15235deb9129	301.34	2024-01-22 00:53:31	credit card	electronics	3	35	Kingstad	tablet	14.73.104.153	5230 Stephanie ForgeinCollinsbury, PR 01853	5230 Stephanie ForgeinCollinsbury, PR 01853
2	f2c14f9d-92df-4aaf-9931-ceaf4e63ed72	ccae47b8-75c7-4f5a-aa9e-957deced2137	340.32	2024-01-22 08:05:03	debit card	toys & games	5	29	North Ryan	desktop	67.58.94.93	195 Cole OvalinPort Larry, IA 58422	195 Cole OvalinPort Larry, IA 58422
3	e99496fa-194d-486b-b4da-9565fca9e5ce	b94960c8-eeee-4987-b1cd-4819816adce1	95.77	2024-01-16 20:34:53	credit card	electronics	5	45	Kaylsville	mobile	202.122.126.216	7609 Cynthia SquareinWest	7609 Cynthia SquareinWest

B. Data Analysis

1) Data Preparation

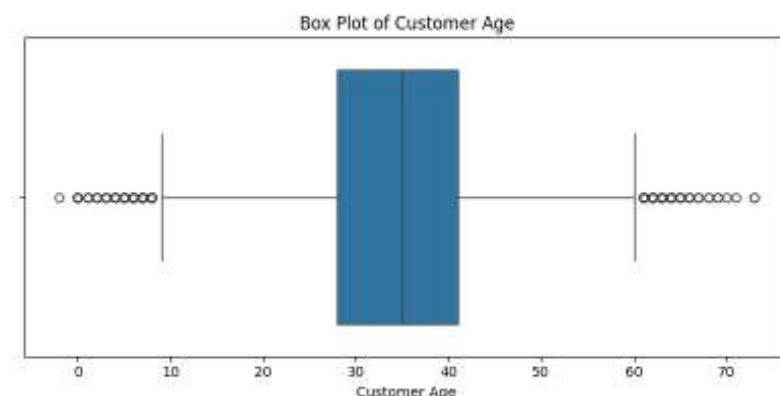
Data preparation is the process of preparing data before conducting the analysis process. The initial stage of data preparation is to check for missing values or empty data in the data set. By using the pandas function “df.isnull().sum()”.

```
df.isnull().sum()
```

	0
Transaction ID	0
Customer ID	0
Transaction Amount	0
Transaction Date	0
Payment Method	0
Product Category	0
Quantity	0
Customer Age	0
Customer Location	0
Device Used	0
IP Address	0
Shipping Address	0
Billing Address	0
Is Fraudulent	0

There is no missing value in the data, which means the data is complete. Next is to check the data whether there is duplicate data in the data set by using the function “df.duplicated().sum()” obtained there is no duplicate data in the data set. To see the value of data distribution in the form of visualization graphs, this method provides a visual description of the data range, including the minimum value, first quartile (Q1), median, third quartile (Q3), and maximum value. Box plot is used because it is an easy-to-understand and widely used method.

At this stage, the variables seen are quantitative variables that have the possibility of wild data distribution. It was found that the variable “Customer age” had data that was outside the normal distribution. The minimum value of Q1 or the lower limit of the data is 10, but there are some data that are below it. The Q2 value or the middle value is 10 to 60 and the Q3 value or the upper limit is 60 but there is still data that is greater than that. Data that is included outside of Q1 and Q3 is data that is not normally distributed or called outlier data to overcome this by deleting data. The existence of data that is not normally distributed can result in the performance of the model to be built less than optimal.



Data deletion using the loc() function with data parameters in the "Customer age" column whose value is less than 10 and whose value is greater than 60. After deletion, it can be seen that the data in the "Customer age" column with the lowest value is 10 and the highest is 60, so the outlier data has been removed and the data has been normally distributed.

```

mean_age = df['Customer Age'].mean()
df.loc[df['Customer Age'] < 10, 'Customer Age'] = mean_age
df.loc[df['Customer Age'] > 60, 'Customer Age'] = mean_age

```

Tampilkan output tersembunyi

```

df.describe()

```

	Transaction Amount	Quantity	Customer Age	Is Fraudulent	Account Age Days	Transaction Hour
count	23634.000000	23634.000000	23634.000000	23634.000000	23634.000000	23634.000000
mean	229.367099	3.000550	34.603326	0.051705	178.660531	11.266015
std	282.046669	1.419663	9.541046	0.221436	107.388682	6.980659
min	10.000000	1.000000	10.000000	0.000000	1.000000	0.000000
25%	69.070000	2.000000	28.000000	0.000000	84.000000	5.000000
50%	151.415000	3.000000	34.560210	0.000000	178.000000	11.000000
75%	296.127500	4.000000	41.000000	0.000000	272.000000	17.000000
max	9716.500000	5.000000	60.000000	1.000000	365.000000	23.000000

In the database there are Shipping Address and Billing Address columns that function to record address details related to each transaction. For the purposes of security analysis and detection of potential fraud, a comparison is made between the two columns to produce a match of addresses that will be entered into the Address Match column. The column will contain a parameter number of 1 for matching shipping and billing addresses and 0 for the opposite. This address match column is used as one of the indicators in identifying potentially suspicious or fraudulent transactions. If there is a significant mismatch between the shipping address and billing address, this can be one of the parameters for detecting transactions that are categorized as risky or fraudulent.

Is Fraudulent	Account Age Days	Transaction Hour	Address Match	Transaction Day	Transaction Month
0	282	23	1	Sunday	March
0	223	0	1	Monday	January
0	360	8	0	Monday	January
0	325	20	1	Tuesday	January
0	116	15	1	Tuesday	January

For analysis purposes, not all columns will be used, so it would be better to delete unused columns for ease of further analysis. The unused columns are: "Transaction ID", "Customer ID", "Customer Location", "Transaction Date", "IP Address", "Shipping Address", "Billing Address". Deleting columns will not reduce the quality of data and analysis results because these columns do not have a significant influence for analysis purposes.

```

df = df.drop(["Transaction ID", "Customer ID", "Customer Location", "Transaction Date", "IP Address", "Shipping Address", "Billing Address"], axis=1)
df.head()

```

	Transaction Amount	Payment Method	Product Category	Quantity	Customer Age	Device Used	Is Fraudulent	Account Age Days	Transaction Hour	Address Match	Transaction Day
0	42.32	PayPal	electronics	1	40.0	desktop	0	282	23	1	Sunday
1	301.34	credit card	electronics	3	35.0	tablet	0	223	0	1	Monday
2	340.32	debit card	toys & games	5	29.0	desktop	0	360	8	0	Monday
3	95.77	credit card	electronics	5	45.0	mobile	0	325	20	1	Tuesday
4	77.45	credit card	clothing	5	42.0	desktop	0	116	15	1	Tuesday

2) Exploratory Data Analyst

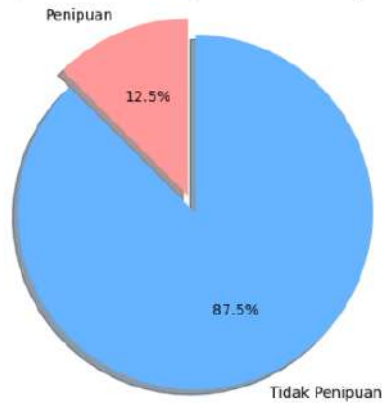
Exploratory data analyst is a crucial step in conducting data analysis. The EDA stage is carried out to determine data characteristics, determine patterns, detect anomalies from the database and determine relationships between variables. This process is not only useful for understanding data in depth, but also as a basis for preparing prediction or classification models in the next analysis stage, especially in the context of fraud detection. The initial stage is to look at the distribution of numeric data such as the number of transactions, the number of items purchased and the distribution of other variables. The distribution of data based on the number of

transactions shows that there are 100 data indicated as fraudulent transaction data and there are 18,658 data that are not included in fraudulent data. From this distribution, it can also be seen that the number of transactions from data indicated as fraud is lower than the number of transactions indicated as non-fraudulent.

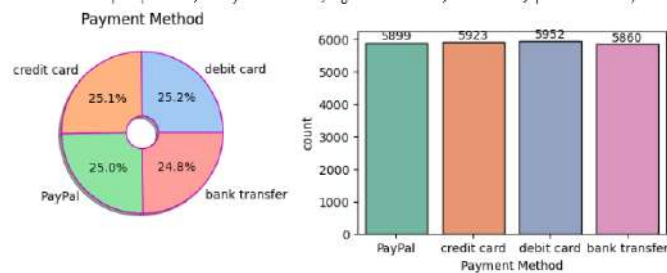
Jumlah total data yang termasuk penipuan: 1000
 Jumlah total data yang tidak termasuk penipuan: 18658

Jumlah total transaksi yang termasuk penipuan: 564013.12
 Jumlah total transaksi yang tidak termasuk penipuan: 3945976.1299999994

Proporsi Transaksi Penipuan vs Tidak Penipuan



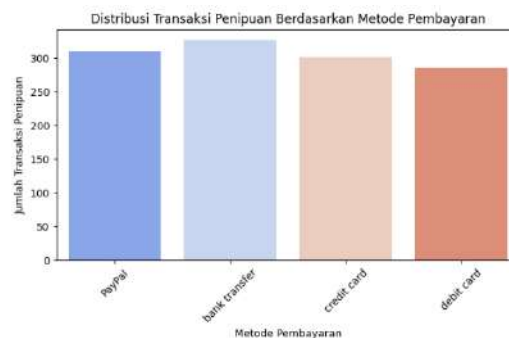
Our advanced analysis needs to know the distribution of each variable that exists both against the amount of data and against data that is indicated as fraud. One way is to look at the distribution of payment methods. What payment methods are most widely used and what payment methods are most indicated as fraudulent transactions.



➡ Jumlah transaksi penipuan per kategori barang:

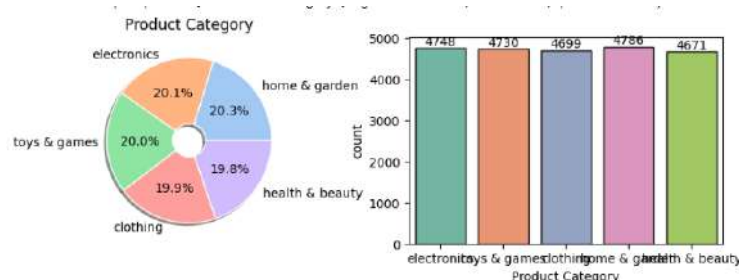
Payment Method	Count
bank transfer	326
PayPal	310
credit card	301
debit card	285

Name: count, dtype: int64

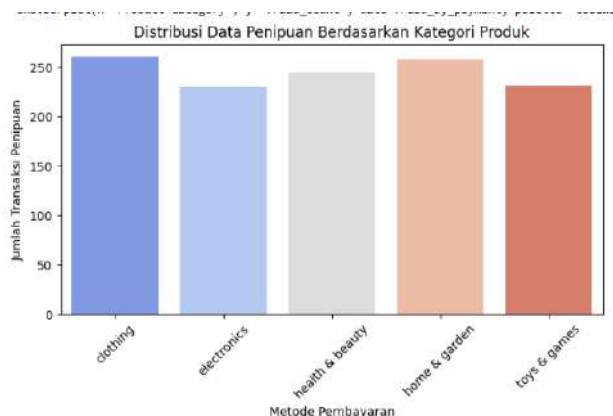


Based on the image above, the distribution of data based on payment methods can be said to be evenly distributed. However, for data that indicates fraudulent transactions, it is found that the payment method using bank transfers is the most indicated fraudulent method, while the lowest is using the debit card method. There are 326 fraudulent transaction data using the bank transfer payment method, this figure is the highest and the lowest is 285 fraudulent transaction data using the debit card payment method. It can be concluded that the payment method using bank transfers has a higher risk of fraudulent transactions.

Next, we also have to look at the distribution of data based on other variables, namely Product category. We will see the distribution of data based on product categories. From this information, we can see the possible anomalies from the Product category variable that may affect data patterns.

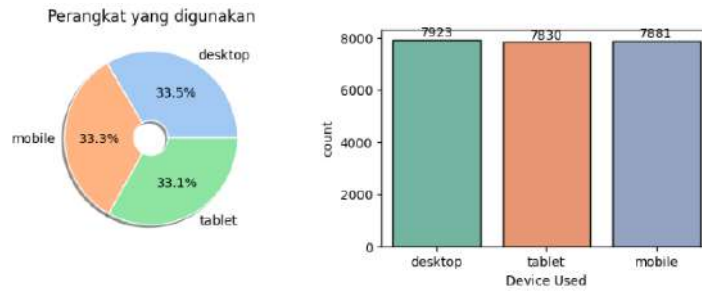


```
Jumlah transaksi penipuan per kategori barang:
Product Category
clothing      260
home & garden  258
health & beauty 244
toys & games   231
electronics   229
Name: count, dtype: int64
```



For the total amount of data, it can be seen that the distribution of data based on product categories is evenly distributed. Meanwhile, for data that indicates fraud, the clothing product category has the highest value with 260 fraudulent transaction data or it can be said that the most fraudulent transactions based on product categories are clothing and the lowest is the electronics product category with 229 fraudulent data. This indicates that the clothing product category has a higher chance of fraudulent transactions.

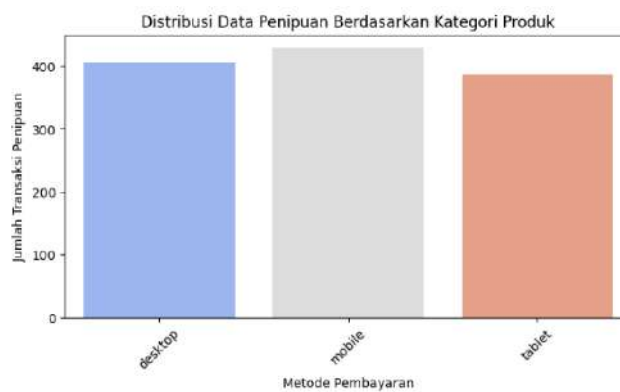
Next is to look at the distribution of data based on the devices used to make purchases or transactions. From this data distribution, we can see what kind of devices are most prone to fraudulent transactions.



```

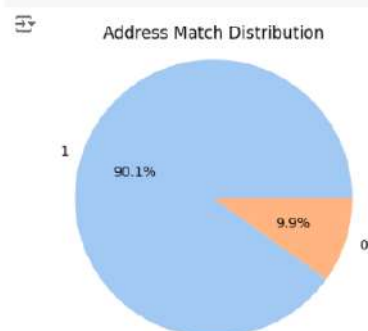
Jumlah transaksi penipuan berdasarkan penggunaan perangkat
Device Used
mobile      429
desktop     406
tablet      387
Name: count, dtype: int64

```



From the graph displayed, it can be seen that the distribution of data based on the device used for the purchase is evenly distributed for the data as a whole. For data that includes fraudulent transactions, the mobile device category is the most fraudulent transaction data with 429 fraudulent transaction data that occurred from mobile devices. This influence can provide insight that the use of mobile devices to make buying and selling transactions is more prone to fraudulent transactions.

The results of the analysis above become new knowledge or useful information for the continuation of the analysis process. Anomalies and patterns in the data can be seen from the analysis. Furthermore, the address match column will be seen for its distribution, this can be an important part in the process of finding data patterns to indicate transactions that are classified as fraudulent transactions.



The address match column is a column that informs the match between the shipping address and the billing address. Its contents are binary 0 or 1 where the number 1 indicates that the shipping address and billing address match and 0 indicates that the shipping address and billing address do not match. The distribution of this data is to see whether a shipping address

(L.A.D. Pasha)

that does not match the billing address has a greater chance of fraudulent transactions or vice versa. From the data above, it can be seen that data that does not match the shipping address and billing address is lower, most of the data distribution does have a match between the shipping address and billing address. Only 9.9% of data has an address that does not match, and there is 90.1% of data whose address matches.

In the existing dataset, there is an account age column that indicates how long the account has been used for transactions. From this, we can see the distribution of existing data and see anomalies based on the age of the account.



From the diagram shown, it can be seen that fraud cases occur in accounts that are still new or around 125 days or less than 4 months. While older accounts do not have fraud cases. The account age diagram is taken based on days and the average of all data.

3) Feature Selection

The feature selection or feature relationship or feature selection stage is the process of selecting the most relevant and informative or most influential subset of features (variables) on the target variable. This stage is carried out to facilitate the model's performance in targeting or focusing analysis on the most influential features on the target variable, namely the Is Fraudulent column.

The feature selection process can only be carried out on numeric columns or variables so that variables containing non-numeric data should not be used or the data should be changed to numeric data. In this process, columns such as "Payment method", "Product Category", "Device Use" and "Transaction Day" are changed to numeric data types. More details can be seen in the image below.

```
df1 = df.copy()
df1['Payment Method'] = df1['Payment Method'].replace(['PayPal'],1)
df1['Payment Method'] = df1['Payment Method'].replace(['credit card'],2)
df1['Payment Method'] = df1['Payment Method'].replace(['debit card'],3)
df1['Payment Method'] = df1['Payment Method'].replace(['bank transfer'],4)
#-----
df1['Product Category'] = df1['Product Category'].replace(['home & garden'],1)
df1['Product Category'] = df1['Product Category'].replace(['health & beauty'],2)
df1['Product Category'] = df1['Product Category'].replace(['toys & games'],3)
df1['Product Category'] = df1['Product Category'].replace(['electronics'],4)
df1['Product Category'] = df1['Product Category'].replace(['clothing'],5)
#-----
df1['Device Used'] = df1['Device Used'].replace(['mobile'],1)
df1['Device Used'] = df1['Device Used'].replace(['desktop'],2)
df1['Device Used'] = df1['Device Used'].replace(['tablet'],3)
#-----
df1['Transaction Day'] = df1['Transaction Day'].replace(['Sunday'],1)
df1['Transaction Day'] = df1['Transaction Day'].replace(['Monday'],2)
df1['Transaction Day'] = df1['Transaction Day'].replace(['Tuesday'],3)
df1['Transaction Day'] = df1['Transaction Day'].replace(['Wednesday'],4)
df1['Transaction Day'] = df1['Transaction Day'].replace(['Thursday'],5)
df1['Transaction Day'] = df1['Transaction Day'].replace(['Friday'],6)
df1['Transaction Day'] = df1['Transaction Day'].replace(['Saturday'],7)
#-----
```

For the data in the "Payment Method" column, the data changes to numeric data with numbers 1 to 4 based on the existing data. For the "PayPal" payment method, it is changed to number 1, for the "credit card" payment, it is changed to number 2, "debit card" is changed to number 3, "bank transfer" is changed to number 4. For the data in the "Product Category" column, the same thing is done by changing the existing data to numbers from 1 to 5 based on different

product categories. Likewise with the other columns. By changing the data into numeric data, the process of selecting the most relevant features can be done by entering the columns above.

Method for selecting the most influential features on data by using the heatmap method. Heatmap is a data visualization that is very useful for understanding the relationship between variables or features in a dataset. By using different colors to represent the intensity level of a value, heatmaps can easily show patterns and trends in data. (Satria et al., 2023). Usually this method is also known as a correlation matrix.

Heatmap, also known as heatmap is a two-dimensional data visualization technique that shows the variation in the magnitude of a phenomenon in terms of the colors represented by various colors. Heatmaps show the shape and direction of different heat values for a number of data points at various temperature levels. Correlation has a positive value, which means that an increase in the value of one feature increases the value of the target variable, or a negative value, which means that an increase in the value of one feature decreases the value of the target variable. (Bengnga & Ishak, 2022)



From the correlation matrix with the heatmap method above, we can see that the distribution of all data or variables is normal. Showing an insignificant number, but it can also be seen that the variable that has the most influence on the target variable is the "Transaction Amount" column or the number of transactions which shows the number 0.29. This indicates that the "Transaction Amount" column is positively correlated with the "Is Fraudulent" column. It can be seen that other data does not have a significant influence other than the account age column and transaction day which each have a value of 0.14, but this number is still smaller than the number of transactions column.

C. Data Modeling

In this study, the data modeling process is carried out to build a predictive model that can predict consumer behavior based on available e-commerce transaction data. This process involves the application of machine learning algorithms, such as Naive Bayes, to identify patterns and relationships between significant variables. Data modeling is a crucial step in the analysis because it allows researchers to simplify the complexity of the data and produce accurate models that can be used for data-based decision making. The following are the steps in data modeling in this study:

1) Import Library

Before doing data modeling, several libraries are needed to support Naïve Bayes data modeling. Some libraries needed are Pandas, numpy, seaborn, matplotlib, stiklearn and others for the data analysis and modeling process. The libraries used have also been used in previous stages such as the EDA and data preparation stages.

```
import numpy as np
import pandas as pd
import matplotlib.pyplot as plt
import seaborn as sns
import time
import pickle
import warnings
warnings.filterwarnings("ignore")
from sklearn.model_selection import train_test_split, RandomizedSearchCV
from sklearn.naive_bayes import GaussianNB
from sklearn.naive_bayes import MultinomialNB
from sklearn.naive_bayes import BernoulliNB
from sklearn.pipeline import Pipeline
from sklearn.compose import ColumnTransformer
from sklearn.preprocessing import MinMaxScaler, LabelEncoder, OneHotEncoder
from sklearn.metrics import classification_report, confusion_matrix, accuracy_score
```

2) Splitting Data

The data split process is to divide the data into training data and test data using the sklearn function "train_test_split". This function will divide the data into training data and test data (validation data). The size of the dataset division is set to 0.2 or 20% which will divide the training data by 80% and the test data by 20%. From this division, the results obtained were that 18,694 data were training data (train_data) and 4,674 were training data (test_data).

```
[62] train_df, test_df = train_test_split(df, test_size=0.2, random_state=42)
# Menampilkan jumlah data latih dan data uji
print(f"Jumlah data latih: {len(train_df)}")
print(f"Jumlah data uji: {len(test_df)}")
```

```
Jumlah data latih: 18694
Jumlah data uji: 4674
```

The next modeling stage is to sort and adjust the features or columns in train_data and test_data. For modeling, the target column "Is Fraudulent" will be discarded because it is not analyzed, but the column will be the target for the model to perform parameters on other columns. In the category column such as "address match" will be identified based on the data type in the column, namely category. This is because the model will read each value in each column differently, if the model is not informed that there is a category column, the results that will be issued will also be different and this will affect the accuracy value of the model created.

```
[63] ## Train Data
train_data = train_df.drop(columns=["Is Fraudulent"])
train_label = train_df["Is Fraudulent"]

## Test Data
test_data = test_df.drop(columns=["Is Fraudulent"])
test_label = test_df["Is Fraudulent"]

cat_col = train_data.select_dtypes(include="O").columns
num_col = []
for col in train_data.columns:
    if col not in cat_col and col != 'Is Address Match':
        num_col.append(col)
```

The next stage in data modeling is to transform the data into values that can be read by the model to be built. At the data transformation stage, normalization is carried out using the MinMaxScaler method to change the numeric feature values into a range between 0 and 1. This method is used to ensure that each feature has the same scale, so that it can prevent the dominance of features with a larger range of values in the model training process.

MinMaxScaler was chosen because it does not produce negative transformation values because the algorithm modeling using Naïve Bayes cannot read negative values. Normalization with MinMaxScaler is done by calculating the minimum and maximum values of each feature, then transforming each value based on the equation:

$$X_{scal} = \frac{X - X_{min}}{X_{max} - X_{min}}$$

In the transformation process, OneHotEncoding is also used which is intended for categorical data. If MinMaxScaler is used for numeric columns, OneHotEncoder is intended for categorical data. Then the use of ColumnTransformer allows all transformations to be applied in one step, thus simplifying the preprocessing process of data that has features with different types. Thus, this process ensures that the data is ready to be used by the machine learning model without losing important information.

```
transformer = ColumnTransformer(transformers=[
    ('encoding', OneHotEncoder(), cat_col),
    ('scaling', MinMaxScaler(), num_col)
], remainder='passthrough')
```

D. Naïve Bayes Algorithm

Naive Bayes is used in this study because it has advantages in handling large datasets and works well for data with simple probability distributions. This algorithm calculates the posterior probability of each class based on the distribution of features in the data, then selects the class with the highest probability to predict.

The Naive Bayes algorithm can be divided into several variants depending on the data distribution, such as Gaussian Naive Bayes for continuous data with a normal distribution, Multinomial Naive Bayes which is suitable for frequency data or event calculations, and Bernoulli Naive Bayes which is used for binary data. In this study, the Naive Bayes variant chosen is adjusted to the characteristics of the dataset used. Naive Bayes was chosen as one of the main algorithms in data modeling in this study because of its efficient ability to handle large datasets and its reliability in solving classification problems with various types of data distributions.

```
classifiers = {
    "GaussianNB": GaussianNB(),
    "MultinomialNB": MultinomialNB(),
    "BernoulliNB": BernoulliNB()
}

classifier_name = []
accuracy = []
for name, classifier in classifiers.items():

    model = Pipeline(steps=[
        ('transformer', transformer),
        ('classifier', classifier)
    ])

    start_time = time.time()

    model.fit(train_data, train_label)

    training_time = time.time() - start_time

    test_pred = model.predict(test_data)
```

E. Model Evaluation

After the modeling process is complete, the next stage is model evaluation to assess the performance of the model that has been built. Evaluation is carried out using three main metrics, namely accuracy score, confusion matrix, and classification report. In model evaluation there are several terms that describe the results of the data itself, here is the explanation:

- 1) Accuracy Score: The accuracy score metric measures the proportion of correct predictions compared to the total predictions made by the model. Accuracy is calculated using the formula:

$$Accuracy = \frac{true\ positif + true\ negatif}{total\ data}$$

In this study, the evaluated model obtained an accuracy score of X%, which indicates the model's ability to classify data well. A high accuracy value indicates the model's ability to classify data well.

- 2) Confusion Matrix: Visually representing the performance of a classification model, the confusion matrix shows the number of correct and incorrect predictions for each class. True positives (TP),

true negatives (TN), false positives (FP), and false negatives (FN) are the four parts of this matrix. We can calculate other metrics such as precision, recall, and F1 score from the confusion matrix.

F. Prediction Result

Prediction results Based on the Naïve Bayes model test used in this study, an accuracy level of 95% was obtained. This value indicates that the model is able to predict fraud risk very well, with 95% of the tested data being predicted correctly by the model. From the analysis of the features used, it was found that the number of transactions feature has the greatest influence on the prediction objective, namely fraud. This feature plays an important role in identifying suspicious transaction patterns.

The model also shows that the greater the number of transactions or the more unusual the transaction, the greater the likelihood of fraud. In addition, balanced data distribution and proper feature selection also contribute to high model accuracy. The prediction results of each data are displayed in a new table "Prediction Results". By using the highest accuracy algorithm model, namely GaussianNB. The column is the prediction result of the GaussianNB model which contains the string values "Fraud" and "Not Fraud" this is obtained from GaussianNB modeling with the target parameter column "Is Fraudulent" and features from other columns. The following is a display of the data prediction results:

```
# Lakukan prediksi pada seluruh dataset menggunakan model terbaik yaitu GaussianNB
model = Pipeline(steps=[
    ('transformer', transformer),
    ('classifier', GaussianNB())
])

model.fit(train_data, train_label)
predictions = model.predict(df.drop(columns=['Is Fraudulent']))

# Tambahkan kolom "Hasil Prediksi" ke dataframe
df['Hasil Prediksi'] = predictions

# Ganti nilai 0 dan 1 pada kolom "Hasil Prediksi" dengan "Tidak Penipuan" dan "Penipuan"
df['Hasil Prediksi'] = df['Hasil Prediksi'].replace({0: 'Tidak Penipuan', 1: 'Penipuan'})

# Tampilkan dataframe dengan kolom "Hasil Prediksi"
df.head()
```

Product Category	Quantity	Customer Age	Device Used	Is Fraudulent	Account Age Days	Transaction Hour	Address Match	Transaction Day	Transaction Month	Hasil Prediksi
electronics	1	40	desktop	0	282	23	1	Sunday	March	Tidak Penipuan
electronics	3	35	tablet	0	223	0	1	Monday	January	Tidak Penipuan
toys & games	5	29	desktop	0	360	8	0	Monday	January	Tidak Penipuan
electronics	5	45	mobile	0	325	20	1	Tuesday	January	Tidak Penipuan
clothing	5	42	desktop	0	116	15	1	Tuesday	January	Tidak Penipuan

4. CONCLUSION

The accuracy, precision, recall and f1-score values of the Naïve Bayes algorithm are classified as good in reading and predicting existing data because the average value of the three Naïve Bayes algorithms used shows a figure of 0.95 or 95%. This shows that Naïve Bayes is an effective method in detecting fraud. The variable that has the most influence on the target variable, namely "Is Fraudulent" is the "Transaction Amount" variable measured by the correlation matrix method showing the highest positive number compared to other variables, namely 0.28. The transaction value has a strong correlation with the risk of fraud, indicating that the nominal amount of the transaction is important to consider in risk analysis. Naïve Bayes is a simple algorithm but still able to provide good performance. This model can be implemented with minimal computing resources, making it the right choice for efficient fraud detection. There are 12.5% of data from the total data that is classified as fraudulent transaction data. Buying and selling through the Facebook marketplace is considered not dangerous

because of the existing data, not even half of it is fraudulent transaction data, so it can be concluded that buying and selling with the Facebook marketplace is considered safe..

REFERENCES

- [1] Acmad. (2018). Pengaruh Pengguna E-Commerce Terhadap Transaksi. *Faktor Exacta*, 11(1), 7–16.
- [2] Agusti, R., & Aravik, H. (2023). Analisis Penggunaan Marketplace Facebook Terhadap Penjualan Mebel Dalam Bauran Pemasaran Syariah Di Supran Mebel Karang Anyar Palembang. *Jurnal Bisnis Dan Manajemen (JURBISMAN)*, 1(2), 275–290.
- [3] Arifiyanti, A. A., & Wahyuni, E. D. (2020). Smote: Metode Penyeimbang Kelas Pada Klasifikasi Data Mining. *SCAN - Jurnal Teknologi Informasi Dan Komunikasi*, 15(1), 34–39. <https://doi.org/10.33005/scan.v15i1.1850>
- [4] Azeez, N. A., Misra, S., Lawal, O. I., & Oluranti, J. (2021). Identification and Detection of Cyberbullying on Facebook Using Machine Learning Algorithms. *Journal of Cases on Information Technology*, 23(4), 1–21. <https://doi.org/10.4018/JCIT.296254>
- [5] Bengnga, A., & Ishak, R. (2022). Implementasi Seleksi Fitur Klasifikasi Waktu Kelulusan Mahasiswa Menggunakan Correlation Matrix with Heatmap. *Jambura Journal of Electrical and Electronics Engineering*, 4(2), 169–174. <https://doi.org/10.37905/jjee.v4i2.14403>
- [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), 602–613.
- [7] Cholissodin, I., & Soebroto, A. A. (2021). AI, MACHINE LEARNING & DEEP LEARNING (Teori & Implementasi). December.
- [8] Damanik, A. R., Sumijan, S., & Nurcahyo, G. W. (2021). Prediksi Tingkat Kepuasan dalam Pembelajaran Daring Menggunakan Algoritma Naïve Bayes. *Jurnal Sistim Informasi Dan Teknologi*, 3, 88–94. <https://doi.org/10.37034/jsisfotek.v3i3.49>
- [9] Fahmi, M., Postingan, F. A. K., Facebook, P., Phising, D., Bayes, N., Fadhilah, M. F., Arsyi, F., & Fadlillah, N. (2023). Klasifikasi Postingan Pengguna Facebook Untuk Deteksi Phising Menggunakan Naive Bayes. *Jurnal Riset Informatika Dan ...*, 1(1), 25–29. <http://ejurnal.jejaringppm.org/index.php/jriti/article/view/49>
- [10] Koh, H. C., & Tan, G. (2011). Data mining applications in healthcare. *Journal of Healthcare Information Management*, 19(2), 65–72.
- [11] Lestari, T. S., & Sirodj, D. A. N. (2022). Klasifikasi Penipuan Transaksi Kartu Kredit Menggunakan Metode Random Forest. *Jurnal Riset Statistika*, 1(2), 160–167. <https://doi.org/10.29313/jrs.v1i2.525>
- [12] Loelianto, I., Thayf, M. S. S., & Angriani, H. (2020). Implementasi Teori Naive Bayes Dalam Klasifikasi Calon Mahasiswa Baru Stmik Kharisma Makassar. *SINTECH (Science and Information Technology) Journal*, 3(2), 110–117. <https://doi.org/10.31598/sintechjournal.v3i2.651>
- [13] Noor, T., Masnun, & Putri, K. G. (2021). *Jurnal Hukum dan Kemasyarakatan Al-Hikmah* Vol. 2, No.3, September 2021 428. *Jurnal Hukum Dan Kemasyarakatan Al-Hikmah*, 2(3), 428–446.
- [14] Purba, K. S., Hartama, D., & Suhada, S. (2022). Analisis Data Mining Pesebaran Siswa Smp Di Pematangsiantar Dengan Metode Algoritma K-Means Clustering. *Kesatria : Jurnal Penerapan Sistem Informasi (Komputer Dan Manajemen)*, 3(1), 1–8. <https://doi.org/10.30645/kesatria.v3i1.91>
- [15] Renaldy, & Putra, D. S. D. (2023). Aplikasi Prediksi Harga Ayam Dengan Metode Naives Bayes Pada Supplier Ayam Potong. *Jurnal Algor*, 4(2), 141–148. <http://repositori.buddhidharma.ac.id/id/eprint/1609>
- [16] Rifai, M. F., Jatnika, H., & Valentino, B. (2019). Penerapan Algoritma Naïve Bayes Pada Sistem Prediksi Tingkat Kelulusan Peserta Sertifikasi Microsoft Office Specialist (MOS). *Petir*, 12(2), 131–144. <https://doi.org/10.33322/petir.v12i2.471>
- [17] Satria, A., Badri, R. M., & Safitri, I. (2023). Prediksi Hasil Panen Tanaman Pangan Sumatera dengan Metode Machine Learning. *Digital Transformation Technology*, 3(2), 389–398. <https://doi.org/10.47709/digitech.v3i2.2852>
- [18] Silalahi, P. R., Salwa Daulay, A., Siregar, T. S., Ridwan, A., Islam, E., Ekonomi, F., & Islam, B. (2022). Analisis Keamanan Transaksi E-Commerce Dalam Mencegah Penipuan Online. *Jurnal Manajemen, Bisnis Dan Akuntansi*, 1(4), 224–235.
- [19] Simatupang, S., Efendi, E., & Putri, D. E. (2021). Facebook Marketplace Serta Pengaruhnya Terhadap Minat Beli. *Jurnal Ekbis*, 22(1), 28. <https://doi.org/10.30736/je.v22i1.695>
- [20] Sunardi, S., Fadlil, A., & Kusuma, N. M. P. (2022). Implementasi Data Mining dengan Algoritma Naïve Bayes untuk Profiling Korban Penipuan Online di Indonesia. *Jurnal Media Informatika Budidarma*, 6(3), 1562. <https://doi.org/10.30865/mib.v6i3.3999>
- [21] Suntoro, J. (2019). 22-DATA MINING Algoritma dan Implementasi Menggunakan Bahasa Pemrograman PHP. *DATA MINING Algoritma Dan Implementasi Menggunakan Bahasa Pemrograman PHP*, 9(9), 259–278.
- [22] Schafer, J. B., Konstan, J., & Riedl, J. (1999). Recommender systems in e-commerce. *Proceedings of the 1st ACM Conference on Electronic Commerce (EC-99)*.
- [23] Tana, M. P., Marisa, F., & Wijaya, I. D. (2018). Penerapan Metode Data Mining Market Basket Analysis Terhadap Data Penjualan Produk Pada Toko Oase Menggunakan Algoritma Apriori. *J I M P - Jurnal Informatika Merdeka Pasuruan*, 3(2), 17–22. <https://doi.org/10.37438/jimp.v3i2.167>
- [24] Twin, A. (2005). Data Mining Data mining. In *Mining of Massive Datasets* (Vol. 2, Issue January 2013). https://www.cambridge.org/core/product/identifier/CB09781139058452A007/type/book_part