

Support Vector Machine and Naïve Bayes for Personality Classification Based on Social Media Posting Patterns

Bayu Seno Nugroho*, Warih Maharani

School of Informatics, Informatics, Telkom University, Bandung, Indonesia

Email: ^{1,*}baysen@student.telkomuniversity.ac.id, ²wmaharani@telkomuniversity.ac.id

Correspondence Author Email: baysen@student.telkomuniversity.ac.id

Submitted: 06/12/2024; Accepted: 12/12/2024; Published: 19/12/2024

Abstract—This research investigates the use of Support Vector Machine (SVM) and Naive Bayes models to classify the personality traits based on the social media posting patterns. This study integrates textual features obtained from the Bag-of-Words (BoW) and Term Frequency-Inverse Document Frequency (TF-IDF) methods, and along with the feature expansion using the Linguistic Inquiry and Word Count (LIWC) tool, to assess their influence on accuracy. Classification Personality characteristics were mapped from social media posts using the Big Five Inventory (BFI-44). The research findings show that the SVM model in which uses the TF-IDF + LIWC feature set, provides the best performance, and achieve 76.60% of accuracy on the base model with a linear kernel. In comparison to the Naive Bayes model performed best with the same feature set, achieving 59.57% accuracy with a smoothing parameter of 1×10^{-2} . Although the oversampling improved recall and precision, the undersampling was found to have a negative effect on model performance. These findings highlight the benefits of combining TF-IDF and LIWC features which improve model effectiveness, with SVM producing the best overall results in personality classification from social media data.

Keywords: Support Vector Machine; Naive Bayes; Personality Classification; Social Media; Text Classification; BFI-44

1. INTRODUCTION

Personality can be identified as the mindset of a person, which determines the behavioral pattern, feelings, and attitude of a person, uniquely distinguishing an individual from others as a unique set of personality traits [1]. Various psychologists have developed psychological theories in order to predict personality. Some of them are Big Five Personality Traits, MBTI, Strengths Finder, and DISC Personality. Among them, Big Five Personality Traits is considered one of the most reliable models for classification because it categorizes personality traits into five broad categories: neuroticism, extraversion, openness, agreeableness, and conscientiousness [2][3]. Those aspects, being structured and quantified for understanding human behavior, have also found applications in different fields like personality classification from text data. Recent works within the domain of social media analytics show that text-based data, which correspond to posts on platforms such as X, are quite indicative of an individual's personality traits, as they actually reflect underlying emotions, attitudes, and behaviors in general [4][5][6].

Growing use of social media, especially X (formerly Twitter), has opened up unique opportunities for personality analysis, whereby the text within posts could become a rich source of behavioral signals. Considering the fact that personality traits (like extraversion and agreeableness) are conveyed in the usage of language-word choices, sentimental feel, and pattern of social interaction—it becomes feasible to increasingly employ machine learning models on these textual signals for making predictions regarding personality. The Big Five Inventory-44 (BFI-44) questionnaire, commonly used in psychological research, can be used to quantify the traits and match them to the corresponding patterns in social media data, as social media posts most often reflect personal experiences and states of emotion [5][6]. Recent studies have shown that text data from social media may serve as a very promising source for the inference of personality traits since it bears rich behavioral signals [7]. The increasing usage of platforms, such as X, which give users the ability to convey their thoughts in short forms of posts, is a very good setting for performing this kind of analysis.

Text classification is widely known to be performed by some machine learning methods such as NB and SVM; these are widely applied in different studies because of their effectiveness in managing text data and their results that have been well established already in several studies [8][9][10]. Naïve Bayes enjoys favorable considerations because of its simplicity and efficiency in handling probabilistic models. Several studies have proved that it reaches high accuracy in text classification tasks, even when the dataset is small. For example, in the study of Natasuwarna [10], it reached 89.86% accuracy with Naïve Bayes in personality classification, thus proving its potential in social media analysis. The performance of NB with small datasets is another important advantage; therefore, it is very popular in research studies where the sample size is small [11].

On the other hand, Support Vector Machine is a powerful technique when it comes to handling big and complex data with high accuracy. It is mainly powerful because it uses kernel functions, which enable it to map the data into higher-dimensional spaces so that it is able to handle non-linear data classification efficiently [12][13]. In this way, SVM has especially been suited for text classification problems where relationships among the features can be complex and nonlinear. For instance, in work by Fikriani et al. [14], personality classification has been made using an SVM on tweet data to arrive at an accuracy value of 86.88%. Widyadhana et al. [15], on the other hand, have illustrated the supremacy of Support Vector Machine over other methods while offering an average accuracy value of 96.43%. This work presents a study using Naïve Bayes and SVM in classifying personalities in social media posts made via X. In the following, the Big Five Inventory-44 (BFI-44) questionnaire is leveraged to collect personality data for

Indonesian users and tweets by analyzing their personality from such data. This present research will look into identifying the better method between the two Naïve Bayes or SVM methods in this context. Additionally, it identifies elements that may influence the outcome of both approaches. The data will be pre-processed, and then both models will be developed and evaluated regarding their performance in classifying personality traits based on textual data. Additionally, the research will emphasize those factors that might influence the success of these methods, such as the length of the posts, the frequency of tweets, or the linguistic styles used by the individuals.

It is expected that the results from this research will give further insights into the effectiveness of Naïve Bayes and SVM in personality classification using textual data from social media posts, hence helping in refining the existing models and contributing to the growing area of personality computing.

2. RESEARCH METHODOLOGY

2.1 Research Stages

In this research, a system was developed to classify personality traits based on social media posts using machine learning methods. The research stages are represented in the system flowchart design shown in Figure 1.

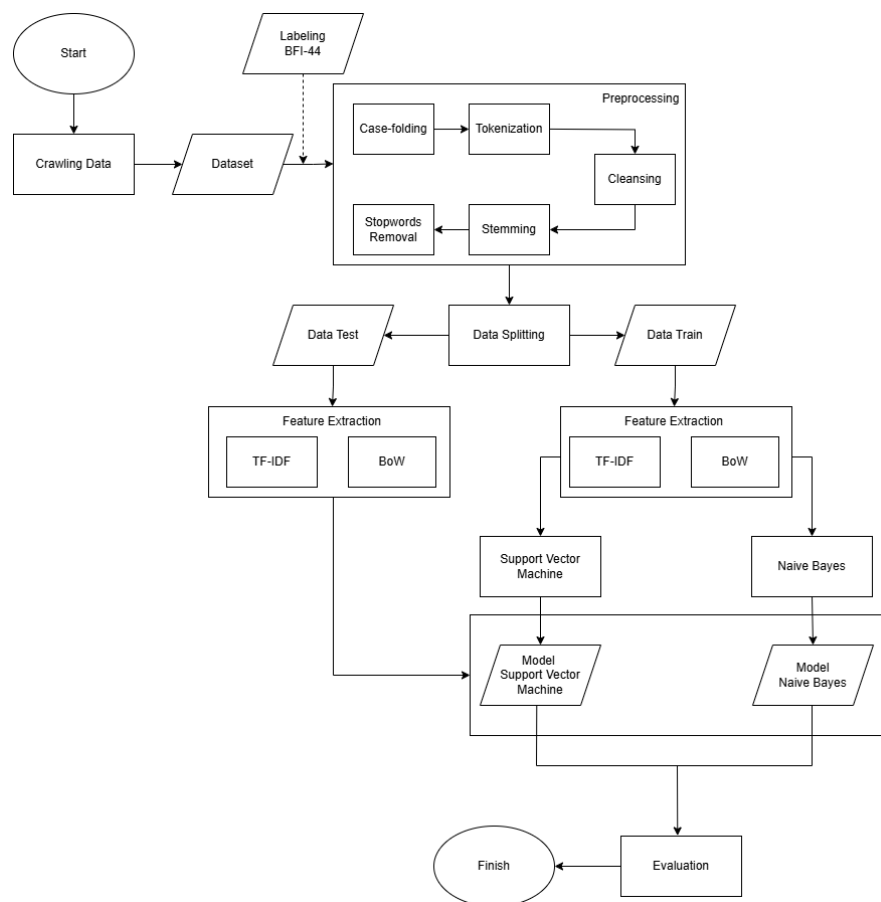


Figure 1. System Flowchart Design

Figure 1 presents a flowchart of system design, covering data collection, preprocessing, feature extraction, model training, and its evaluation. It begins by crawling texts from the social media platform X, wherein posts from people who allow their data to be used have been taken into account. The BFI-44 Big Five Inventory-44 then maps data from each dataset, labeled thereafter, into each one of the Big Five personality traits, including neuroticism, extraversion, openness, agreeableness, and conscientiousness.

The data preprocessing follows in line with the collection of data in the analysis. Here, preprocessing involves some series of case folding to lowercase, tokenization into individual words, cleansing-undesirable characters or symbols and incomplete data, removal of stopwords, and removal of very common and less essential words; stemming is used to get back roots. These steps ready the text data for good extraction of features.

Following preprocessing, data splitting splits the dataset into a training and testing set. These shall be used to develop the machine learning models, while the testing set is needed for performance evaluation of machine learning models. Each of the subsets then undergoes feature extraction, the end product of which will be representations in numerical forms for purposes of machine learning. In this paper, the two feature extraction techniques explored are TF-IDF of terms frequency-inverse document frequency and BoW.