



Feature Selection Analysis for Diagnosing Narcissistic Personality Disorder (NPD) Using Principal Component Analysis and the Naïve Bayes Model

¹Sarwadi 

Department of Computer Science, Universitas Potensi Utama, Medan, Indonesia

²Rika Rosnelly 

Department of Computer Science, Universitas Potensi Utama, Medan, Indonesia

³Budi Triandi 

Department of Computer Science, Universitas Potensi Utama, Medan, Indonesia

Article Info

Article history:

Accepted, 28 May 2025

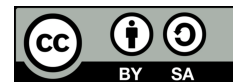
Keywords:

Classification Models;
Clinical Diagnosis;
Feature Selection;
Gaussian Naïve Bayes;
Machine Learning;
Mental Health;
Narcissistic Personality Disorder;
Principal Component Analysis;
Psychological Assessment;
Predictive Modeling.

ABSTRACT

The mental health illness known as narcissistic personality disorder (NPD) affects a person's capacity to preserve harmonious social interactions. Early diagnosis plays a crucial role in providing timely intervention and treatment. This study examines the effectiveness of Principal Component Analysis (PCA) for feature selection in diagnosing NPD using the Naïve Bayes algorithm. The dataset utilized in this research was sourced from Open Psychometrics via Kaggle, followed by preprocessing, including data cleaning and dimensionality reduction through PCA. Feature selection reduced the number of attributes from 44 to 10 principal features. Three types of Naïve Bayes models—Gaussian, Bernoulli, and Multinomial—were compared to identify the most suitable classification approach. The findings reveal that Gaussian Naïve Bayes combined with PCA achieved the % classification accuracy of 91%, significantly outperforming Bernoulli Naïve Bayes (80%) and Multinomial Naïve Bayes (69%). Moreover, the computational time was only 0.01 seconds, indicating a substantial improvement in processing efficiency. These results suggest that integrating PCA with Gaussian Naïve Bayes improves accuracy and reduces irrelevant data, leading to faster diagnosis. This study demonstrates the potential of machine learning for automating mental health evaluation, particularly for personality disorders such as NPD. It provides a foundation for future exploration using hybrid models to enhance predictive accuracy.

This is an open access article under the [CC BY-SA](#) license.



Corresponding Author:

Sarwadi,
Department of Computer Science,
Universitas Potensi Utama Medan, Indonesia
Email: sarwadi69@gmail.com

1. INTRODUCTION

Narcissistic Personality Disorder (NPD) is one of ten psychological disorders identified in mental health studies. People with NPD frequently display an overwhelming feeling of self-importance and superiority towards others. They may display irrational behaviors to maintain their image of superiority [1]. Based on personality disorder classifications, NPD is categorized under Cluster B, which encompasses Antisocial, Borderline, Histrionic, and narcissistic personality disorders. This category is characterized by dramatic, emotional, and unstable behavior [2].

Given the negative impacts of NPD, early diagnosis becomes crucial. Early identification of NPD is also essential to prevent comorbid conditions such as depression, anxiety, and substance abuse, which are often reported in clinical populations [3], to ensure timely intervention [4]. Individuals with NPD are often obsessed with their achievements and abilities, considering themselves exceptional and struggling to show empathy toward others. This condition negatively impacts their ability to build healthy social relationships in personal, social, and professional environments [5]. Therefore, early diagnosis and more effective detection methods are crucial to ensure that individuals with NPD receive appropriate intervention. Alternative approaches utilising technical developments, especially in the areas of artificial intelligence and machine learning, are required in light of these constraints [6]. The difficulty of treating this condition is made worse in Indonesia by the scarcity of mental health specialists. According to data from the Indonesian Ministry of Health [7], only 1,053 psychiatrists serve the entire country, meaning that each psychiatrist must cater to approximately 250,000 people. Additionally, the reliance on interviews and verbal patient reports as the primary diagnostic method poses a challenge in achieving fast and accurate diagnoses. Advances in machine learning and artificial intelligence enable new methods. The integration of AI and ML in mental health diagnostics has shown great promise in improving scalability and personalization of care [8], to more objectively and methodically increase the accuracy and efficacy of NPD detection [9].

However, implementing machine learning brings challenges, especially in selecting the most relevant features. A significant advancement in data processing and computer science, machine learning enhances several services, including healthcare [10]. Through machine learning, hidden patterns in data can be uncovered, facilitating the identification and prediction of mental health conditions [11]. However, selecting relevant features or variables is a key challenge in machine learning implementation. Datasets used in machine learning often contain many irrelevant or redundant features, which can reduce model performance. Therefore, proper feature selection is essential to ensure optimal prediction results [12].

Principal Component Analysis (PCA) is one of the feature selection methods. PCA has been successfully applied in psychometric evaluations to simplify datasets and uncover latent psychological constructs [13]. That have been put forth to deal with this problem. Feature selection aims to prevent problems such as overfitting, improve system performance, and accelerate model training. Large datasets often include unnecessary features that do not contribute significantly to classification [14]. There are numerous feature selection strategies accessible, such as filter, wrapper, and embedding approaches. Several methods are commonly applied, including Principal Component Analysis (PCA), genetic algorithms, backward elimination, and forward selection. Each method has advantages and disadvantages, and its application depends on dataset characteristics and feature selection objectives [15]. PCA, for instance, is widely utilized for dimensionality reduction while preserving essential information by projecting data onto a new coordinate system where variables remain uncorrelated [16]. Prior studies show that PCA effectively improves classification performance in mental health screening tasks, such as autism spectrum disorder detection [17]. Various feature extraction techniques, including PCA, have been tested in psychological disorder classifications, showing promising improvements in model accuracy [18].

Although PCA has been widely applied in mental health classification, its use in diagnosing NPD remains explicitly underexplored. Previous research has shown that applying PCA in IDS systems can enhance attack detection efficiency while preserving the key information in the dataset. The implementation of PCA on the UNSW-NB15 dataset has been proven to reduce data dimensionality to 30 principal components without losing significant information. Combining this method with the Naïve Bayes model resulted in an accuracy of 96.65%, with a perfect recall (1.00) for the threat class, although precision remained limited (0.49). Additionally, using PCA contributed to reducing computation time from 1 minute to 30 seconds, thereby improving the efficiency of the Intrusion Detection System (IDS) [19]. Additionally, even with its ease of use and efficiency, Naïve Bayes classifiers' performance is heavily reliant on the features used, especially when working with high-dimensional data [20]. The three main categories of Naïve Bayes classifiers are Gaussian, Bernoulli, and Multinomial. Gaussian Naïve Bayes is appropriate for continuous data, such as test scores or measurement values, and requires that the features have a normal distribution. Bernoulli Naïve Bayes is optimized for binary or boolean features, making it effective when the presence or absence of a feature (e.g., yes/no, true/false) is the primary consideration. In text classification problems, where features indicate the frequency of categorical events, like word occurrences, multinomial Naïve Bayes is frequently utilised for count data [21].

Each of these variants applies Bayes' theorem differently according to the assumed distribution of input data. Therefore, depending on the structure and features of the dataset, choosing the right variation has a big impact on the classification results [22]. Understanding how PCA influences the performance of these Naïve Bayes variants is critical in optimizing predictive models for mental health disorders such as NPD. Given these considerations and the lack of specific studies targeting NPD using PCA-optimized Naïve Bayes models, a focused investigation is warranted. Therefore, this study aims to investigate how PCA can optimize the performance of various Naïve Bayes classifiers in detecting NPD.

Previous studies on NPD diagnosis have primarily relied on traditional clinical interviews or basic statistical methods, resulting in diagnostic accuracy and generalizability limitations. Most of these methods depend on subjective clinician assessments, prone to bias and human error [23]. Additionally, statistical models used in earlier studies often fail to capture complex, nonlinear relationships within high-dimensional psychological data [24].

These limitations emphasize the need for more advanced computational approaches to extract meaningful patterns from psychological assessments and improve the robustness and reproducibility of diagnostic models.

With the increasing number of personality disorder cases and the limited availability of mental health professionals, technology-based solutions are needed to aid the diagnostic process [25]. Machine learning offers a faster and more objective approach to identifying patients at risk of NPD. However, challenges remain in determining the most relevant features to enhance model accuracy.

While previous research has successfully applied PCA for feature selection in psychological disorder classification, its specific role in diagnosing NPD remains underexplored [26]. Much of the existing research focuses on disorders like depression and autism spectrum disorder, where dimensionality reduction has improved classification accuracy and efficiency [27]. However, there is a lack of studies that directly evaluate the effectiveness of PCA in optimizing Naïve Bayes classifiers for NPD detection.

Naïve Bayes is known for its efficiency in probabilistic classification, but its performance can be sensitive to feature selection methods [28]. Comparative studies of Gaussian, Bernoulli, and Multinomial Naïve Bayes for NPD classification are still limited. Understanding how PCA affects each of these variants can provide valuable insights into optimizing diagnostic models for NPD.

This study aims to fill this gap by systematically evaluating PCA as a feature selection method in NPD classification using Naïve Bayes models. By comparing the performance of Gaussian, Bernoulli, and Multinomial Naïve Bayes, this research contributes to developing a more efficient and accurate machine learning approach for early NPD diagnosis. The findings may also lay the groundwork for future research on hybrid models or alternative feature selection strategies to enhance classification performance further.

2. RESEARCH METHOD

This work employs a machine learning methodology to diagnose Narcissistic Personality Disorder (NPD) through a quantitative approach. The main objective is to build a predictive model that can classify individuals based on questionnaire responses, incorporating feature selection techniques to improve model accuracy and interpretability.

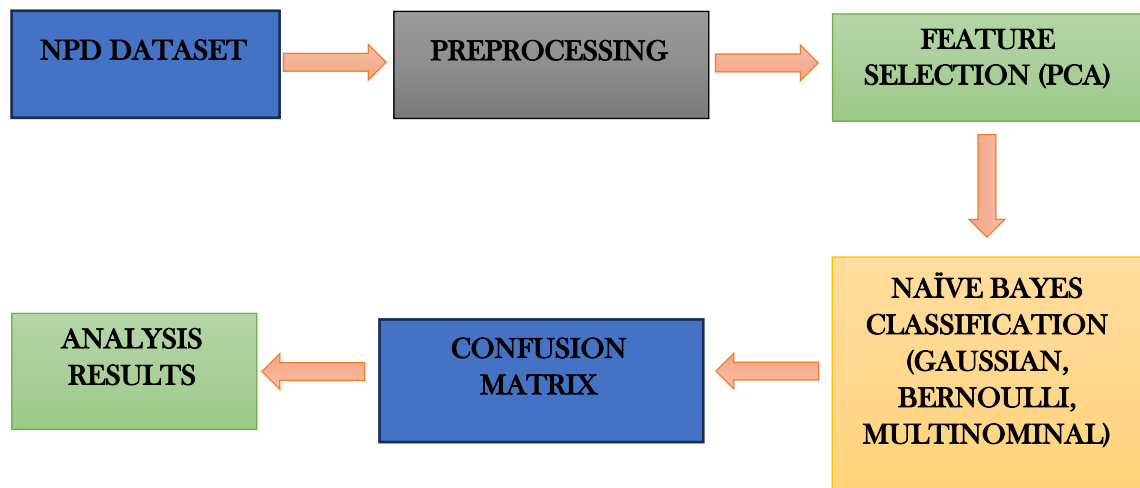


Figure 1. Research process stages

2.1. Data

The dataset used in this research was obtained from the [Open Psychometrics](#) dataset available on Kaggle. The data includes answers to 40 questions related to narcissistic personality traits, along with demographic information such as gender and age. The total number of records in the dataset is 4338. Data was collected through secondary sources, utilizing publicly available psychological assessment data. The dataset was downloaded and stored in CSV format for preprocessing and further analysis.

2.2. Preprocessing

This phase consists of multiple steps designed to refine and structure the data for practical analysis. The preprocessing procedure includes:

1. Irrelevant Gender Data Removal:
2. The initial dataset of 11,243 respondents was filtered to include only Male (1) and Female (2) categories. Fifty-two records with the Other (3) or Unspecified (0) category were removed, reducing the dataset to 11,191.
3. Removal of Data Outside the Age Range (14–50 Years):
4. The dataset was focused on respondents aged 14–50, based on research evidence indicating the prevalence of personality disorders within this range. Of the 10,391 respondents, 1,517 records outside this age range were removed, resulting in 8,874 records.

5. Removal of Missing Values (NaN):
6. Data in the Q1-Q40 question columns with a value of 0 were removed to maintain data quality. From 11,191 respondents, 800 records were deleted, leaving 10,391 records.

2.3. Feature Selection

Principal Component Analysis (PCA) is employed in this study to minimize data dimensionality and improve model performance and accuracy [29]. The feature selection process using PCA involves the following steps:

1. Standard Scaling / Z-score Normalization

To prevent bias that could arise from differences in the range or scale of features, data standardization is used. In order to adjust each characteristic, the mean is subtracted and then divided by the standard deviation using the following formula:

$$X(\text{new}) = \frac{Xi - \mu}{\sigma} \quad (1)$$

Here, (X) the new standardized value of feature Xi , (Xi) Original value of the feature (μ) denotes the mean of all feature values, and (σ) represents the standard deviation of the feature values. To calculate the standard deviation (σ) , the following formula is used:

$$\sigma^2 = \frac{\sum (X - \mu)^2}{N} \quad (2)$$

2. Calculating the Covariance-Variance Matrix variance measures data dispersion, while covariance measures the relationship between variables. The covariance matrix is used to compute the relationships between features in the dataset. The formula below is used to compute variance:

$$Var(x) = \sigma^2 = \frac{1}{n} \sum_{i=1}^n (Z_{ij} - \mu_j)^2 \quad (3)$$

The covariance between two features, x and y , is calculated using the following equation:

$$Cov(x, y) = \frac{1}{n-1} \sum_{i=1}^n (x_{ij} - \mu_{xj})(y_{ij} - \mu_{yj}) \quad (4)$$

3. Calculating Eigenvalues and Eigenvectors: The covariance matrix is used to compute the eigenvalues (λ) and eigenvectors, which help transform the dataset into a new feature space.
4. In determining the number of principal components (K) using Kaiser's stopping rule, all principal components (PC) with an eigenvalue greater than 1 are selected. The value of " K " determines how many principal components should be retained.
5. Transforming Data into the New Feature Space: The projection matrix (PM) is obtained by selecting the top " K " eigenvectors.
6. Result Interpretation: The newly dimension-reduced dataset can be utilized after the transformation. The selected features have the highest eigenvector weights in the principal components that exhibit significant correlations.

2.4. Naïve Bayes Classification

At this stage, the Naïve Bayes classification method analyses the dataset and predicts the class (NPD or Normal) based on the available features. This algorithm assumes independence between features. The steps involved are as follows [30].

1. Data Splitting: The dataset is divided into two sets:
 - Training set - Used to train the model.
 - Testing set - Used to evaluate model performance.
2. Probability Calculation: The probability of each class is calculated using Bayes' Theorem. The probability $P(H)$ is determined based on the dataset, where:
 - (H) represents the hypothesis (the target class).
 - (X) represents the features in the dataset.

$$P(H|X) = \frac{P(X|H).P(H)}{P(X)} \quad (5)$$

Types of Naïve Bayes Classifiers in Machine Learning. Three varieties of Naïve Bayes classifiers are frequently employed in machine learning and data science training [31]:

1. Gaussian Naïve Bayes. Assumes the data features adhere to a normal distribution pattern.
2. Bernoulli Naïve Bayes. Ideal for qualities that are binary or Boolean.
3. Multinomial Naïve Bayes - Used for discrete count-based data, such as text classification.

2.5. Evaluation

The performance of the constructed classification method is evaluated in this study utilizing a Confusion Matrix. The Confusion Matrix gives specific details on how the model categorizes each class and how many predictions were right and wrong. Based on the Confusion Matrix, key metrics can be calculated to evaluate model performance, including Accuracy, Precision, Recall, and F1-score. The Confusion Matrix, as shown in Table 1, is a crucial tool for assessing the classification algorithm's performance [32].

Table 1. Confusion Matrix for Binary Classification Results

	Predicted Positives	Predicted Negatives
Actual Positives	True Positives (TP)	False Negatives (FN)
Actual Negatives	False Positives (FP)	True Negatives (TN)

TP: Correctly predicted positive cases. FP: Incorrectly predicted as positive. FN: Missed positive cases. TN: Correctly predicted negative cases.

The following formulas can be used to calculate each of these evaluation metrics. Calculation of Accuracy. The percentage of instances (both positive and negative) that are correctly identified out of all instances is known as accuracy. The formula below is used to compute it:

$$Accuracy = \frac{(TP+TN)}{(TP+TN+FP+FN)} \quad (6)$$

The precision metric quantifies the percentage of accurately predicted NPD cases (TP) among all cases predicted as NPD (TP + FP). It is computed using the formula that follows:

$$Precision = \frac{TP}{TP+FP} \quad (7)$$

Recall quantifies the percentage of real NPD cases (TP) that were accurately recognised. Other names for it include true positive rate and sensitivity. It is calculated as follows:

$$Recall = \frac{TP}{TP+FN} \quad (8)$$

By considering the harmonic mean of precision and recall, the F1-Score strikes a balance between the two. It is computed using the formula below:

$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (9)$$

2.6. Tools Used

This study utilizes the Python programming language for its extensive data analysis and machine learning capabilities. The main libraries used include:

1. Pandas: For data manipulation and analysis.
2. NumPy: For numerical computation and array processing.
3. Scikit-learn: For data preprocessing, PCA implementation, Naïve Bayes modeling, and evaluation metrics.
4. Matplotlib and Seaborn: For data visualization and result interpretation.

All processes were carried out using Google Colaboratory (Google Colab), a cloud-based platform that supports Python programming with access to GPU/TPU for faster computation. It also allows easy collaboration and integration with Google Drive for dataset storage.

3. RESULT AND ANALYSIS

The following section presents the results and discusses each step of this study.

3.1. Data

Times This study utilizes secondary data obtained from the www.kaggle.com website, specifically a dataset related to Narcissistic Personality Disorder (NPD), which is publicly available through Open Psychometrics. The collected dataset is stored in CSV format, allowing for further processing. The dataset consists of 11,243 records, which serve as the basis for analysis. A dataset visualisation can be found below. The goal of analysis is to describe the data in terms of variables.

	gender	age	Q1	Q2	Q3	Q4	Q5	Q6	Q7	Q8	...	Q34	Q35	Q36	Q37	Q38	Q39	Q40	elapsed	Score	Class
0	1	50	2	2	2	2	1	2	1	2	...	1	1	2	2	2	1	2	211	18	no
1	1	40	2	2	2	1	2	2	1	2	...	2	1	2	2	2	2	1	149	6	yes
2	1	28	1	2	2	1	2	1	2	1	...	1	2	1	1	2	1	2	168	27	yes
3	1	37	1	1	2	2	2	1	2	1	...	1	2	1	2	2	1	1	230	29	yes
4	1	50	1	2	1	1	1	2	1	2	...	2	1	2	2	2	0	1	389	6	no
...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...
8335	2	45	2	2	2	1	1	2	1	1	...	2	1	2	1	1	1	2	643	10	no
8336	1	19	1	1	2	1	2	1	2	1	...	2	1	2	2	2	1	2	226	23	yes
8337	2	36	1	2	2	1	1	2	1	2	...	2	1	2	2	2	2	1	232	3	no
8338	2	21	2	2	2	2	1	2	2	1	...	2	2	2	1	2	2	1	521	15	yes
8339	1	23	1	2	2	1	2	1	2	2	...	2	1	2	2	2	2	1	254	11	no

8340 rows × 45 columns

Figure 2. Dataset

3.2. Preprocessing

The preprocessing stage ensures data quality and relevance before analysis. First, irrelevant gender data was removed by filtering the dataset to include only respondents categorized as Male (1) and Female (2). Fifty-two records labelled Other (3) or None/Unspecified (0) were eliminated, reducing the dataset from 11,243 to 11,191. Next, missing values (NaN) were addressed by removing responses in the Q1-Q40 columns with a value of 0, resulting in the deletion of 800 records and leaving 10,391 records. Lastly, to focus the analysis on the most affected age group, data outside the age range of 14–50 years was excluded based on research findings indicating the prevalence of personality disorders within this range. This process removed 1,517 records, bringing the final dataset to 8,874 for further analysis.

3.3. Feature Selection

This section discusses Feature Selection, specifically the application of Principal Component Analysis (PCA) to enhance model efficiency and accuracy. Key points include:

- PCA is used for dimensionality reduction while preserving significant features.
- It helps optimize model performance by minimizing multicollinearity.
- Feature selection improves model efficiency by identifying important features for analysis.

The key stages in PCA implementation include:

- a. Data Standardization – To guarantee that every feature contributes equally, the data scale should be adjusted. The results of Z-score Normalization can be seen in Figure 3.

```
Data setelah scaling:
[[ 1.17182268  0.12434312 -0.77907958  0.53878235  0.46843554 -0.46143484
 -0.88131241  0.79594949 -0.67365855 -0.77686463 -0.95127718 -1.01367193
 0.88846152  0.88978645  0.72873654  0.59831991 -0.54329696  0.94735114
 -0.59484193  1.34381486 -0.64936354 -0.68594912  0.53465851 -0.77722802
 -0.43458083  0.68070134  0.55927696 -1.03456697  0.61828651 -0.51494996
 0.93222798  0.78447206  0.88237212 -0.63628824  0.91974195  0.74847556
 -0.59426762  0.48552132  0.68322885  0.43954551  0.80508096 -0.86532384
 2.6612697 -1.27548763]
 [ 1.17182268  1.81645298  1.27246151 -1.81388638  0.46843554 -0.46143484
 -0.88131241 -1.24838478 -0.67365855 -0.77686463  1.03801293 -0.95761946
 -1.10111974 -1.09380263 -1.34487952 -1.61523626 -0.54329696 -1.03566642
 1.62981355  1.34381486 -0.64936354  1.59812005 -1.01844777  1.25560999
 -0.43458083  0.68070134  0.55927696  1.03456697 -1.57738756  1.88814985
 0.93222798  0.78447206 -1.10271813  1.58964929 -1.0623607 -1.32099162
 1.59182839  0.48552132  0.68322885  0.43954551 -1.20941124  1.12277556
 -0.07835808  1.19621445]
 [ 1.17182268  2.02796671 -0.77907958  0.53878235  0.46843554 -0.46143484
 -0.88131241  0.79594949 -0.67365855 -0.77686463  1.03801293 -1.01367193
 0.88846152  0.88978645  0.72873654  0.59831991 -0.54329696 -1.03566642
 1.62981355  1.34381486  1.50893954 -0.68594912  0.53465851  1.25560999
 2.24449839  0.68070134  0.55927696  0.95220999  0.61828651  1.88814985
 0.93222798  0.78447206 -1.10271813  1.58964929 -1.0623607 -1.32099162
 1.59182839  0.48552132  0.68322885  0.43954551  0.80508096  1.12277556
 -0.04425191  0.37231376]
 [-0.05337071  0.33558585 -0.77907958 -1.11380656  0.46843554  2.09440633
 1.11911599  0.79594949  1.45449805 -0.77686463 -0.95127718  0.95761946
 -1.10111974 -1.09380263 -1.34487952  0.59831991  1.79576474  0.94735114
 1.62981355 -0.72451922  1.50893954  1.59812005  0.53465851  1.25560999
 -0.43458083 -1.43380357  0.55927696  0.95220999 -1.57738756 -0.51494996
 -1.04538824 -1.36533563  0.88237212 -0.63628824 -1.0623607 -1.32099162
 -0.59426762 -1.99607979  0.68322885  0.43954551  0.80508096  1.12277556
 -0.40590882  1.19621445]
 [ 1.17182268 -1.14473927 -0.77907958  0.53878235  0.46843554 -0.46143484
 1.11911599  0.79594949 -0.67365855 -0.77686463  1.03801293  0.95761946
 0.88846152  0.88978645  0.72873654  0.59831991 -0.54329696  0.94735114
 1.62981355 -0.72451922 -0.64936354 -0.68594912  0.53465851 -0.77722802
 -0.43458083  0.68070134  0.55927696 -1.03456697  0.61828651 -0.51494996
 0.93222798  0.78447206  0.88237212 -0.63628824 -1.0623607 -1.32099162
 -0.59426762  0.48552132  0.68322885  0.43954551  0.80508096 -0.86532384
 -0.33715738 -0.68698713]]
```

Figure 3. Output Matrix After Z-score Standardization of Feature Values

This figure illustrates the outcome of Z-score normalization, demonstrating how the data has been standardized to ensure that all features contribute equally by transforming them to a common scale with a mean of 0 and a standard deviation of 1.

- b. Covariance Matrix – Measures the relationships between features to understand data variance.
- c. Eigenvalues and Eigenvectors – Identifies principal components based on their contribution to the total variance.
- d. Determining the Number of Principal Components – Uses Kaiser's Rule to select components with eigenvalues > 1.
- e. Creating the Projection Matrix – Combines the selected principal components to form a projection matrix.
- f. Data Transformation – Transforms the original dataset into a new feature space using the projection matrix.

As shown in Figure 4, the outcome of the Principal Component Analysis (PCA) feature selection procedure is as follows.

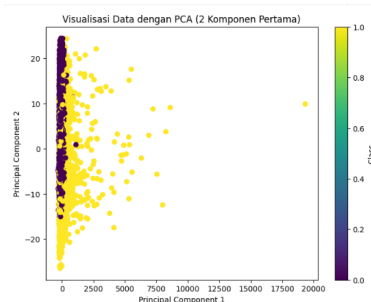


Figure 4. Visual Representation of PCA Feature Selection Results

The visualization results display the PCA transformation using two principal components (PC1 and PC2). PCA functions to reduce data dimensionality while retaining relevant information for classifying between the two classes: NPD (Yellow) and Normal (Purple). After conducting a Principal Component Analysis (PCA), the 10 most influential features in the principal components were identified. The following data presents the contribution values of each feature to the corresponding principal components, as illustrated in the following figure. 5.

	Feature	Contribution	Component
0	elapse	1.000000	1
1	age	0.000203	1
2	Q29	0.000037	1
3	Q19	0.000031	1
4	Q24	0.000031	1
..	...	...	...
5	Q7	0.222764	10
6	Q15	0.215078	10
7	Q40	0.178516	10
8	Q8	0.174519	10
9	Q16	0.170863	10

[100 rows x 3 columns]

Figure 5. Top 10 Most Influential Features based on PCA.

Principal Component 1 indicates that the "elapse" feature contributes the most, with a value of 1.000000, followed by other features such as age (0.000203), Q29 (0.000037), Q19 (0.000031), and Q24 (0.000031). This finding suggests that completion time (elapse) is the most significant factor in forming the first principal component. In Principal Component 10, the feature with the highest contribution is Q7 (0.222764), followed by Q15 (0.215078), Q40 (0.178516), Q8 (0.174519), and Q16 (0.170863). These features exhibit a significant pattern in shaping the tenth principal component.

Principal Component Analysis (PCA) was used in this study to choose features. This process successfully reduced the initial set of 44 features, which included demographic factors and 40 questionnaire questions, to a more manageable set of 10 primary features. These selected features included *elapse*, *age*, *Q29*, *Q19*, *Q24*, *Q7*, *Q15*, *Q40*, *Q8*, and *Q16*. Notably, several of these features directly relate to the clinical characteristics of narcissistic personality disorder (NPD), as outlined in the Diagnostic and Statistical Manual of Mental Disorders (DSM). The DSM states that a person is diagnosed with NPD if they meet at least five of the nine criteria, which include a lack of empathy, a sense of entitlement, a need for continual praise, and a sense of superiority [33].

The features selected through PCA are strongly aligned with these narcissistic traits. For example, *Q29*, which asks about the individual's interest in looking at themselves in the mirror, and *Q19*, which asks about their perception of their own body, reflect the need for self-admiration and external validation—key behaviors commonly associated with narcissism. Additionally, questions such as *Q7*, which relates to the individual's desire to be the center of attention, and *Q15*, which explores the tendency to show off one's body, further exemplify attention-seeking and exhibitionistic tendencies often seen in individuals with NPD. *Q24* and *Q40*, which pertain to expectations of special treatment and a grandiose self-image, also align with the narcissistic traits of entitlement and grandiosity. Furthermore, *Q16*, which addresses the ability to read people's emotions, may hint at the interpersonal manipulation often observed in individuals with narcissistic behaviors.

These selected features highlight the psychological markers of NPD and reinforce the PCA technique's psychological relevance in feature selection. By isolating these traits, PCA has allowed for a more efficient and focused analysis, directly contributing to the model's ability to predict the disorder accurately.

3.4. Naïve Bayes Classification

This study focuses on implementing the Naïve Bayes classification algorithm to predict respondent categories based on the preprocessed data. The research compares the algorithm's performance in two scenarios: using PCA as a feature selection technique and without PCA. This approach aims to evaluate the impact of dimensionality reduction on classification accuracy and efficiency.

The initial step in the Naïve Bayes classification process involves splitting the dataset into two main parts: training data (70%) and testing data (30%). This partitioning uses stratified sampling to ensure that the class distribution remains balanced between both datasets. This data-splitting process is essential to maintain the validity and reliability of the classification results. By preserving a balanced class proportion, the model can learn optimally from the training data while being tested on a dataset that accurately represents real-world scenarios. The success of this process is crucial in ensuring that the results obtained reflect the actual performance of the model. After dividing the dataset, the next step is to initialize the classification models, which include Gaussian Naïve Bayes, Bernoulli Naïve Bayes, and Multinomial Naïve Bayes.

The Gaussian Naïve Bayes algorithm is employed when the dataset's characteristics are continuous and exhibit a normal (Gaussian) distribution. This algorithm operates based on Bayes' theorem, assuming that each feature is independent of the others (naïve assumption). The probability distribution of the features is calculated using the Gaussian probability distribution function, making this algorithm highly suitable for continuous data that exhibits typical distribution characteristics. Meanwhile, Bernoulli Naïve Bayes is applied to binary data, where

variables have only two possible values (e.g., 0 or 1). This model is commonly used in cases where the presence or absence of a feature is essential for classification. On the other hand, Multinomial Naïve Bayes is suitable for count-based data, where features represent the frequency of occurrences within a given sample. This method is frequently applied to jobs involving text categorisation, where characteristics are correlated with the frequency of a word's occurrences in a document.

The selection of the Naïve Bayes variant used in this study depends on the dataset's characteristics. By applying all three methods, this study aims to evaluate the performance of each algorithm in classifying respondent categories based on the available features. The last phase of this study is the model evaluation procedure, which evaluates the effectiveness of the model and the algorithms employed. A confusion matrix and a categorisation report are used in the evaluation.

A thorough summary of the model's classification performance is given in the classification report, which includes important assessment metrics for every class, including precision, recall, F1-score, and support. Recall evaluates how well the model captures all positive cases, and precision gauges how accurate optimistic projections are. F1-score is the harmonic mean of precision and recall. Especially in clinical datasets with class imbalance, F1-score provides a better metric for model reliability than accuracy alone [34].

The model's prediction outcomes are visually represented by the confusion matrix, which contrasts accurate predictions (True Positive and True Negative) with inaccurate predictions (False Positive and False Negative). This matrix helps researchers gain deeper insights into how the model classifies data in each class, allowing for the identification of areas that require improvement.

By combining these two evaluation methods, the study ensures a more comprehensive and accurate assessment of the model's performance.

Table 2. Comparison of Naïve Bayes Model Accuracy with and without PCA.

Model	Accuracy Without PCA	Accuracy With PCA
Gaussian Naïve Bayes	88%	91%
Multinomial Naïve Bayes	76%	69%
Bernoulli Naïve Bayes	68%	80%

This table compares the classification accuracy of the three Naïve Bayes models with and without Principal Component Analysis (PCA). The results indicate that Gaussian Naïve Bayes performs best with PCA, increasing accuracy from 88% to 91%.

Bernoulli Naïve Bayes also benefits from PCA, improving from 68% to 80%. However, Multinomial Naïve Bayes experiences a decrease in accuracy, dropping from 76% to 69% when PCA is applied, suggesting that PCA may not be suitable for this type of data representation.

The decrease in accuracy for Multinomial Naïve Bayes (76% to 69%) suggests that PCA may not be effective for this model due to its reliance on discrete feature distributions. Unlike Gaussian Naïve Bayes, which benefits from PCA's continuous transformation, Multinomial Naïve Bayes works better with raw frequency-based categorical data, where dimensionality reduction might remove critical information needed for classification [35]. This emphasises how crucial it is to choose the best feature selection methods depending on the underlying properties of the data.

Beyond accuracy, performance was also assessed using precision, recall, and F1-score, as shown in the classification report. These metrics are critical in medical or psychological diagnosis: Precision reflects how many predicted NPD cases were correct, helping to avoid false positives. Recall (Sensitivity) measures the model's ability to detect actual NPD cases, which is crucial for minimizing false negatives. F1-score balances both precision and recall, making it ideal for imbalanced datasets.

Table 3. Evaluation Metrics of Naïve Bayes Classifiers Based on Precision, Recall, and F1-Score

Model	Pesisi	Recall	F-1 Score
Gaussian Naïve Bayes	90%	91%	90%
Bernoulli Naïve Bayes	83%	80%	81%
Multinomial Naïve Bayes	47%	69%	56%

This table presents the classification performance of three Naïve Bayes models—Gaussian, Bernoulli, and Multinomial—based on precision, recall, and F1-score. The Gaussian Naïve Bayes model achieved the highest scores across all metrics, indicating its superior suitability for NPD diagnosis.

The Gaussian Naïve Bayes classifier showed superior performance across all evaluation metrics. It was especially effective in consistently classifying both positive and negative cases. Meanwhile, Multinomial Naïve Bayes performed poorly in classifying the positive class (NPD), as reflected in its low precision and F1-score, rendering it less suitable for clinical application.

This study presents a more efficient approach than a previous study employing Information Gain and Gain Ratio for feature selection [36]. used 37–38 features and achieved a maximum accuracy of 91% with the Naïve Bayes classifier. In contrast, this study reduced the feature set to only 10 principal components using PCA while maintaining the same classification accuracy. Additionally, the processing time dropped significantly — from 0.22 seconds to 0.01 seconds — indicating that PCA retains predictive power and improves computational efficiency.

These results align with findings from [37] and [38], who highlight that dimensionality reduction techniques like PCA can enhance performance by removing noise and redundancy in high-dimensional psychological data.

Despite these encouraging outcomes, several limitations should be acknowledged. The dataset was based on self-reported assessments, which may lack the objectivity of clinical diagnoses and introduce potential bias. Moreover, the evaluation was conducted on a single dataset without external validation, limiting the generalizability of the findings. Finally, the study focused solely on Naïve Bayes classifiers and did not explore more complex models such as deep learning or ensemble approaches.

Future research should explore testing the PCA-Naïve Bayes framework on diverse datasets, including clinically verified cases. Additionally, combining PCA with more advanced models could improve prediction performance. Applying the model to real-world mental health platforms may also enhance early intervention strategies and increase accessibility to psychological care.

The classification report obtained from testing using the Naïve Bayes algorithm is attached. This report presents key evaluation metrics, including precision, recall, F1-score, and support, for assessing the model's classification performance.

The Gaussian Naïve Bayes algorithm demonstrated the best performance, achieving an accuracy of 91%, followed by Bernoulli Naïve Bayes with 80% accuracy. Both algorithms performed well on the tested dataset, particularly in classifying negative cases. In contrast, Multinomial Naïve Bayes recorded the lowest accuracy (69%), indicating that this algorithm is less suitable for the dataset. This may be due to the dataset's characteristics, which are more compatible with Gaussian Naïve Bayes, designed for continuous data. In contrast, Multinomial Naïve Bayes is better suited for text-based or categorical data.

In addition to accuracy results, this study also includes an evaluation using the confusion matrix for each Naïve Bayes algorithm, namely Gaussian Naïve Bayes, Bernoulli Naïve Bayes, and Multinomial Naïve Bayes. A thorough summary of the model's performance is given by the confusion matrix, which also shows how many predictions were right and wrong for each class and how well the model classified positive and negative cases.

Based on the evaluation using the confusion matrix, the Gaussian Naïve Bayes algorithm not only excels in accuracy but also demonstrates more consistent performance in classifying both positive and negative classes compared to the other two algorithms. In contrast, the Multinomial Naïve Bayes algorithm struggles more significantly, particularly in classifying negative cases, as indicated by its higher False Negative (FN) value than the other models.

This study confirms that Gaussian Naïve Bayes achieves the best performance for the tested dataset, followed by Bernoulli Naïve Bayes, while Multinomial Naïve Bayes performs the worst. The confusion matrix evaluation provides additional insights that reinforce the accuracy analysis, allowing researchers to understand each algorithm's strengths and weaknesses better. Future research could extend this analysis by testing different datasets or combining the Naïve Bayes algorithm with other machine learning methods to enhance classification performance further.

Compared to the previous study conducted by [36] the approach used in this research demonstrates comparable classification performance while offering significantly improved computational efficiency. [36] applied the Information Gain and Gain Ratio for feature selection from 44 initial features, including demographic variables and 40 questionnaire items. Their selection process resulted in 37 features using Information Gain and 38 features using Gain Ratio. They employed three classification algorithms—Random Forest, Support Vector Machine (SVM), and Naïve Bayes—with the highest accuracy of 91% achieved using the Naïve Bayes model.

In contrast, this study utilized Principal Component Analysis (PCA) for feature selection, successfully reducing the feature set from 44 to only 10 principal features. Despite the significant reduction in feature count, the model maintained the same classification accuracy of 91% using the Naïve Bayes algorithm. Moreover, the processing time was significantly shorter—only 0.01 seconds—compared to the 0.22 seconds [36]. These findings highlight that PCA not only preserves high predictive performance but also greatly enhances computational efficiency, making it a more effective approach for Diagnosing Narcissistic Personality Disorder (NPD).

4. CONCLUSION

This study explores the effectiveness of Principal Component Analysis (PCA) in feature selection for diagnosing Narcissistic Personality Disorder (NPD) using the Naïve Bayes algorithm. The results show that PCA significantly improves the model's classification accuracy and computational efficiency. Specifically, PCA reduced the original 44 features to only 10 principal components while maintaining a high classification accuracy of 91% using the Gaussian Naïve Bayes model. The processing time was also reduced to just 0.01 seconds, demonstrating the potential of PCA in building efficient diagnostic tools.

These findings support the growing body of research indicating that dimensionality reduction techniques can enhance model accuracy by eliminating noise and redundancy in psychological assessment data [39]. In particular, this aligns with [40], who emphasize the practical value of PCA in streamlining personality disorder diagnosis through more interpretable and scalable models.

From a practical standpoint, this approach could assist professionals in providing faster, more scalable, and more accurate screening tools for personality disorders such as NPD. However, challenges remain in real-world implementation, including the need for high-quality labeled data, model interpretability for clinicians, and ethical concerns regarding automated mental health diagnosis.

Future research should explore hybrid approaches combining PCA with more advanced models such as deep learning or ensemble techniques to improve accuracy and generalizability. Moreover, integrating such diagnostic models into digital mental health platforms can enhance accessibility, support early detection strategies, and contribute to more data-driven psychological care.

5. REFERENCES

- [1] T. Safaria and P. D. Psi, *Psikologi Abnormal: Dasar-Dasar, Teori, dan Aplikasinya*. Uad Press, 2021.
- [2] P. Mitra, T. J. Torrico, and D. Fluyau, "Narcissistic Personality Disorder,," Treasure Island (FL), 2025.
- [3] E. Ronningstam, "Pathological Narcissism and Narcissistic Personality Disorder: Recent Research and Clinical Implications," *Curr. Behav. Neurosci. Reports*, vol. 3, no. 1, pp. 34–42, 2016, doi: 10.1007/s40473-016-0060-y.
- [4] M. E. Raygoza-L., J. H. Orduño-Osuna, R. Jimenez-Sanchez, and F. N. Murrieta-Rico, "Innovative Artificial Intelligence approaches for identifying and managing DSM Cluster B personality disorders in mental health: A case study on the dark triad," *Explor. Psychol. Soc. Innov. Adv. Appl. Mach. Learn.*, pp. 1–20, 2025, doi: 10.4018/979-8-3693-6910-4.ch001.
- [5] D. S. Stolz, A. Vater, B. H. Schott, S. Roepke, F. M. Paulus, and S. Krach, "Reduced frontal cortical tracking of conflict between self-beneficial versus prosocial motives in Narcissistic Personality Disorder,," *NeuroImage. Clin.*, vol. 32, p. 102800, 2021, doi: 10.1016/j.nicl.2021.102800.
- [6] H. Taheri and A. Salimi Beni, "Artificial Intelligence, Machine Learning, and Smart Technologies for Nondestructive Evaluation BT - Handbook of Nondestructive Evaluation 4.0,," N. Meyendorf, N. Ida, R. Singh, and J. Vrana, Eds., Cham: Springer Nature Switzerland, 2025, pp. 1–29. doi: 10.1007/978-3-030-48200-8_70-1.
- [7] R. I. Kemenkes, "Kemenkes Beberkan Masalah Permasalahan Kesehatan Jiwa di Indonesia,," *Diambil dari https://sehatnegeriku.kemkes.go.id/baca/rilis-media/20211007/1338675/kemenkes-beberkan-masalah-permasalahan-kesehatan-jiwa-di-indonesia*, 2021.
- [8] A. B. R. Shatte, D. M. Hutchinson, and S. J. Teague, "Machine learning in mental health: a scoping review of methods and applications," *Psychol. Med.*, vol. 49, no. 9, pp. 1426–1448, 2019, doi: DOI: 10.1017/S0033291719000151.
- [9] T. Jain, A. Jain, P. Hada, H. Kumar, V. Verma, and A. Patni, *Machine Learning Techniques for Prediction of Mental Health*. 2021. doi: 10.1109/ICIRCA51532.2021.9545061.
- [10] M. Soori, B. Arezoo, and R. Dastres, "Digital twin for smart manufacturing, A review,," *Sustain. Manuf. Serv. Econ.*, vol. 2, no. June, p. 100017, 2023, doi: 10.1016/j.smse.2023.100017.
- [11] J. Chung and J. Teo, "Mental Health Prediction Using Machine Learning: Taxonomy, Applications, and Challenges,," *Appl. Comput. Intell. Soft Comput.*, vol. 2022, no. 1, p. 9970363, Jan. 2022, doi: https://doi.org/10.1155/2022/9970363.
- [12] A. Taha, B. Cosgrave, and S. McKeever, "Using Feature Selection with Machine Learning for Generation of Insurance Insights,," *Appl. Sci.*, vol. 12, no. 6, 2022, doi: 10.3390/app12063209.
- [13] B. Dettmar, C. Peltier, and P. Schlich, "Beyond principal component analysis (PCA) of product means: Toward a psychometric view on sensory profiling data,," *J. Sens. Stud.*, vol. 35, no. 2, p. e12555, Apr. 2020, doi: https://doi.org/10.1111/joss.12555.
- [14] S. J. Yim *et al.*, "The utility of smartphone-based, ecological momentary assessment for depressive symptoms,," *J. Affect. Disord.*, vol. 274, pp. 602–609, 2020, doi: https://doi.org/10.1016/j.jad.2020.05.116.
- [15] A. Qtaish, M. Braik, D. Albashish, M. T. Alshammari, A. Alreshidi, and E. J. Alreshidi, "Enhanced coati optimization algorithm using elite opposition-based learning and adaptive search mechanism for feature selection,," *Int. J. Mach. Learn. Cybern.*, vol. 16, no. 1, pp. 361–394, 2025, doi: 10.1007/s13042-024-02222-3.
- [16] R. Adhao and V. Pachghare, "Feature selection using principal component analysis and genetic algorithm,," *J. Discret. Math. Sci. Cryptogr.*, vol. 23, no. 2, pp. 595–602, Feb. 2020, doi: 10.1080/09720529.2020.1729507.
- [17] L. Amatullah, Y. Widiastiwi, and N. Chamidah, "Penerapan Klasifikasi Random Forest Terhadap Data Gangguan Spektrum Autisme (ASD) Pada Anak-Anak Menggunakan Seleksi Fitur Principal Component Analysis,," *Semin. Nas. Mhs. Ilmu Komput. dan Apl.*, pp. 356–364, 2022.
- [18] A. Javeed, P. Anderberg, A. N. Ghazi, A. Javeed, A. L. Dallora, and J. S. Berghlund, "Optimizing Depression Prediction in Older Adults: A Comparative Study of Feature Extraction and Machine Learning Models,," *2024 Int. Conf. Control. Autom. Diagnosis, ICCAD 2024*, 2024, doi: 10.1109/ICCAD60883.2024.10553890.
- [19] K. S. Arlandy *et al.*, "Mengoptimalkan Kinerja Naïve Bayes Pada Ancaman Modern Dengan Menggunakan PCA Pada Data Intrusion Detection System (IDS),," vol. 8, no. 1, 2025.
- [20] J. S. Aguilar-Ruiz, C. Romero, and A. Cicconardi, "XNB: Explainable Class-Specific Naïve-Bayes Classifier,," 2024, [Online]. Available: <http://arxiv.org/abs/2411.01203>
- [21] İ. Ö. Bucak, M. Veziroğlu, and E. Veziroğlu, "Performance Comparison between Naive Bayes and Machine Learning Algorithms for News Classification,," İ. Ö. Bucak, Ed., Rijeka: IntechOpen, 2024. doi: 10.5772/intechopen.1002778.
- [22] P. Sarang, "Naive Bayes BT - Thinking Data Science: A Data Science Practitioner's Guide,," P. Sarang, Ed., Cham: Springer International Publishing, 2023, pp. 143–152. doi: 10.1007/978-3-031-02363-7_7.
- [23] D. Tang and X. Tong, "Evaluation of Missing Data Analytical Techniques in Longitudinal Research:

- Traditional and Machine Learning Approaches,” pp. 1–47, 2024, [Online]. Available: <http://arxiv.org/abs/2406.13814>
- [24] E. F. Finch and E. J. Mellen, “‘Labeled, Criticized, Looked Down On’: Characterizing the Stigma of Narcissistic Personality Disorder,” *Personal. Ment. Health*, vol. 19, no. 2, p. e70015, May 2025, doi: <https://doi.org/10.1002/pmh.70015>.
- [25] D. M. Hilty, Y. Cheng, and D. D. Luxton, “Artificial Intelligence and Predictive Modeling in Mental Health BT - Digital Mental Health: The Future is Now,” D. Mucić and D. M. Hilty, Eds., Cham: Springer Nature Switzerland, 2024, pp. 323–350. doi: [10.1007/978-3-031-59936-1_13](https://doi.org/10.1007/978-3-031-59936-1_13).
- [26] E. Odhiambo Omuya, G. Onyango Okeyo, and M. Waema Kimwele, “Feature Selection for Classification using Principal Component Analysis and Information Gain,” *Expert Syst. Appl.*, vol. 174, p. 114765, 2021, doi: <https://doi.org/10.1016/j.eswa.2021.114765>.
- [27] G. H. Martono and N. Sulistianingsih, “Optimizing Autisme Spectrum Disorder Identification with Dimentionality Reduction Technique and K-Medoid,” *5th Int. Conf. Comput. Sci. Inf. Manag.*, pp. 837–854, 2024.
- [28] D. S. W. Nguyen and M. M. Shaik, “Impact of Artificial Intelligence on Corporate Leadership,” *J. Comput. Commun.*, vol. 12, no. 04, pp. 40–48, 2024, doi: [10.4236/jcc.2024.124004](https://doi.org/10.4236/jcc.2024.124004).
- [29] P. Elisa and A. R. Isnain, “COMPARISON OF RANDOM FOREST, SUPPORT VECTOR MACHINE AND NAIVE BAYES ALGORITHMS TO ANALYZE SENTIMENT TOWARDS MENTAL HEALTH STIGMA,” *J. Tek. Inform.*, vol. 5, no. 1 SE-Articles, pp. 321–329, Feb. 2024, doi: [10.52436/1.jutif.2024.5.1.1817](https://doi.org/10.52436/1.jutif.2024.5.1.1817).
- [30] Shixiao Li, “Employment trend prediction and innovative talent training strategy optimization using naive Bayes classifier,” *J. Comput. Methods Sci. Eng.*, p. 14727978251337980, Apr. 2025, doi: [10.1177/14727978251337979](https://doi.org/10.1177/14727978251337979).
- [31] P. Apricia *et al.*, “Kinerja Naïve Bayes Classifier Pada Penyaringan Short Message Service (SMS) Spam,” *J. Siger Mat.*, vol. 04, no. 02, pp. 59–66, 2023.
- [32] M. Heydarian, T. E. Doyle, and R. Samavi, “MLCM: Multi-Label Confusion Matrix,” *IEEE Access*, vol. 10, pp. 19083–19095, 2022, doi: [10.1109/ACCESS.2022.3151048](https://doi.org/10.1109/ACCESS.2022.3151048).
- [33] R. Chen, “Narcissistic personality disorder: A general review,” *SHS Web Conf.*, vol. 193, p. 03012, 2024, doi: [10.1051/shsconf/202419303012](https://doi.org/10.1051/shsconf/202419303012).
- [34] D. Chicco and G. Jurman, “The advantages of the Matthews correlation coefficient (MCC) over F1 score and accuracy in binary classification evaluation,” *BMC Genomics*, vol. 21, no. 1, pp. 1–13, 2020, doi: [10.1186/s12864-019-6413-7](https://doi.org/10.1186/s12864-019-6413-7).
- [35] C. Naulak and P. R. Jindal, “A comparative study of Naive Bayes Classifiers with improved technique on Text Classification,” *TechRxiv*, pp. 0–8, 2022, doi: [10.36227/techrxiv.19918360.v1](https://doi.org/10.36227/techrxiv.19918360.v1).
- [36] H. Sulistiani, A. Syarif, K. Muludi, and Warsito, “Performance evaluation of feature selections on some ML approaches for diagnosing the narcissistic personality disorder,” *Bull. Electr. Eng. Informatics*, vol. 13, no. 2, pp. 1383–1391, 2024, doi: [10.11591/eei.v13i2.6717](https://doi.org/10.11591/eei.v13i2.6717).
- [37] S. L. K. Gruijters, “Gruijters 2019 - Using Principal Component Analysis to Validate Psychological Scales,” vol. 20, no. 5, pp. 544–549, 2003.
- [38] M. Greenacre, P. J. F. Groenen, T. Hastie, A. I. D’Enza, A. Markos, and E. Tuzhilina, “Principal component analysis,” *Nat. Rev. Methods Prim.*, vol. 2, no. 1, p. 100, 2022, doi: [10.1038/s43586-022-00184-w](https://doi.org/10.1038/s43586-022-00184-w).
- [39] F. L. Gewers, S. Paulo, and S. Paulo, “Supplementary Material for : Principal Component Analysis : A Natural Approach,” vol. 54, no. 4, 2021.
- [40] D. B. Dwyer, P. Falkai, and N. Koutsouleris, “Machine Learning Approaches for Clinical Psychology and Psychiatry,” *Annu. Rev. Clin. Psychol.*, vol. 14, pp. 91–118, 2018, doi: [10.1146/annurev-clinpsy-032816-045037](https://doi.org/10.1146/annurev-clinpsy-032816-045037).