

Application of the Naïve Bayes Algorithm for Employee Performance Prediction Based on SIMPEG at TVRI East Kalimantan Station

Ishmah Hanani, Siti Lailiyah*, Yulindawati*

Department of Informatics Engineering, STMIK Widya Cipta Dharma, Samarinda, Indonesia

Email: ¹ 2243909@wicida.ac.id, ^{2*} lail.59a@gmail.com, ^{3*} yulindawati@wicida.ac.id

(Corresponding author: 2243909@wicida.ac.id)

Abstract - Employee performance evaluation is a crucial aspect of public organizational management, including at the public broadcasting institution TVRI East Kalimantan Station. To date, attendance indicators obtained from the Employee Management Information System (SIMPEG) have often been used as the primary benchmark, as the data are objectively and structurally available. However, a single attendance-based approach risks overlooking more substantive aspects of work achievement. Therefore, this study integrates attendance data with the Employee Performance Targets (SKP) to construct a more representative performance label. The method employed is a classification approach using the Naïve Bayes (GaussianNB) algorithm. The research dataset consists of attendance records (normal attendance, leave, official duty, study assignment, early departure, absence, and total working days) and quantized SKP scores. Performance labels were generated using a composite score ($0.30 \times \text{attendance percentage} + 0.70 \times \text{normalized SKP}$), which was then categorized into three classes: Excellent, Good, and Needs Improvement. The model was trained using SIMPEG and SKP data that had undergone preprocessing, data partitioning, and class balancing. Experimental results show that the model achieved an accuracy of 0.83, with a precision of 0.86, recall of 0.84, and F1-score of 0.83 on the test data. These results indicate that the model can consistently recognize employee performance patterns across all categories. Practically, this study offers a simple, efficient, and easily implementable predictive framework to support more objective processes of coaching, monitoring, and reward allocation within TVRI East Kalimantan Station.

Keywords: Employee performance, SIMPEG, Employee Performance Targets, Naïve Bayes, Performance prediction.

1. INTRODUCTION

Human resource management (HRM) is a strategic factor determining organizational success in both public and private sectors. Competent, disciplined, and adaptive human resources support the achievement of organizational goals through effective processes and measurable outcomes. In public broadcasting institutions such as TVRI East Kalimantan Station, the quality of broadcast services and operational support is strongly influenced by the consistency and productivity of employees. In practice, many departments use attendance discipline indicators as the primary proxy for performance because such data are easily obtained from the Employee Management Information System (SIMPEG), well-structured, and relatively objective [2], [11]. The availability of detailed attendance data makes it a convenient and quick metric for evaluation. However, relying solely on this indicator presents several conceptual and practical limitations.

Performance assessments based solely on attendance risk obscuring the true dimensions of work achievement. Employees with perfect attendance do not necessarily meet their performance targets; conversely, those who occasionally take official leave or field duties may still exhibit high performance by achieving measurable output indicators. To address this gap, the government encourages the use of the Employee Performance Target (*Sasaran Kinerja Pegawai* – SKP) as a goal-based evaluation instrument. SKP captures aspects not reflected in attendance records, such as output quality, timeliness, and target realization, thus better representing the substantive dimension of performance rather than mere attendance discipline. Based on this perspective, this study asserts that employee performance cannot be evaluated solely through attendance but must also incorporate elements of work achievement represented by SKP.

In the analytical domain, the advancement of data mining and machine learning provides lightweight yet effective predictive approaches to support data-driven decision-making [1]. The Naïve Bayes (NB) algorithm is a probabilistic classification method widely recognized for its simplicity, computational efficiency, and competitive performance across various domains [1]. It is also frequently applied in Decision Support Systems (DSS) within the education sector because of its transparent computation process and straightforward implementation [16], [17]. Nonetheless, prior literature indicates that in certain complex datasets, Naïve Bayes may underperform compared to nonlinear models such as Neural Networks (NN), making the quality and relevance of features crucial [17]. In this study, this practical implication motivates the integration of attendance and SKP data to enrich the learning signals for the model rather than relying on a single source of indicators.

Naïve Bayes is a classification algorithm grounded in Bayes' Theorem, a mathematical approach for estimating the probability of an event based on prior knowledge or observed evidence [1]. Its core characteristic lies in the assumption of conditional independence among features, meaning that each variable contributes independently to the classification decision. Although this assumption is simple, it enables the algorithm to perform efficiently, be easily trained, and require minimal computational resources.

In practical terms, the algorithm estimates the probability of each class based on the combination of observed feature values and assigns the class with the highest probability as the predicted outcome. Its probabilistic nature allows

predictions to be interpreted quantitatively and transparently, making it particularly suitable for public institutions that demand accountability and explainability in decision-making processes [1], [4]. Another advantage is that Naïve Bayes remains competitive when applied to small- to medium-scale tabular datasets, such as personnel administration data, without requiring complex training procedures. However, its effectiveness highly depends on the relevance of input features—the more informative the attributes, the better the model's ability to identify data patterns [4]. For these reasons, Naïve Bayes is deemed appropriate for this study, which combines attendance (SIMPEG) and SKP data to predict employee performance objectively and efficiently.

Beyond the HR or education domains, the flexibility of Naïve Bayes has been demonstrated across diverse applications [1], [12]. Prior studies on employee performance evaluation and prediction have utilized Naïve Bayes and its variants for performance appraisal, employee eligibility classification, and other personnel decisions [4]–[7], [9]–[12]. In the academic sector, Naïve Bayes has also been widely employed to model student performance and graduation prediction [8], [13]–[17]. However, most HR-oriented studies still focus on attendance-related attributes (normal attendance, leave, official duty, or absence) or other administrative features, while explicit integration with SKP—either as a key feature or as a basis for constructing performance labels—remains relatively uncommon. Within governmental institutions, several studies highlight the importance of leveraging SIMPEG data to promote transparency and support bureaucratic reform [2], [11].

Other studies further illustrate the algorithm's versatility. Yusnita et al. [16] applied Naïve Bayes to a student admission system to support selection decisions, while Azahari et al. [17] explored its use in predicting undergraduate study duration. Both studies emphasize that Naïve Bayes can be effectively implemented in educational and academic management contexts. Building upon these findings, the present study positions the application of Naïve Bayes in the public sector, specifically through the integration of SIMPEG and SKP data to assess employee performance more comprehensively.

This research extends such utilization from monitoring to predictive classification of performance categories that can be directly applied for employee development and reward allocation. The case of TVRI provides a unique context since, as a public broadcasting institution, it must fulfill both governmental obligations and public service demands. Consequently, employee performance evaluation not only concerns internal administrative accountability but also the quality of public information services delivered [3]. The nature of broadcasting tasks—requiring on-time program delivery, production quality, and cross-functional coordination—makes performance indicators more diverse than those in purely administrative sectors. This condition reveals a research gap: few studies have integrated SIMPEG and SKP data within public broadcasting institutions, despite the distinctive operational demands. This gap thus forms the foundation of the present study.

Based on the above background, this study takes the case of TVRI East Kalimantan Station and establishes the following objectives: (1) to formulate a composite performance label integrating attendance and SKP components; and

(2) to build a Naïve Bayes model to predict the resulting performance label.

From a policy perspective, this study adopts a three-class framework—Excellent, Good, and Needs Improvement. Academically, it expands the application of Naïve Bayes in the public sector, which has been more frequently reported in educational and industrial contexts [9], [11], [16], [17], by positioning SKP as a key determinant of performance. Practically, this approach can be readily adopted by HR units because it is transparent (the formulas and thresholds are easily explainable to stakeholders), computationally efficient (suitable for limited infrastructure), and flexible in aligning with institutional policies. Therefore, this study offers a more representative and operationally feasible framework for predicting employee performance at TVRI East Kalimantan Station.

2. RESEARCH METHODOLOGY

2.1 Research Stages

This study is an applied quantitative research employing a supervised classification approach to predict employee performance categories using the Naïve Bayes (NB) algorithm at TVRI East Kalimantan Station. The choice of NB is motivated by: (i) its simplicity, speed, and ease of institutional deployment; (ii) competitive performance on tabular data; and (iii) transparent probability computations that facilitate stakeholder communication [1], [4], [16]. The overall framework follows the Knowledge Discovery in Databases (KDD) process comprising six main stages: (1) data collection from SIMPEG and SKP; (2) attribute selection; (3) preprocessing (cleaning, encoding, normalization); (4) transformation and train–test partitioning; (5) application of the Naïve Bayes classifier; and (6) model evaluation using confusion matrix, accuracy, precision, recall, and F1-score [1].

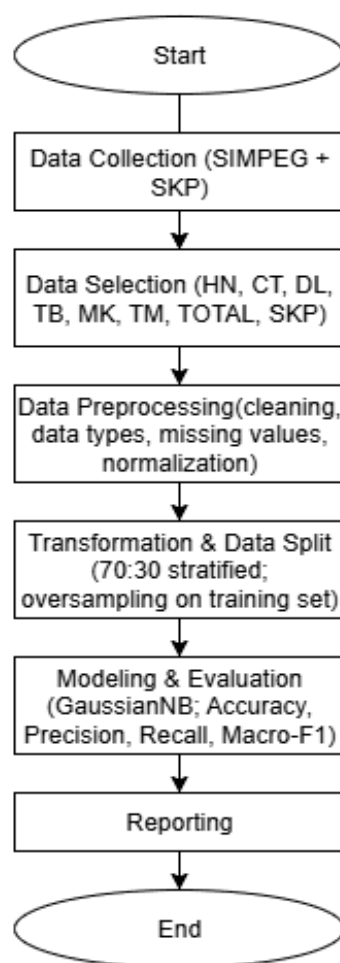


Figure 1. KDD-Based Research Flow

At the same time, emphasizing feature quality (combining attendance + SKP) because the literature shows that for certain data, NB can lag behind nonlinear models such as Neural Networks, so attribute selection is crucial [17]. This approach is in line with common practices that utilize NB in the domains of education and selection, as demonstrated by Yusnita et al. [16], who utilized NB in a new student admission system, and Azahari et al. [17], who used NB to predict student study periods. Both studies demonstrate the simplicity and adaptability of NB in processing structured administrative data.

Data sources comprise the SIMPEG (employee management information) system at TVRI East Kalimantan Station and SKP (Employee Performance Target) evaluation sheets for the same period. The workflow includes data extraction, cleaning and normalization, construction of a composite performance label (Attendance + SKP) into three classes (Excellent, Good, Needs Improvement), feature engineering, training with GaussianNB, and evaluation on a 30% hold-out test set using accuracy, precision, recall, F1-score, and confusion matrix. This scheme reflects concise, organization-friendly data-mining practice [1], [4], and is consistent with NB implementations in HR/education domains [5]–[8], [10]–[12], [15]. Policy and data-governance considerations follow public-sector practices of leveraging SIMPEG to support transparency and bureaucratic reform [2], [11].

Tools & environment: Python with *scikit-learn* for modeling, *pandas* for data handling, and standard visualization (*matplotlib*). GaussianNB uses default parameters, as the focus is on composite-label design and feature construction rather than extensive hyperparameter tuning [4].

2.2 Population and Sample

The population of this study consists of employees registered in the Personnel Management Information System (SIMPEG) at TVRI East Kalimantan Station who had Employee Performance Targets (SKP) assessments during the observation period. This population is relevant because SIMPEG provides standardized administrative/attendance records, while SKP represents target-based performance achievements [2], [11].

The research sample consists of all data rows (employee-period) that passed the pre-processing stage, namely the numerical attendance columns HADIRNORMAL_HN, CUTI_CT, DINASLUAR_DL, TUGASBELAJAR_TB, MENINGGALKANKANTOR_MK, TIDAKMASUK_TM, TOTAL. The SKP EVALUATION column is available and has been successfully mapped. And the TOTAL value is > 0 (to avoid division by zero when calculating the attendance percentage). After cleaning, the data is divided into 70% for training and 30% for testing with a stratified split so that the proportions of the three classes (Very Good/Good/Needs Improvement) are relatively balanced in both parts. Given the tendency for the Good class to dominate, RandomOverSampler is applied to the training data to reduce model bias towards the majority class. The test data is left untouched (no resampling) so that the evaluation reflects performance on unseen data.

This total sample approach (census of available rows) is common in SIMPEG-based prediction studies, maximizing available data while maintaining internal validity through strict pre-processing procedures [2], [11]. It is also in line with practices in many Naïve Bayes studies on employee performance [5]–[8], [10]–[12], [15].

Thus, the evaluation results obtained from the test data can provide an objective picture of the performance of the Naïve Bayes algorithm in predicting employee performance based on attendance and other administrative attributes from SIMPEG.

2.3 Research Variables

Dependent variable (label): *Predikat Kinerja* $\in$ {Excellent, Good, Needs Improvement}, derived from a composite score combining attendance discipline and work achievement (SKP). The three-class mapping supports operational coaching/reward decisions and reduces ambiguity in mid-range labels.

Independent variables (features): Combined SIMPEG attendance attributes and SKP indicator: HADIRNORMAL_HN (normal attendance days), CUTI_CT (leave days), DINASLUAR_DL (official duty days), TUGASBELAJAR_TB (training/study days), MENINGGALKANKANTOR_MK (early-leave events), TIDAKMASUK_TM (absence days), Operational definitions & scaling: *persen_hadir* = HADIRNORMAL_HN / TOTAL, clipped to [0,1] to prevent data artifacts. SKP mapping $\rightarrow$ *skp_percent*: *excellent* = 150, *good* = 100, *needs improvement* = 75, *Below Expectation* = 50, *Unsatisfactory* = 25; then *skp_norm* = *skp_percent* / 150 for scale compatibility. Composite score (for label construction): *score* = $0.30 \times \text{persen_hadir} + 0.70 \times \text{skp_norm}$. Three-class policy mapping: Excellent ($\text{score} \geq 0.85$), Good ($0.70 \leq \text{score} < 0.85$), Needs Improvement ($\text{score} < 0.70$, including Below Expectation/Unsatisfactory).

Methodological note. *skp_percent* is used as an input feature (besides attendance components), while *Predikat Kinerja* is the 3-class label predicted by the model. This design ensures the model learns from combined signals (attendance + SKP)—a proven approach in HR/education NB applications [5]–[8], [10]–[12], [15].

Table 1. Variabel Independen (Fitur Model)

No	Variable	Code	Operational Definition	Computation/Values	Scale	Unit	Source
1	Present	HN	Days present as scheduled	Numeric from SIMPEG	Ratio	Days	SIMPEG
2	Leave	CT	Leave days in the period	Numeric from SIMPEG	Ratio	Days	SIMPEG
3	Official Duty	DL	Official duties outside office	Numeric from SIMPEG	Ratio	Days	SIMPEG
4	Study/Training	TB	Training/study days	Numeric from SIMPEG	Ratio	Days	SIMPEG
5	Leaving the office	MK	Frequency/events of leaving during work hours	Numeric from SIMPEG	Ratio	Event/ Days	SIMPEG
6	Absent	TM	Days absent	Numeric from SIMPEG	Ratio	Days	SIMPEG
7	Total Working Days	TOTAL	Total relevant working entries	Numeric from SIMPEG	Ratio	Days	SIMPEG
8	SKP (quantized)	Skp_percent	Quantized SKP score (category mapping)	150=Excellent; 100=Good; 75=Needs Improvement; 50=Below Expectation; 25=Unsatisfactory	Ratio	Percent	SKP

The first seven variables represent the documented attendance patterns of employees in SIMPEG (present, leave, official duty, study/training, leaving the office, absent, and total working days). These variables were selected because they are available, standardized, and auditable in the personnel administration process [2], [11]. The `skp_percent` variable is the main feature that represents work achievement derived from the Employee Performance Target (SKP) assessment and quantified to the set {150, 100, 75, 50, 25} for consistency between entries. The inclusion of `skp_percent` confirms that performance is not only about attendance discipline, but also about achieving outputs methodologically, incorporating meaningful features that improve the signal for the Naïve Bayes algorithm [1], [4], [16], while also responding to the supervisor's directive that assessments should not be “absence-only”.

Table 2. Intermediate Variables (for Label Formation; not model features except `skp_percent`)

No	Variable	Code	Operational Definition	Computation/Values	Scale	Unit	Source
1	Attendance Ratio	<code>persen_hadir</code>	Proportion of presence	HN / TOTAL, clipped to [0,1]	Ratio	-	SIMPEG
2	Normalized SKP	<code>skp_norm</code>	SKP score in [0,1]	<code>skp_percent</code> / 150	Ratio	-	SKP
3	Composite Score	<code>score</code>	Composite for label construction	$0.30 \times \text{persen_hadir} + 0.70 \times \text{skp_norm}$	Ratio	-	Derived

These three intermediate variables are not used directly as features (except for `skp_percent` in Table 1), but are used to form labels that will be predicted by the model. `persen_hadir` normalizes attendance to [0,1], while `skp_norm` normalizes SKP to a comparable scale. `Score` combines the two signals with weights of 0.30 : 0.70 (leaning towards SKP) because the substance of work performance must stand out without negating the importance of attendance discipline. This composite strategy is commonly adopted in NB/HR studies to maintain the balance of heterogeneous quantitative aspects [5]–[8], [10]–[12], [15].

Table 3. Dependent Variable (Performance Label)

No	Variable	Code	Operational Definition	Class Rule	Scale	Domain
1	Performance Category	<code>Predikat_Kinerja</code>	Category mapped from the composite Attendance+SKP score	Excellent if $\text{score} \geq 0.85$; Good if $0.70 \leq \text{score} < 0.85$; Needs Improvement if $\text{score} < 0.70$	Nominal (3 classes)	{SB, B, BP}

The Performance Predicate label is mapped into three classes to facilitate coaching and reward policies. Good (representing the majority of operations) is combined, while Below Expectation/Unsatisfactory is combined into Needs Improvement so that the intervention plan is clear and measurable. This three-class approach is in line with managerial needs in agencies and maintains the stability of estimates in the Naïve Bayes model (reducing label sparsity), as recommended in applied classification practices [1], [4], [16].

Since all features used are numerical and have a continuous distribution, the Gaussian Naïve Bayes (GNB) variant is used as the classification algorithm. This model calculates the probability of each feature for each class based on a normal (Gaussian) distribution, as shown in Equation (1):

$$P(x_j | C_k) = \frac{1}{\sqrt{2\pi\sigma_{jk}^2}} \exp\left(-\frac{(x_j - \mu_{jk})^2}{2\sigma_{jk}^2}\right) \quad (1)$$

dengan:

- x_j = the j -th feature value of an employee (e.g., number of normal attendance days, HN = 20);
- C_k = the performance class being evaluated (e.g., *Good*);
- μ_{jk} = mean of feature j for all employees belonging to class C_k ;
- σ_{jk}^2 = variance of feature j in class C_k ;
- $\exp(.)$ = exponential function;
- $\sqrt{2\pi\sigma_{jk}^2}$ = normalization factor ensuring that total probability sums to one.

This equation measures how distinctive an employee's feature values are for each performance class under the assumption that the data follow a normal distribution centered at the class mean (μ), with the spread determined by the standard deviation (σ). Values near the mean obtain higher likelihood scores, while those further away decrease exponentially. The model then multiplies the likelihoods of all features and selects the class with the highest overall score. Based on the formula above, for each feature (HN, CT, DL, TB, MK, TM, TOTAL, SKP), the likelihood $P(x_j | C_k)$ is calculated. For each performance class C_k , all feature probabilities are multiplied—assuming conditional independence—to obtain the joint probability $P(C_k | X)$, which represents the likelihood that an employee belongs to class C_k . The class with the highest posterior probability is then selected as the prediction result.

Thus, the model does not guess randomly but systematically compares how well each employee's feature profile fits the "characteristic pattern" of each performance category. All probability computations are automatically handled by the scikit-learn library (*GaussianNB*). Systematically, the GaussianNB formulation provides an appropriate representation for continuously distributed data such as those in the SIMPEG and SKP datasets, ensuring stable and interpretable predictive outcomes

2.4 Naïve Bayes Classification Pipeline

The classification process was designed to be simple, transparent, and easily replicable within institutional environments: (i) input of SIMPEG (attendance) and SKP data; (ii) preprocessing and normalization; (iii) construction of three-class performance labels from the composite Attendance+SKP score; (iv) feature assembly; (v) 70:30 stratified split with oversampling applied to the training set; (vi) training of the Gaussian Naïve Bayes classifier; (vii) testing on the hold-out test set; and (viii) evaluation using accuracy, precision, recall, F1-score, and the confusion matrix.

The choice of Naïve Bayes is based on its computational efficiency, simplicity, and established performance in Decision Support Systems and educational data-mining applications [1], [4], [16]. Particular attention is given to feature quality, as Naïve Bayes may underperform compared to nonlinear models when feature relevance is weak [17]. The integration of SIMPEG and SKP data follows best practices in governmental HR information systems aimed at promoting transparency and bureaucratic accountability [2], [11]. Model performance is evaluated using standard classification metrics—accuracy, precision, recall, and F1-score—defined as follows:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (2)$$

$$\text{Precision} = \frac{TP}{TP + FP} \quad (3)$$

$$\text{Recall} = \frac{TP}{TP + FN} \quad (4)$$

$$\text{F1-score} = \frac{2 \times (\text{Precision} \times \text{Recall})}{\text{Precision} + \text{Recall}} \quad (5)$$

With:

- TP (True Positive): Number of positive samples correctly predicted.
- TN (True Negative): Number of negative samples correctly predicted.
- FP (False Positive): Number of negative samples incorrectly predicted as positive.
- FN (False Negative): Number of positive samples not detected by the model.

For the multi-class case (three classes), the values of TP , TN , FP , and FN are computed using a one-vs-rest scheme for each class. Subsequently, both macro average (the unweighted mean across classes) and weighted average (the mean weighted by class sample size) are reported to provide a balanced evaluation. This research followed the Naïve Bayes classification workflow as illustrated in Figure 2.

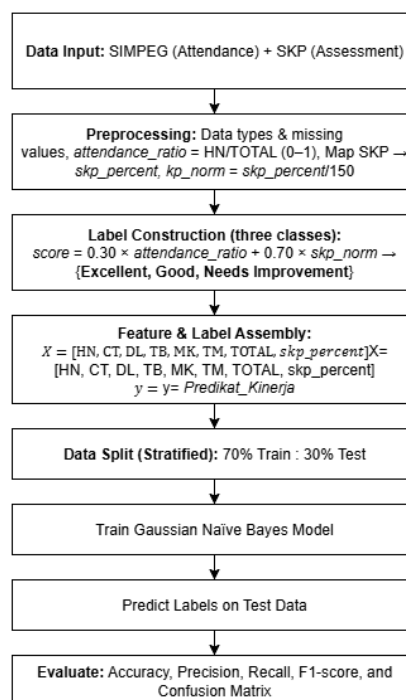


Figure 2. Naïve Bayes Classification Flows

The employee performance classification workflow using the Naïve Bayes algorithm consists of seven main stages:

- a. **Data Input and Preprocessing:**
Retrieve data from SIMPEG and SKP, including attendance attributes (HN, CT, DL, TB, MK, TM, TOTAL) from SIMPEG and SKP evaluations for the corresponding period [2], [11].
- b. **Data Cleaning and Normalization:**
Convert data types, handle missing or anomalous values; compute $attendance_ratio = HADIRNORMAL_HN / TOTAL$ and clip the result to the [0,1] range. Map the “SKP Evaluation” into $skp_percent$ with the following quantization: *Excellent* = 150, *Good* = 100, *Needs Improvement* = 75, *Below Expectation* = 50, and *Unsatisfactory* = 25. Then calculate $skp_norm = skp_percent / 150$ to ensure scale comparability.
- c. **Label Construction (Three Classes):**
Compute the composite score $score = 0.30 \times attendance_ratio + 0.70 \times skp_norm$, then map it into three categories: *Excellent* ($score \geq 0.85$), *Good* ($0.70 \leq score < 0.85$), and *Needs Improvement* ($score < 0.70$, including Below Expectation and Unsatisfactory).
- d. **Feature and Label Assembly:**
 $X = [HN, CT, DL, TB, MK, TM, TOTAL, skp_percent]$,
 $y = Predikat_Kinerja \in \{Excellent, Good, Needs Improvement\}$
Including $skp_percent$ as a feature enriches the representation of achievement signals alongside attendance patterns.
- e. **Data Splitting and Balancing:**
Apply a 70:30 stratified split between training and test sets. To reduce bias toward the majority class (*Good*), apply RandomOverSampler only to the training data. The test data remain untouched (unseen).
- f. **Model Training:**
The study employs Gaussian Naïve Bayes (GaussianNB) for three-class performance classification. Conceptually, the model learns class-specific feature distributions (mean and variance) from the training data. During prediction, it evaluates how well an employee’s feature values fit each class. The overall match from all features is combined with class priors, and the class with the highest posterior probability is chosen as the prediction result.
This approach is efficient, transparent, and well-suited for tabular data such as SIMPEG and SKP, as it requires minimal computation and offers easily interpretable logic for decision-makers. Feature quality is emphasized since simple models like Naïve Bayes rely heavily on relevant information; in some datasets, more complex models may outperform, making attribute selection a key factor.

g. Prediction and Evaluation:

Generate predictions on the test set and compute *accuracy*, *precision*, *recall*, and *F1-score* per class, along with *macro* and *weighted averages* and the confusion matrix. The desired model performance target is $0.79 \leq \text{accuracy} < 1.00$.

3. RESULT AND DISCUSSION

This section presents the research findings and analytical discussion related to the implementation of the Naïve Bayes algorithm for predicting employee performance based on data obtained from SIMPEG at TVRI East Kalimantan Station. The discussion is divided into three subsections: (1) data processing and model training, (2) model evaluation results, and (3) system implementation and practical implications.

3.1 Data Processing and Model Training

a. Research Dataset

The dataset used in this study integrates administrative records from the *Employee Management Information System* (SIMPEG) and the *Employee Performance Targets* (SKP) data for the same evaluation period of employees at TVRI East Kalimantan Station. From SIMPEG, attendance-related attributes consistently recorded by the HR unit were extracted: HADIRNORMAL_HN (normal attendance days), CUTI_CT (leave days), DINASLUAR_DL (official duty days), TUGASBELAJAR_TB (training/study days), MENINGGALKANTOR_MK (early-leave events), TIDAKMASUK_TM (absence days), and TOTAL working days or valid entries. The SKP data were obtained from performance assessment forms based on target achievement and work quality, then recoded into a numerical variable *skp_percent* using the following mapping: *Excellent* = 150, *Good* = 100, *Needs Improvement* = 75, *Below Expectation* = 50, and *Unsatisfactory* = 25. The inclusion of these two complementary data sources underscores the main purpose of this study: employee performance should not be evaluated solely on attendance. Attendance reflects discipline and work availability, whereas SKP represents goal attainment and quality of output. Their combination yields a more comprehensive and operational representation of employee performance, suitable for both coaching and recognition purposes. This approach aligns with recommendations for leveraging personnel data to enhance transparency and support bureaucratic reform within public-sector institutions [2], [11].

Conceptually, the data handled in this study are tabular, consisting of discrete or quasi-continuous numerical features. Such characteristics are well-suited for lightweight and explainable applied machine-learning approaches such as Naïve Bayes [1], [4], [9], [18]. Models of this type are easily operationalized in institutional environments because they require low computational resources and provide transparent reasoning that can be clearly communicated to stakeholders..

b. Data Processing

The preprocessing stage was conducted to ensure that the data were clean, consistent, and scaled comparably before being trained in the model. The following steps were applied:

1. Data Type Alignment : All attendance-related columns from SIMPEG were converted into numeric types, and non-numeric entries were cleaned. Textual SKP scores were mapped to the numeric variable *skp_percent*.
2. Handling Missing or Anomalous Values: Missing values in attendance features were treated using a *minimal-distortion* principle—assigning rational zeros for non-occurring events or removing rows that failed to meet fundamental conditions, such as invalid *TOTAL* values.
3. Attendance Normalization: The variable *attendance_ratio* was defined as HN/TOTAL and clipped to the range $[0,1]$ to prevent artifacts (e.g., division by zero). This normalization made attendance signals comparable across employees and work periods.
4. SKP Normalization: After categorical mapping to *skp_percent*, a normalized variable $\text{skp_norm} = \text{skp_percent} / 150$ was created to ensure that SKP values fall within the range $[0,1]$, comparable to the attendance ratio.
5. Feature and Label Assembly: The feature matrix (*X*) consists of [HN, CT, DL, TB, MK, TM, TOTAL, *skp_percent*], and the label (*y*) represents the three-class performance category (*Excellent*, *Good*, *Needs Improvement*), constructed from the composite score described in the Methodology section.
6. Quality Gate Filtering: Rows with $\text{TOTAL} \leq 0$ were removed to avoid division errors, and entries lacking SKP records for the given period were excluded because a valid label could not be constructed.

The decision to include *skp_percent* as both a predictive feature and a component of the composite label is substantively grounded. Organizationally, performance evaluation emphasizes target achievement as a determinant of overall performance; thus, SKP signals are logically integrated with attendance patterns. The literature also

highlights that feature quality is critical to the performance of Naïve Bayes on tabular data—simple models can underperform compared to nonlinear approaches (e.g., neural networks) when features lack discriminative strength [1], [4], [17], [18].

c. Data Splitting

The dataset was divided into a training set (70%) and a test set (30%) using stratified sampling, ensuring that the class proportions in the labels were preserved in both subsets. The test set contained 36 instances, distributed as follows according to the classification report: *Good* = 20, *Needs Improvement* = 12, and *Excellent* = 4. Stratification was crucial to prevent evaluation bias toward the majority class.

Since the natural class distribution placed *Good* as the dominant category, a *RandomOverSampler* was applied only to the training data. This technique balanced the number of examples in each class during model learning, reducing the model’s tendency to favor the majority class. The test data were left untouched (no resampling) to ensure that the evaluation metrics truly reflected the model’s generalization capability on unseen data.

This design follows standard evaluation practices for classification models in HR and educational domains [5]–[8], [10]–[12], [15], [16].

d. Model Training

The model employed in this study is the Gaussian Naïve Bayes (*GaussianNB*) classifier from the *scikit-learn* library. Naïve Bayes learns class-specific summary statistics of features—namely the mean and variance—from the training data. During prediction, the model calculates how well an employee’s feature values fit each class profile. The compatibility scores across all features are then combined with class priors, and the class with the highest combined probability is selected as the predicted outcome. The rationale for choosing Naïve Bayes includes: Efficiency and speed: It can be trained quickly and is well-suited for institutional environments with limited computational resources. Transparency: The decision-making process is easily interpretable by management, as it is based on straightforward statistical summaries and additive feature contributions [1], [4], [9]. Competitiveness on tabular data: When features are relevant, Naïve Bayes performs competitively with more complex models. Incorporating *skip_percent* alongside attendance patterns enriches the predictive signal.

The *GaussianNB* model used default parameters, as the research focus was on composite label design, feature construction, and performance validation rather than extensive hyperparameter tuning. Nevertheless, the model could be further enhanced through cross-validation or lightweight parameter optimization in future operational stages if required by the organization.

3.2 Model Evaluation Result

a. Model Evaluation

The model evaluation was conducted using precision, recall, F1-score, and support metrics, with an overall accuracy of 0.83. The test results are presented in Table 4.

Table 4. Model Evaluation Matriks

Class	Precision	Recall	F1-Score	Support
Good	0.79	0.95	0.86	20
Needs Improvement	1.00	0.58	0.74	12
Excellent	0.80	1.00	0.89	4
Accuracy			0.83	36
Macro avg	0.86	0.84	0.83	36
Weighted avg	0.86	0.83	0.82	36

The evaluation utilized the macro average as the primary reference metric. On the test data, the model achieved an overall accuracy of 0.83, with a macro-averaged precision of 0.86, recall of 0.84, and F1-score of 0.83. These results indicate that the model performs relatively consistently across all three target classes—*Good*, *Needs Improvement*, and *Excellent*. The use of macro averaging was considered appropriate because it computes the mean performance across classes with equal weighting, regardless of sample distribution. This is particularly relevant for employee performance datasets where class imbalance naturally occurs: the *Good* category tends to dominate, while *Excellent* and *Needs Improvement* classes contain fewer samples. By employing macro averages, the evaluation reflects not only the model’s performance on the majority class but also its ability to correctly identify minority classes.

The macro-averaged precision value of 0.86 indicates that the model’s predictions are highly reliable, with a low error rate. In other words, when the model classifies an employee into a particular performance category, there is a high

probability that the prediction corresponds correctly to the true class. The macro-averaged recall of 0.84 confirms that most employees across all performance categories were successfully identified, reducing the likelihood of misclassification or missed detection of employees who should belong to a specific class. The balance between precision and recall is reflected in the macro-averaged F1-score of 0.83, which demonstrates the model's stable performance in both predictive accuracy and sensitivity. From an organizational standpoint, such consistent macro-average metrics are essential because they ensure fairness in performance evaluation across all employee categories. Employees with *Excellent* ratings remain accurately identified, those categorized as *Good* are recognized with high accuracy, and employees requiring improvement are reliably detected despite their smaller representation in the dataset. With these results, the model meets the established quality threshold (accuracy ≥ 0.79 and < 1.00), signifying that it is sufficiently robust to serve as a baseline decision-support model for employee performance evaluation.

b. Bar Chart

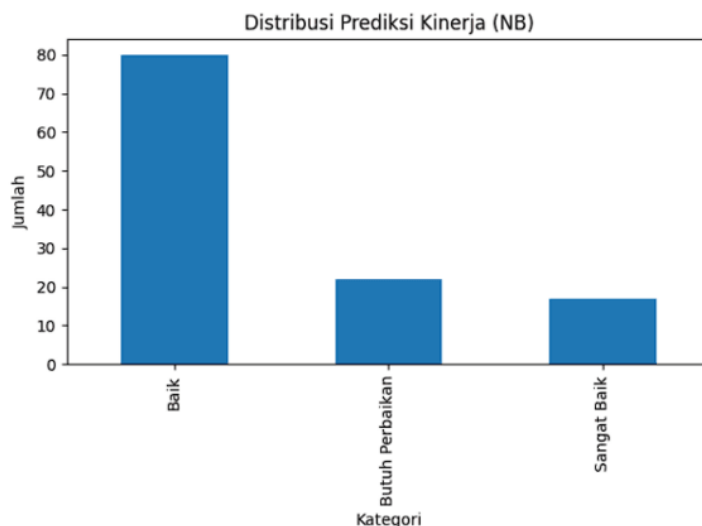


Figure 3. Distribution of Naïve Bayes Prediction

Figure 3 shows the overall distribution of model predictions across the full dataset: Good (80), Needs Improvement (22), and Excellent (17). This information is useful for managerial planning—e.g., sizing capacity for developmental programs targeting the *Needs Improvement* group and designing recognition strategies for the *Excellent* group each period. These counts are not used as scientific evaluation metrics; formal evaluation relies on the 30% hold-out test set, but the distribution remains operationally relevant for HR planning.

c. Confusion Matrix

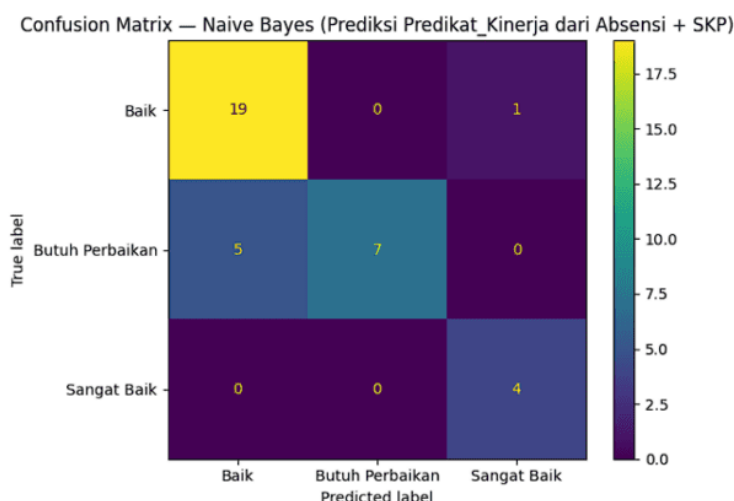


Figure 4. Naïve Bayes Classification Confusion Matrix

The confusion matrix on the test set exhibits the following patterns. Row “Good”: 19 correctly predicted as *Good*, 1 upgraded to *Excellent*, 0 predicted as *Needs Improvement*—explaining the high recall for *Good* (0.95). Row “Needs Improvement”: 7 correctly predicted as *Needs Improvement*, 5 misclassified as *Good*, 0 as *Excellent*. There is no over-promotion to *Excellent* (i.e., *Needs Improvement* → *Excellent*), but leakage into *Good* lowers recall to 0.58. Row

“Excellent”: 4 correctly predicted as *Excellent*, 0 to other classes—consistent with a recall of **1.00**. The precision for *Excellent* (0.80) indicates that a portion of predictions labeled *Excellent* may originate from other classes (typically *Good*) with scores near the 0.85 decision boundary.

Operationally, these patterns suggest that the model is conservative toward *Excellent* (high recall, moderate precision) and highly reliable for *Good*, while borderline cases between *Needs Improvement* and *Good* warrant attention (e.g., review of feature thresholds or targeted coaching criteria).

3.3 Analysis of Result

The analysis was conducted according to the Naïve Bayes stages described in the methodology.

a. Preprocessing Stage

Normalization of attendance ($attendance_ratio = HN/TOTAL$) and SKP ($skp_norm = skp_percent / 150$) ensured that all features were comparable on a similar scale. Without normalization, the wider numerical range of SKP values could overshadow the attendance signal. The composite score weighting of 0.30 : 0.70 was chosen for discriminative reasons—SKP provides a stronger signal for differentiating performance levels, while attendance serves as a supporting indicator. This weighting aligns with the organizational context, where attendance is necessary but actual work achievement remains the primary determinant of employee performance.

b. Label Construction

The three performance classes—*Excellent* (SB), *Good* (B), and *Needs Improvement* (BP)—were derived from the composite score. The class distribution revealed that *Good* dominated, while *Excellent* and *Needs Improvement* appeared as minority categories. This imbalance implies that, without adjustment, the model could be biased toward the majority class. Therefore, oversampling was applied only to the training data, enabling the model to learn minority-class patterns without distorting evaluation integrity. The test data were kept in their original distribution so that evaluation metrics would fairly represent generalization performance.

The GaussianNB classifier performed well on numerical data; however, the broader dispersion of the *Needs Improvement* class made boundary separation more challenging.

c. GaussianNB Training Stage (Mechanism and Impact)

The Gaussian Naïve Bayes model learns the mean and variance of each feature per class, then computes the likelihood of an observation given each class distribution, weighted by the class prior. In this dataset, the *Needs Improvement* class exhibited a wider spread and overlapped with the *Good* class, particularly near the composite score threshold of 0.70–0.85. Statistically, this overlap reduced the posterior probability of *Needs Improvement* relative to *Good*, causing some *Needs Improvement* cases to be misclassified upward as *Good*.

This explains the mechanistic reason behind the recall disparity: the *Needs Improvement* class obtained a recall of 0.58, while the *Good* and *Excellent* classes achieved much higher recall scores (0.95 and 1.00, respectively). The *Excellent* class achieved perfect recall because its high SKP scores distinctly separated it from other categories.

d. Evaluation and Confusion Matrix Stage

The summarized test results yielded an accuracy of 0.83, macro precision of 0.86, macro recall of 0.84, and macro F1-score of 0.83, indicating balanced model performance across classes despite the inherent data imbalance. The confusion matrix further demonstrated that the *Good* class achieved a recall of 0.95 (stable and mostly correct predictions), the *Excellent* class achieved a recall of 1.00 (no instances were missed), and the *Needs Improvement* class obtained a recall of 0.58, primarily due to five out of twelve cases being misclassified upward into the *Good* category.

To improve performance, particularly in the *Needs Improvement* class, several enhancement strategies can be implemented. These include: performing threshold tuning within the 0.70–0.85 range to make class boundaries in the overlapping region more decisive; applying model-based class balancing using *class weight* (increasing the misclassification penalty for the *Needs Improvement* class) as a complement to data-level oversampling in the training set; enriching the SKP feature with additional indicators such as timeliness, target–achievement gap, or output quality to sharpen the signal of the *Needs Improvement* class; and optionally conducting probability calibration prior to threshold application, enabling more consistent score-based decisions.

These findings are consistent with previous literature highlighting the flexibility of Naïve Bayes for structured administrative data. Yusnita *et al.* [16] demonstrated the effectiveness of NB in supporting student admission decisions through transparent classification mechanisms, while Azahari *et al.* [17] emphasized its application for predicting undergraduate study duration. Both studies confirm that Naïve Bayes can yield reliable results when the input data are representative. In line with these insights, integrating attendance and SKP as key features has proven to enhance the model’s relevance for employee performance assessment in the public sector.

Furthermore, the results have clear managerial implications. First, the list of employees classified under *Needs Improvement* can be prioritized for coaching and training programs. Second, employees categorized as *Excellent* can be identified as candidates for recognition and reward programs. Third, the consistent macro-average metrics provide assurance that the model operates without systemic bias, supporting objective and transparent performance monitoring processes.

For future improvement, steps such as k-fold cross-validation can be applied to stabilize performance estimates on relatively small datasets, and enrichment of SKP-related features (e.g., timeliness indicators, target–realization gap, or work quality metrics) can further strengthen the discriminative power of the model. Nonetheless, even without algorithmic modification, the current results demonstrate that Naïve Bayes can provide an accurate, fair, and management-ready representation of employee performance to support data-driven decision-making.

4. CONCLUSION

This study demonstrates that employee performance prediction is more accurate when combining *attendance information (SIMPEG)* with *SKP achievement*, rather than relying on a single indicator. Using a composite score ($0.30 \times \text{attendance percentage} + 0.70 \times \text{normalized SKP}$) and a three-class scheme (*Excellent, Good, Needs Improvement*), the proposed *Gaussian Naïve Bayes* model achieved an accuracy of 0.83 on the test set. Macro-averaged metrics yielded *precision* = 0.86, *recall* = 0.84, and *F1-score* = 0.83, indicating balanced performance across all categories. These findings have practical implications for *TVRI East Kalimantan Station*. High precision indicates trustworthy predictions; stable recall ensures that most employees in each category are correctly identified; and consistent F1-scores reflect a sound balance between precision and recall. Consequently, the model can serve as a *lightweight, transparent baseline decision-support tool*. For future work, we recommend *enriching SKP features* (e.g., quality indicators, timeliness, target–realization gap) and applying *k-fold cross-validation* to stabilize performance estimates. These steps can improve the recall of minority classes without sacrificing precision. Overall, this research contributes a practical approach to strengthening *data-driven governance* through the integration of SIMPEG and SKP, and shows that *macro-average evaluation* offers a fair and balanced basis for managerial decision-making.

REFERENCES

- [1] A. Géron, *Machine Learning with Scikit-Learn and TensorFlow*. Jakarta, Indonesia: Elex Media, 2020.
- [2] D. W. Lubis and J. Veri, “Pengaruh Sistem Informasi Manajemen Kepegawaian terhadap Kualitas Pelayanan Administrasi Kepegawaian: Systematic Review,” *J. Manajemen Informatika Jayakarta*, vol. 5, no. 2, pp. 135–141, Apr. 2025.
- [3] N. Insani, *Signifikansi Budaya Organisasi, Semangat Kerja dan Kepuasan Kerja Terhadap Kinerja Pegawai pada LPP TVRI Sulawesi Selatan*, M.S. thesis, Univ. Muhammadiyah Makassar, 2024.
- [4] B. H. Pangestu, “Data Mining Menggunakan Algoritma Naïve Bayes Classifier untuk Evaluasi Kinerja Karyawan,” *J. Riset Matematika*, vol. 3, no. 2, pp. 177–184, Dec. 2023.
- [5] M. R. Khoirudin, M. Hasbi, B. Widada, K. Akhyar, and K. Sandradewi, “Klasifikasi Kelayakan Pegawai Kontrak Menjadi Pegawai Tetap Menggunakan Metode Naïve Bayes,” *J. TIKomSiN*, vol. 12, no. 2, pp. 7–15, Oct. 2024.
- [6] F. Ramadhan and H. D. Bhakti, “Klasifikasi Penilaian Kinerja Karyawan Menggunakan Algoritme Naïve Bayes (Studi Kasus PT. As Sabar Sukses Berkah),” *Kohesi: J. Multidisiplin Saintek*, vol. 4, no. 2, pp. 33–42, Jul. 2024.
- [7] S. Rukmini, Nurchim, and D. Hartanti, “Pendekatan Naïve Bayes dalam Klasifikasi Penilaian Kinerja Pegawai ASN di Sragen,” *JATI: J. Mahasiswa Tek. Inform.*, vol. 9, no. 2, pp. 2811–2820, Apr. 2025.
- [8] A. Arifin and N. Nuraini, “Implementasi Naïve Bayes Classifier dalam Memprediksi Kelulusan Mahasiswa,” *J. Teknol. Inf. dan Ilmu Komputer*, vol. 8, no. 3, pp. 245–253, Sep. 2022.
- [9] Y. S. Sari, L. Fitriani, and M. Yusuf, “Perbandingan Algoritma Naïve Bayes dan C4.5 dalam Prediksi Keputusan Karyawan untuk Meninggalkan Perusahaan,” *J. Ilm. Tek. & Rekayasa*, vol. 10, no. 1, pp. 55–62, 2023.
- [10] D. Handayani and R. Pratama, “Studi Perbandingan Algoritma Naïve Bayes dan Decision Tree pada Penilaian Kinerja Karyawan,” *J. Komput. dan Inform.*, vol. 6, no. 1, pp. 44–52, 2023.
- [11] R. Wati and M. Hidayat, “Pemanfaatan Sistem Informasi Kepegawaian dalam Peningkatan Kinerja Aparatur Sipil Negara,” *J. Inform. Publik*, vol. 7, no. 1, pp. 55–64, 2023.
- [12] G. P. Santos and M. Oliveira, “Predicting Employee Productivity with Naïve Bayes: A Case Study,” *Int. J. Comput. Appl.*, vol. 184, no. 16, pp. 12–18, 2022.
- [13] M. Oliveira and J. Costa, “Application of Naïve Bayes Classifier in Academic Performance Evaluation,” in *Proc.Int. Conf. Educ. Data Mining*, 2022, pp. 221–228.
- [14] L. Ningsih and R. Pratama, “Penerapan Algoritma Naïve Bayes untuk Prediksi Prestasi Akademik Mahasiswa,” *J. Sist. Inf. dan Komputasi*, vol. 12, no. 2, pp. 88–96, 2021.



- [15] M. Yusuf and R. Arifin, "Model Prediksi Kelulusan Mahasiswa Menggunakan Naïve Bayes Berbasis Web," *J. Inform. dan Teknol.*, vol. 9, no. 2, pp. 133–140, 2021.
- [16] A. Yusnita, S. Lailiyah, and K. Saumahudi, "Penerapan Algoritma Naïve Bayes untuk Penerimaan Peserta Didik Baru," *Informatika*, vol. 11, no. 1, pp. 11–16, 2020.
- [17] A. Azahari, Yulindawati, D. Rosita, and S. Mallala, "Komparasi Data Mining Naïve Bayes dan Neural Network Memprediksi Masa Studi Mahasiswa S1," *J. Teknol. Inf. dan Ilmu Komputer (JTIK)*, vol. 7, no. 3, pp. 443–452, Jun. 2020.
- [18] D. Nuryansah and M. Ary, "Implementasi Algoritma Naïve Bayes Classifier untuk Memprediksi Tingkat Produktivitas Kinerja Karyawan," *Jurnal Informatika UMT (JIKA)*, 2024.