

KLASIFIKASI KELAYAKAN SISWA PENERIMA PIP MENGGUNAKAN ALGORITMA NAIVE BAYES SMP NEGERI 1 LEMAHABANG

Siti Nur Fadilah¹, Willy Prihartono², Nining Rahaninsih³, Kaslani⁴

Program Studi Komputerisasi Akuntansi¹²³⁴

STMIK IKMI Cirebon
<https://ikmi.ac.id/page/18/?lang=de>
fdlhdila17@gmail.com

(*) Corresponding Author : fdlhdila17@gmail.com
Published : 30 Januari 2026

Abstract—The Smart Indonesia Program (PIP) is a government assistance program aimed at helping students from underprivileged families continue their education. However, inaccurate targeting is still found in its implementation due to the manual selection process. Based on observations at SMP Negeri 1 Lemahabang, the determination of PIP recipient eligibility still relies on administrative data and interviews without data-based analysis, resulting in a mismatch between parents' economic conditions and the status of recipients. The root of this problem is the suboptimal use of students' socioeconomic data as a basis for decision-making. Therefore, this study applies a data mining method using the Naïve Bayes algorithm to classify the eligibility of PIP recipients. The data used were 454 student data, which were then processed through the Knowledge Discovery in Databases (KDD) stages, including data selection, data cleaning, data transformation, data mining, and evaluation, with the help of the RapidMiner application. The purpose of this study was to objectively predict the eligibility of PIP recipients and to determine the accuracy level of the Naïve Bayes algorithm. The results showed that the Naïve Bayes algorithm produced an accuracy value of 61.06%, a recall value of 68.25% for the eligible class and 52.00% for the ineligible class, and an AUC value of 0.633, which is included in the poor classification category. These results indicate that the Naïve Bayes algorithm can be used as a decision support tool in determining the eligibility of PIP recipients, although it still needs further development to improve its performance.

Keyword: Smart Indonesia Program, Data Mining, Naïve Bayes, Classification

Abstrak—Program Indonesia Pintar (PIP) merupakan bantuan pemerintah yang bertujuan membantu siswa dari keluarga kurang mampu agar dapat melanjutkan pendidikan, namun dalam pelaksanaannya masih ditemukan ketidaktepatan sasaran akibat proses seleksi yang dilakukan secara manual. Berdasarkan hasil observasi di SMP Negeri 1 Lemahabang, penentuan kelayakan penerima PIP masih bergantung pada data administrasi dan wawancara tanpa analisis berbasis data, sehingga menimbulkan ketidaksesuaian antara kondisi ekonomi orang tua dengan status penerimaan bantuan. Akar permasalahan tersebut adalah belum optimalnya pemanfaatan data sosial ekonomi siswa sebagai dasar pengambilan keputusan. Oleh karena itu, penelitian ini menerapkan metode data mining menggunakan algoritma Naïve Bayes untuk mengklasifikasikan kelayakan siswa penerima PIP. Data yang digunakan sebanyak 454 data siswa yang kemudian diproses melalui tahapan Knowledge Discovery in Databases (KDD), meliputi seleksi data, pembersihan data, transformasi data, data mining, dan evaluasi, dengan bantuan aplikasi RapidMiner. Tujuan penelitian ini adalah untuk memprediksi kelayakan penerima PIP secara objektif serta mengetahui tingkat akurasi algoritma Naïve Bayes. Hasil penelitian menunjukkan bahwa algoritma Naïve Bayes menghasilkan nilai akurasi sebesar 61,06%, nilai recall sebesar 68,25% untuk kelas layak dan 52,00% untuk kelas tidak layak, serta nilai AUC sebesar 0,633 yang termasuk dalam kategori *poor classification*. Hasil tersebut menunjukkan bahwa algoritma Naïve Bayes dapat digunakan sebagai alat bantu pendukung keputusan dalam menentukan kelayakan penerima PIP, meskipun masih perlu pengembangan lebih lanjut untuk meningkatkan kinerjanya.

Kata Kunci : Program Indonesia Pintar, Data Mining, Naïve Bayes, Klasifikasi

INTRODUCTION

Pendidikan memiliki peran penting dalam meningkatkan kualitas hidup masyarakat dengan menyediakan akses terhadap pengetahuan, keterampilan, dan peluang ekonomi. Namun, keterbatasan kondisi ekonomi keluarga yang memburuk telah menjadi hambatan bagi siswa untuk menyelesaikan Pendidikan mereka. Pemerintah Indonesia berupaya mengatasi masalah tersebut melalui Program Indonesia Pintar (PIP), yaitu bantuan pendidikan bagi siswa dari keluarga berpenghasilan rendah agar tetap bersekolah. Permasalahan ketidaktepatan sasaran dalam bantuan pendidikan yang tidak akurat juga telah diidentifikasi dalam berbagai studi beasiswa, termasuk penelitian [1] seleksi bantuan pendidikan tanpa dukungan sistem terkomputerisasi dapat menyebabkan bantuan tidak sepenuhnya diberikan kepada pihak yang benar-benar berhak, sehingga diperlukan penggunaan mekanisme seleksi berbasis data untuk meningkatkan akurasi dan objektivitas penetapan penerima bantuan [2], [3].

Data mining adalah proses menganalisis data dalam jumlah besar untuk menemukan pola, informasi, dan pengetahuan sebagai dasar pengambilan keputusan. Salah satu metode data mining paling populer adalah Naïve Bayes, yaitu algoritma klasifikasi berbasis probabilitas yang efektif karena menghasilkan prediksi cepat, sederhana, dan memiliki tingkat akurasi baik, seperti yang ditunjukkan dalam penelitian [4], [5] yang mencapai akurasi klasifikasi sebesar 76,39%. Dalam konteks penelitian ini, metode Naïve Bayes digunakan untuk mengklasifikasikan siswa yang layak atau tidak layak dalam menerima PIP berdasarkan variabel pekerjaan orang tua, dan penghasilan orang tua. Permasalahan utama yang diangkat dalam tugas akhir ini adalah masih banyaknya ketidaktepatan dalam proses kelayakan penerima PIP di SMP Negeri 1 Lemahabang akibat metode seleksi manual yang belum didukung oleh analisis berbasis data.

Penelitian oleh [5], [6] menunjukkan bahwa algoritma Naïve Bayes dapat secara akurat mengklasifikasikan kelayakan penerima bantuan Program Indonesia Pintar (PIP) berdasarkan variabel kondisi sosial ekonomi keluarga, seperti status keluarga, kepemilikan KIP, jumlah tanggungan, dan penghasilan orang tua. Hasil evaluasi penelitian menunjukkan tingkat akurasi yang tinggi yang mengindikasikan bahwa metode ini dinilai efektif dalam mengurangi ketidaktepatan seleksi yang sebelumnya dilakukan secara manual. Penelitian serupa dilakukan oleh [7], [8] yang mengklasifikasi calon penerima PIP menggunakan metode Naive Bayes, dan menemukan bahwa algoritma ini dapat meningkatkan objektivitas dan

konsistensi dalam proses penentuan kelayakan bantuan. Selain itu, penelitian [9], [10] berjudul *"Penentuan Penerima Bantuan Beasiswa KIP Kuliah Menggunakan Naive Bayes Classifier"* menunjukkan bahwa Naive Bayes dapat memberikan akurasi tinggi dalam proses seleksi KIP Kuliah. Ketiga penelitian tersebut menegaskan bahwa metode Naive Bayes merupakan algoritma yang cocok, efektif, dan relevan untuk digunakan dalam penilaian penerima PIP di SMP.

Tabel 1 Data Siswa Penerima PIP

Nama Siswa	Peke rjaa n Ayah	Peke rjaa n Ibu	Peng hasil an Ayah	Pengh asilan Ibu	Pen eri ma PIP
SHEILA NURFI TRIYANI	Pedagang Kecil	Tidak bekerja	Rp. 1,000,000 - Rp. 1,999,999	Tidak Berpenghasilan	Laya k
CINDY ANA PUTRI AISYAH	Lain nya	Tidak bekerja	Kurang dari Rp. 500,000	Tidak Berpenghasilan	Tidak Laya k
MAWAR LESTARI	Karyawan Swasta	Tidak bekerja	Rp. 500,000 - Rp. 999,999	Tidak Berpenghasilan	Laya k
FARHAN HABIBI	Wiraswasta	Wiraswasta	Rp. 2,000,000 - Rp. 4,999,999	Rp. 500,000 - Rp. 999,999	Tidak Laya k
ALIF SYAHBANI FAIRUZ	Karyawan Swasta	Tidak bekerja	Rp. 1,000,000 - Rp. 1,999,999	Tidak Berpenghasilan	Tidak laya k
NAYLA TRI APRILIA	Karyawan Swasta	Tidak bekerja	Rp. 1,000,000 - Rp. 1,999,999	Tidak Berpenghasilan	Tidak laya k
ANNISA ZEIDIL LA ULANI	Pedagang Kecil	Pedagang Kecil	Rp. 1,000,000 - Rp.	Rp. 1,000,000 - Rp.	Tidak laya k

			1,999,999	1,999,999	
MUHAMMAD FIRDA NUGROHO DIKRY PRASE TIO	Karyawan Swastara	Tidak bekerja	Rp. 500,000 - Rp. 999,999	Tidak Berpenghasilan	Tidak layak
MUHAMMAD BILAL FAQIH	Karyawan Swastara	Tidak bekerja	Rp. 500,000 - Rp. 999,999	Tidak Berpenghasilan	Tidak layak

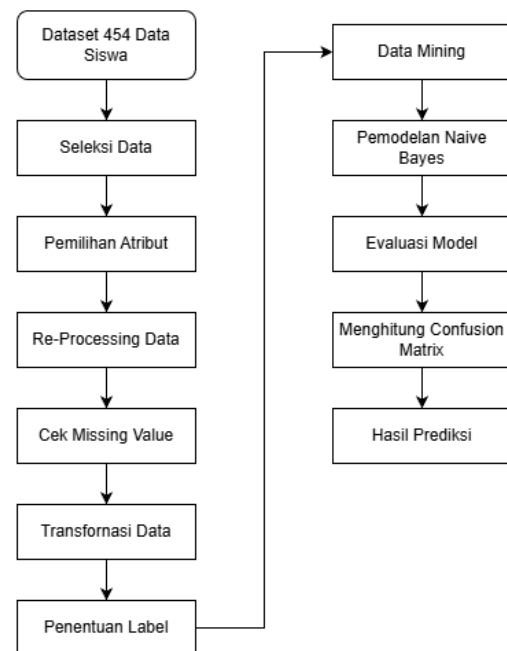
Sumber : Data SMP Negeri 1 Lemahabang

Berdasarkan data table 1.1 tentang data siswa penerima PIP adalah sebagian data siswa yang diajukan penerima PIP di SMP Negeri 1 Lemahabang dengan spesifikasi Nama Siswa, Pekerjaan Ayah, Pekerjaan Ibu, Penghasilan Ayah, Penghasilan Ibu, dan Penerima PIP. Masih banyak siswa yang belum menerima PIP dengan tepat sasaran.

Status kelayakan penerima PIP tidak selalu sesuai dengan kondisi ekonomi yang ditunjukkan oleh data pekerjaan dan penghasilan orang tua. Misalnya, beberapa siswa dengan penghasilan ayah di bawah Rp 1.000.000 masih disebut "Tidak Layak", sementara siswa yang kondisi keuangannya bisa disebut "Layak". Perbedaan ini menunjukkan adanya inkonsistensi dalam proses seleksi, baik karena evaluasi manual maupun kurangnya standar berbasis data. Fakta ini menyoroti perlunya metode klasifikasi yang sistematis agar penyaluran PIP dapat lebih objektif, tepat, dan adil bagi siswa yang membutuhkan.

MATERIALS AND METHODS

Ada beberapa langkah dalam penerapan Knowledge Discovery in Database (KDD), termasuk:



Gambar 1 Proses KDD

1. Seleksi Data (*Data Selection*)

Seleksi data adalah tahap awal yang melibatkan pemilihan operasional data yang relevan sebelum dimulainya proses *mining* atau penggalian informasi dalam tahapan KDD. Data yang telah diseleksi kemudian akan digunakan dalam proses *data mining* dan disimpan dalam suatu file terpisah sehingga tidak mengganggu operasional data yang ada.

2. Pembersihan (*Preprocessing*)

Tahap *preprocessing* mencakup proses pembersihan data dengan cara menghapus data yang duplikat, memeriksa dan memperbaiki data yang tidak konsisten, serta memperbaiki kesalahan seperti kesalahan ketik. Selain itu, tahap ini juga melibatkan **pengayaan** atau penambahan data yang ada dengan informasi tambahan yang relevan dan dibutuhkan dalam proses KDD.

3. Transformasi (*Transformation*)

Pada tahap ini, data yang belum memiliki format atau entitas yang jelas diubah menjadi format yang sesuai dan sah untuk dianalisis lebih lanjut dalam proses *data mining*. Transformasi bertujuan mengonversi data ke dalam bentuk format yang dapat diproses oleh algoritma yang digunakan. Proses ini sangat bergantung pada jenis data dan model yang digunakan dalam penelitian.

Menurut penelitian yang dilakukan oleh (Oktavia et al., 2024), tahap transformasi dalam proses KDD memungkinkan transformasi data mentah ke dalam format yang lebih terstruktur dan mudah dianalisis. Hal ini mencakup konversi data

ke format yang diperlukan agar dapat digunakan dengan teknik *data mining* seperti klasifikasi Naïve Bayes. Selain itu, data yang sudah melalui tahap transformasi juga bisa mencakup proses **penggabungan** atribut atau informasi baru yang relevan dengan tujuan analisis, guna memberikan pemahaman yang lebih komprehensif tentang pola-pola dalam kumpulan data.

4. Data Mining

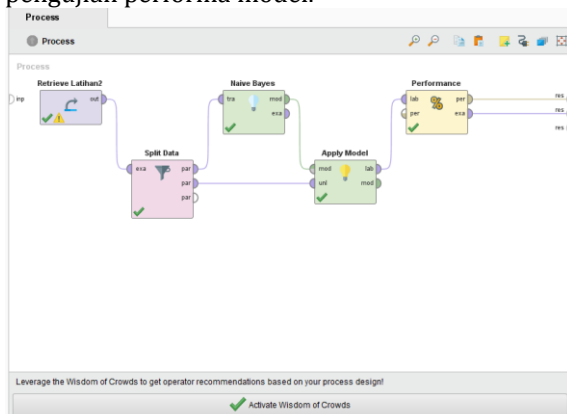
Pada tahap ini, diterapkan metode atau algoritma yang sesuai untuk memperoleh informasi atau pengetahuan yang diinginkan. Proses ini melibatkan penggunaan beberapa teknik seperti klasifikasi, estimasi, prediksi, keterhubungan, klastering, dan berbagai teknik lainnya, tergantung pada jenis informasi yang perlu diekstraksi dari data yang tersedia.

5. Evaluasi (*Evaluation*)

Tahap ini merupakan langkah terakhir dalam proses KDD, di mana hasil atau output dari proses *data mining* dievaluasi. Tujuan evaluasi adalah untuk memastikan bahwa informasi atau kebijakan yang diperoleh dapat dipahami dengan jelas dan mudah oleh pihak yang terkait. Pada tahap ini, hasil dari analisis akan dianalisis lebih detail untuk memastikan relevansi dan akurasi sebelum digunakan untuk pengambilan keputusan atau langkah selanjutnya.

RESULTS AND DISCUSSION

Proses ini terdiri dari beberapa tahapan yang saling terhubung, mulai dari pengambilan data, pembagian data, pembuatan model, hingga pengujian performa model.



Gambar 2 Model Naïve Bayes

Gambar 2 menunjukkan alur pemodelan klasifikasi penerima PIP menggunakan algoritma Naïve Bayes pada aplikasi RapidMiner. Proses diawali dengan operator Retrieve yang berfungsi untuk mengambil data latih yang telah melalui tahap seleksi,

pembersihan, dan transformasi data. Selanjutnya, dataset dibagi menjadi data latih dan data uji menggunakan operator Split Data. Data training digunakan pada operator Naïve Bayes untuk membangun model klasifikasi, sedangkan data testing diproses pada tahap Apply Model untuk menghasilkan prediksi kelayakan penerima PIP. Hasil prediksi tersebut kemudian dievaluasi menggunakan operator Performance untuk mengetahui kinerja model, termasuk nilai akurasi yang dihasilkan.

	true Layak	true Tidak layak	class precision
pred. Layak	43	24	64.18%
pred. Tidak layak	20	26	56.52%
class recall	68.25%	52.00%	

Gambar 3 Hasil Hitung Rapidminer

Gambar 3 menyajikan hasil evaluasi kinerja model klasifikasi penerima PIP menggunakan algoritma Naïve Bayes yang ditampilkan dalam bentuk confusion matrix. Confusion matrix digunakan untuk membandingkan hasil prediksi model dengan kondisi data sebenarnya pada masing-masing kelas.

Berdasarkan hasil tersebut, pada kelas Layak terdapat 43 data yang berhasil diprediksi dengan benar oleh model, sedangkan 24 data lainnya diprediksi sebagai Layak namun pada kondisi sebenarnya termasuk ke dalam kelas Tidak Layak. Pada kelas Tidak Layak, model mampu memprediksi 26 data secara tepat, namun masih terdapat 24 data yang diprediksi Tidak Layak meskipun pada kondisi sebenarnya termasuk kelas Layak.

Nilai precision untuk kelas Layak sebesar 64,18% sementara itu, nilai precision pada kelas Tidak Layak sebesar 56,52%. Adapun nilai recall pada kelas Layak sebesar 68,25%. Namun, nilai recall pada kelas Tidak Layak yang hanya sebesar 52,00%.

Evaluasi kinerja model klasifikasi pada penelitian ini dilakukan menggunakan confusion matrix yang terdiri dari empat komponen utama, yaitu True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN), berikut perhitungan manual pada confusion matrix.

Tabel 1 Pengukuran Confusion Matrix

Hasil Prediksi	Kondisi Aktual: Layak	Kondisi Aktual: Tidak Layak
Diprediksi Layak	True Positive (TP)	False Positive (FP)
Diprediksi Tidak Layak	False Negative (FN)	True Negative (TN)

Berikut keterangan yang ada pada tabel 1. True Positive (TP) Adalah data siswa yang sebenarnya layak menerima PIP dan memang benar diprediksi layak dalam menerima PIP. False Positive (FP) Adalah data siswa yang sebenarnya tidak layak menerima PIP namun saat diprediksi ternyata layak dalam menerima PIP. False Negative (FN) Adalah data siswa yang sebenarnya layak menerima PIP namun diprediksi ternyata tidak layak dalam menerima PIP. True Negative (TN) Adalah data siswa yang sebenarnya tidak layak menerima PIP dan memang benar saat diprediksi tidak layak dalam menerima PIP.

Riwayat proses ini menunjukkan bahwa seluruh tahapan pemodelan, mulai dari pembagian data, penerapan algoritma Naïve Bayes, hingga evaluasi kinerja model, telah berhasil dijalankan tanpa mengalami kesalahan. Informasi waktu eksekusi yang tercatat juga 0 menunjukkan bahwa proses berjalan secara efisien.

Criterion	accuracy
accuracy	61.06%

Gambar 2 Hasil Accuracy

Gambar 2 menampilkan hasil evaluasi akurasi model klasifikasi penerima Program Indonesia Pintar (PIP) menggunakan algoritma Naïve Bayes pada RapidMiner. Berdasarkan hasil pengujian terhadap data uji, diperoleh nilai akurasi sebesar 61,06%. Nilai ini menunjukkan bahwa dari seluruh data yang diuji, sekitar 61,06% data berhasil diklasifikasikan dengan benar oleh model.

Precision merupakan salah satu metrik evaluasi yang digunakan untuk mengukur tingkat ketepatan hasil klasifikasi, yaitu sejauh mana data yang diprediksi oleh sistem benar-benar sesuai dengan kelas sebenarnya. Nilai precision pada penelitian ini ditunjukkan pada Gambar 4.10.

Criterion	precision
precision	56.52% (positive class: Tidak Layak)

Gambar 3 Hasil Precision

Gambar 3 menunjukkan hasil pengukuran precision dari model klasifikasi penerima PIP menggunakan RapidMiner. Pada pengujian ini, nilai precision yang ditampilkan sebesar 56,52% dengan kelas tidak layak sebagai kelas positif. Nilai precision menggambarkan tingkat ketepatan model

dalam memberikan prediksi terhadap suatu kelas tertentu.

Hasil ini berarti bahwa dari seluruh data yang diprediksi sebagai tidak layak oleh model, sekitar 56,52% di antaranya benar-benar termasuk ke dalam kategori tidak layak, sedangkan sisanya merupakan data yang salah diklasifikasikan. Dengan demikian, model masih memiliki kesalahan prediksi pada kelas tidak layak, namun tetap mampu memberikan hasil klasifikasi dengan tingkat ketepatan yang cukup baik.

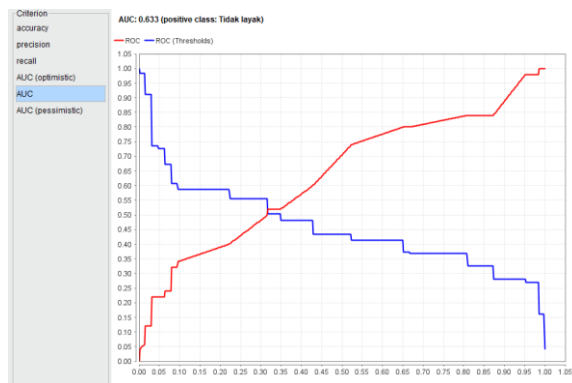
Recall adalah salah satu metrik evaluasi yang digunakan untuk menilai seberapa baik model klasifikasi dalam mengidentifikasi data yang benar di setiap kelas. Nilai recall menggambarkan seberapa besar proporsi data yang berhasil dikenali dengan akurat dibandingkan dengan seluruh data di kelas tersebut. Gambar 4.11 menunjukkan hasil nilai recall dalam penelitian ini.

Criterion	recall
recall	52.00% (positive class: Tidak Layak)

Gambar 3 Hasil Recall

Hasil data yang telah diolah dan ditunjukkan pada Gambar 3 memperlihatkan nilai recall yang didapat dari proses klasifikasi dengan RapidMiner. Nilai recall untuk kelas layak sebesar 68,25%, menunjukkan bahwa sistem dapat mengidentifikasi sebagian besar siswa yang benar-benar layak menerima PIP dengan tingkat keberhasilan yang cukup baik. Di sisi lain, nilai recall untuk kelas tidak layak sebesar 52,00%. Ini menunjukkan bahwa sistem masih kurang optimal dalam mendeteksi siswa yang seharusnya tidak menerima PIP, karena masih terdapat banyak data yang tidak layak tetapi terklasifikasi sebagai layak.

AUC (Area Under Curve) digunakan sebagai ukuran performa model klasifikasi dalam menilai kemampuan pemisahan antar kelas melalui kurva ROC. Semakin besar nilai AUC yang dihasilkan, maka semakin baik kemampuan model dalam membedakan data pada masing-masing kelas.



Gambar 4 Hasil AUC

Algoritma Naïve Bayes digunakan untuk memprediksi kelayakan calon penerima Program Indonesia Pintar (PIP) dan menghasilkan nilai akurasi sebesar 61,06%, serta nilai AUC sebesar 0,633.

CONCLUSION

Berdasarkan hasil penelitian maka dapat ditarik beberapa kesimpulan sebagai berikut :

Pengolahan data dalam penelitian ini dilakukan menggunakan data siswa calon penerima Program Indonesia Pintar (PIP) di SMP Negeri 1 Lemahabang. Data yang digunakan merupakan data administrasi siswa yang diperoleh dari pihak sekolah dan telah melalui tahapan seleksi, pembersihan, serta transformasi data sebelum dilakukan proses pemodelan.

Hasil pemodelan klasifikasi menggunakan algoritma Naïve Bayes menunjukkan bahwa sistem mampu memprediksi kelayakan penerima PIP dengan nilai akurasi sebesar 61,06%. Berdasarkan hasil confusion matrix, diperoleh 43 data yang diprediksi layak dan sesuai dengan kondisi sebenarnya TP, serta 26 data yang diprediksi tidak layak dan sesuai dengan kondisi aktual TN. Sementara itu, terdapat 24 data yang diprediksi layak namun sebenarnya tidak layak FP dan 20 data yang diprediksi tidak layak namun sebenarnya layak FN.

REFERENCE

- [1] A. Gusderia, M. Ramadhan, and I. Perangin-Angin, "Data Mining Untuk Klasifikasi Data Penjualan Alat Teknik Menggunakan Metode Naive Bayesian Clacifier," *Jurnal Sains Manajemen Informatika dan Komputer*, vol. 21, no. 2, pp. 73–79, 2022, [Online]. Available: <https://ojs.trigunadharma.ac.id/index.php/jis>
- [2] Nadia Tiara Rahman, "ANALISA ALGORITMA DECISION TREEDAN NAÏVE BAYESPADA PASIEN PENYAKIT LIVER," *Jurnal Fasilkom*, 2020.
- [3] Nadia Tiara Rahman, "ANALISA ALGORITMA DECISION TREEDAN NAÏVE BAYESPADA PASIEN PENYAKIT LIVER," *Jurnal Fasilkom*, 2020.
- [4] I. Iwandini, A. Triayudi, and G. Soepriyono, "Analisa Sentimen Pengguna Transportasi Jakarta Terhadap Transjakarta Menggunakan Metode Naives Bayes dan K-Nearest Neighbor," *Journal of Information System Research (JOSH)*, vol. 4, no. 2, pp. 543–550, Jan. 2023, doi: 10.47065/josh.v4i2.2937.
- [5] I. Iwandini, A. Triayudi, and G. Soepriyono, "Analisa Sentimen Pengguna Transportasi Jakarta Terhadap Transjakarta Menggunakan Metode Naives Bayes dan K-Nearest Neighbor," *Journal of Information System Research (JOSH)*, vol. 4, no. 2, pp. 543–550, Jan. 2023, doi: 10.47065/josh.v4i2.2937.
- [6] E. Mardiani *et al.*, "Analisis Kompleksitas Password Dengan Metode KNN, Naïve Bayes, Decision Tree, Ensemble Methods Dan Linear Regression," *Digital Transformation Technology*, vol. 3, no. 2, pp. 955–966, Jan. 2024, doi: 10.47709/digitech.v3i2.3513.
- [7] E. Mardiani *et al.*, "Analisis Kompleksitas Password Dengan Metode KNN, Naïve Bayes, Decision Tree, Ensemble Methods Dan Linear Regression," *Digital Transformation Technology*, vol. 3, no. 2, pp. 955–966, Jan. 2024, doi: 10.47709/digitech.v3i2.3513.
- [8] A. Kaharudin, A. Agus Supriyadi, H. Baitika, and M. Derryanur, "OKTAL: Jurnal Ilmu Komputer dan Science Analisis Sentimen pada Media Sosial dengan Teknik Kecerdasan Buatan Naïve Bayes: Kajian Literatur Review," vol. 2, no. 6, 2023, [Online]. Available: <https://harzing.com/resources/publish-or-perish>
- [9] A. Kaharudin, A. Agus Supriyadi, H. Baitika, and M. Derryanur, "OKTAL: Jurnal Ilmu Komputer dan Science Analisis Sentimen pada Media Sosial dengan Teknik Kecerdasan Buatan Naïve Bayes: Kajian Literatur Review," vol. 2, no. 6, 2023, [Online]. Available: <https://harzing.com/resources/publish-or-perish>
- [10] Zairi saputra, H. A. Supahri, R. Ismanizan, and R. Rahmadden, "Perbandingan Algoritma KNN (K-Nearest Neighbors),

Naïve Bayes, Dan SVM (Support Vector Machine) Untuk Klasifikasi Pemberian Pinjaman Nasabah," *Jurnal Ilmiah Sistem Informasi dan Teknik Informatika (JISTI)*,

vol. 7, no. 1, pp. 67–75, Apr. 2024, doi:
10.57093/jisti.v7i1.182.