



SNESTIK

Seminar Nasional Teknik Elektro, Sistem Informasi,
dan Teknik Informatika

<https://ejurnal.itats.ac.id/snestik> dan <https://snestik.itats.ac.id>



Informasi Pelaksanaan :

SNESTIK V - Surabaya, 26 April 2025

Fakultas Teknik Elektro dan Teknologi Informasi, Institut Teknologi Adhi Tama Surabaya

Informasi Artikel:

DOI : 10.31284/p.snestik.2025.7055

Prosiding ISSN 2775-5126

Fakultas Teknik Elektro dan Teknologi Informasi-Institut Teknologi Adhi Tama Surabaya
Gedung A-ITATS, Jl. Arief Rachman Hakim 100 Surabaya 60117 Telp. (031) 5945043
Email : snestik@itats.ac.id

OPTIMASI MODEL REKOMENDASI TOPIK SKRIPSI BERDASARKAN PERFORMA AKADEMIK MAHASISWA MENGUNAKAN SMOTE

Nelly Oktavia Adiwijaya, Muhammad Farhan Al Abror, Tio Dharmawan,
Muhamad Arief Hidayat

Fakultas Ilmu Komputer, Universitas Jember

Email : nelly.aa@unej.ac.id

ABSTRACT

About 68% of students at the Faculty of Computer Science, University of Jember, experience delays in completing their thesis, mainly due to difficulty choosing a research topic that suits their interests and expertise. The main problem is the mismatch between students' academic competence and the chosen thesis topic. Many students choose topics based on external factors such as research trends, friends' recommendations, or lecturers' suggestions without considering their academic abilities. This problem is exacerbated by the absence of systematic guidance in choosing topics, thus increasing the risk of research stagnation and even changing topics midway. This research aims to develop a thesis topic recommendation system based on students' academic data, specifically transcripts from semesters 1 to 6, as an indicator of academic ability. The system uses Support Vector Machine (SVM) and Naïve Bayes classification algorithms, with the SMOTE method applied to overcome class imbalance and improve model performance. The analysis results show that the best model is achieved with SVM kernel RBF, producing the highest accuracy of 96.81%, while Naïve Bayes Categorical produces an accuracy of 83.75%.

Keywords: Topic Recommendations, Support vector Machine, Naïve Bayes, SMOTE

ABSTRAK

Sekitar 68% mahasiswa di Fakultas Ilmu Komputer Universitas Jember mengalami keterlambatan dalam menyelesaikan skripsi, yang mencerminkan kesulitan dalam menentukan topik penelitian yang sesuai dengan minat dan keahlian. Salah satu masalah utama yang ditemukan adalah ketidaksesuaian antara kompetensi akademik mahasiswa dan topik skripsi yang dipilih. Banyak mahasiswa memilih topik berdasarkan faktor

eksternal seperti tren penelitian, rekomendasi teman, atau dorongan dosen, tanpa mempertimbangkan kesesuaian dengan kemampuan akademik yang dimiliki, dan diperburuk dengan kurangnya panduan yang sistematis dalam pemilihan topik, sehingga meningkatkan risiko stagnasi dalam proses penelitian bahkan mengganti topik di tengah jalan. Penelitian ini bertujuan untuk membangun sistem rekomendasi topik skripsi berbasis pada data akademik mahasiswa, khususnya transkrip nilai mata kuliah yang diambil dari semester 1 hingga semester 6, sebagai indikator kemampuan akademik mahasiswa. Sistem ini diharapkan dapat membantu mahasiswa dalam memilih topik yang lebih sesuai dengan keahlian yang dimiliki, sehingga meningkatkan keberhasilan dalam penyelesaian skripsi. Untuk membangun sistem rekomendasi, digunakan algoritma klasifikasi Support Vector Machine (SVM) dan Naïve Bayes, dengan metode SMOTE diterapkan untuk menangani masalah ketidakseimbangan data antar kelas dan meningkatkan akurasi model. Hasil analisis menunjukkan bahwa model terbaik diperoleh dengan SVM kernel RBF, yang menghasilkan akurasi tertinggi sebesar 96,81%, sedangkan Naïve Bayes tipe Categorical menghasilkan akurasi 83,75%.

Kata kunci: Rekomendasi Topik Skripsi, Support Vector Machine, Naïve Bayes, SMOTE

PENDAHULUAN

Skripsi merupakan karya ilmiah yang dihasilkan dari penelitian mahasiswa secara mandiri [1]. Dalam proses pengerjaan skripsi, beberapa mahasiswa membutuhkan waktu dua semester atau lebih untuk menyelesaikan skripsinya [2]. Mahasiswa sering mengalami kesulitan dalam menentukan topik skripsi. Sekitar 68% mahasiswa mengalami keterlambatan dalam menyelesaikan skripsinya, hal ini berarti mahasiswa dalam menentukan topik penelitian tidak sesuai dengan minat dan keahlian yang dimiliki, kebanyakan dari mereka hanya mengikuti tren bidang penelitian dan topik skripsi [3]. Fenomena ini akan mempengaruhi kinerja mahasiswa dalam menyelesaikan studinya tepat waktu.

Permasalahan tersebut dapat diatasi dengan membangun sebuah model klasifikasi yang dapat membantu mahasiswa untuk menentukan topik skripsi berdasarkan kemampuan yang dimilikinya [3]. Indikator yang digunakan dalam menentukan kemampuan mahasiswa dapat menggunakan data akademik [1]. Data akademik berupa transkrip nilai mata kuliah yang diambil oleh mahasiswa dari semester 1 hingga semester 6.

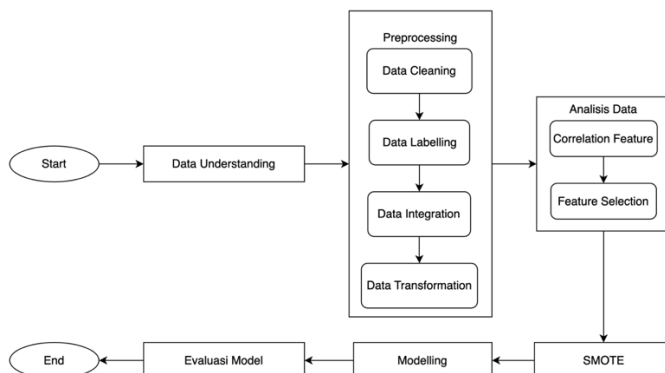
Klasifikasi merupakan pendekatan data mining yang paling sering digunakan dalam pendidikan tinggi untuk memprediksi kinerja, prestasi, pengetahuan, dan tingkat putus sekolah [4]. Algoritma yang dapat digunakan dalam teknik prediksi antara lain Support Vector Machine (SVM) dan Naïve Bayes. Algoritma SVM memiliki performa yang baik dalam mengklasifikasikan suatu pola dan mencegah terjadinya masalah dimensionalitas [5]. Penerapan algoritma SVM untuk prediksi kelulusan mahasiswa tepat waktu menghasilkan model dengan akurasi sebesar 94,4% dengan perbandingan data training dan data testing sebesar 90% dan 10% [6]. Algoritma Naïve Bayes merupakan metode dengan algoritma yang sederhana namun memiliki kecepatan dan akurasi pemodelan yang tinggi [7]. Sistem pendukung keputusan rekomendasi topik skripsi menggunakan Naïve Bayes Classifier menghasilkan performa model dengan akurasi sebesar 69,27%, hasil akurasi tersebut masih belum optimal dikarenakan jumlah data yang tidak seimbang [8]. Maka dari itu pada penelitian ini data yang digunakan akan diolah terlebih dahulu menggunakan metode SMOTE sehingga didapatkan dataset yang seimbang untuk setiap kelas.

Penelitian ini bertujuan untuk menerapkan dua algoritma klasifikasi yaitu SVM dan Naive Bayes dalam memprediksi topik skripsi mahasiswa. Namun pada tahap pra pemrosesan datanya akan diterapkan *feature selection* dan SMOTE untuk menangani ketidakseimbangan dataset. Dengan adanya pengembangan model ini, diharapkan dapat membantu para mahasiswa untuk lebih mudah menemukan topik skripsi yang sesuai dengan minat dan kemampuannya.

METODE

Dalam penelitian ini, ada enam proses utama yang akan dilakukan untuk mengklasifikasikan dan memodelkan, yaitu: pemahaman data, persiapan data, analisis data,

SMOTE, pemodelan, dan evaluasi model. Proses yang dilakukan dalam penelitian ini dapat dilihat pada Gambar 1.



Gambar 1. Metodologi dalam Flowchart

Dataset yang digunakan terdiri dari data akademik dan topik skripsi mahasiswa Fasilkom Unej, yang diperoleh dari UPTTI Universitas Jember. Data akademik mencakup 2.493 mahasiswa (angkatan 2009–2021), berisi transkrip nilai yang mencerminkan capaian akademik. Data topik skripsi terdiri dari 2.121 mahasiswa (angkatan 2009–2018), yang mencatat topik penelitian yang dipilih dan dikerjakan mahasiswa. Data understanding bertujuan untuk memahami karakteristik data, mendeskripsikan setiap atribut data, kemudian mengeksplorasi data dengan melakukan visualisasi data dan ringkasan statistik. Pada tahap data processing terdapat empat proses utama yaitu *Data Cleaning* untuk membersihkan data yang tidak sesuai, *Data Labeling* untuk memberi label kelas pada data mahasiswa berdasarkan durasi lulus tepat waktu, *Data Integration* untuk menggabungkan data mahasiswa dan topik skripsi, dan *Data Transformation* untuk mengubah bentuk data menyesuaikan dengan algoritma yang digunakan.

Pada tahap data analysis, data yang telah diintegrasikan dianalisis melalui korelasi dan seleksi fitur untuk menentukan atribut relevan dalam pemodelan topik skripsi mahasiswa. Korelasi fitur dihitung menggunakan korelasi Spearman, di mana nilai ≥ 0.6 menunjukkan hubungan signifikan, sementara nilai ≤ 0.2 dianggap tidak relevan. Fitur dengan korelasi ≥ 0.4 dipilih untuk Feature Selection guna membangun model yang akurat. Analisis ini mengidentifikasi mata kuliah dari semester 1–6 yang mempengaruhi pemilihan topik skripsi mahasiswa Sistem Informasi angkatan 2009–2016. Karena seleksi fitur dapat menyebabkan ketidakseimbangan data antar topik, dilakukan balancing data dengan SMOTE untuk meningkatkan performa model klasifikasi. SMOTE (Synthetic Minority Over-sampling Technique) digunakan untuk menangani ketidakseimbangan data antara kelas mayoritas dan minoritas untuk meningkatkan performa model. Prinsip SMOTE adalah dengan memperbanyak jumlah data pada kelas minoritas agar seimbang dengan kelas mayoritas, melalui pembuatan data sintesis. Data sintesis pada SMOTE dibuat berdasarkan prinsip K-Nearest Neighbor (k-NN) melalui persamaan 1 [10][11].

$$\chi_{new} = \chi_i + (\chi_j + \chi_i) \times \alpha \quad (1)$$

χ_i = sampel minoritas yang dipilih secara acak dari kelas minoritas

χ_j = salah satu K tetangga terdekat dari sampel minoritas

α = bilangan acak antara $[0, 1]$

Selanjutnya model dibangun menggunakan algoritma Naïve Bayes (tipe Multinomial, Complement, Gaussian, Categorical) dan Support Vector Machine (kernel Linier dan Radial Basis Function). Pemodelan akan dilakukan dengan 72 skenario uji coba berdasarkan hasil seleksi fitur. Algoritma naïve bayes memprediksi kemungkinan masa depan berdasarkan pengalaman masa lalu. Naïve bayes memiliki beberapa tipe antara lain yaitu *multinomial*, *gaussian*, *complement*, dan

categorical [15]. Pada penelitian ini digunakan *naïve bayes categorical* dikarenakan semua variabel independent pada data penelitian ini adalah kategorikal. Data kategorikal dapat digunakan untuk mengklasifikasikan dua kelas atau lebih. Sedangkan algoritma Support Vector Machine memperluas konsep pemisahan *hyperplane* ke data dengan memperluas vector fitur ke dimensi yang lebih tinggi [16].

Setelah model dibangun berikutnya tahap pengujian model yang telah dibangun guna memperoleh model yang paling akurat dan andal. Evaluasi dilakukan dengan mengukur akurasi menggunakan model *k-fold cross validation* dan *F1-score*. Dalam *cross-validation*, akan dilakukan beberapa pemisahan untuk set latih dan validasi, setiap bagian yang dipisah biasanya disebut dengan *Fold*. *Cross Validation* memiliki beberapa variasi salah satunya disebut *k-fold cross-validation* [17]. Variasi ini bekerja dengan melakukan *k* kali pelatihan dan evaluasi dengan perbandingan (*k-1*) untuk set pelatihan dan 1 untuk validasi. Berikut formula untuk *F1-score* (3)-(5) [18]:

$$F1 = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (3)$$

$$\text{Precision} = \frac{\text{True Positive}}{\text{True Positive} + \text{False Positive}} \quad (4)$$

$$\text{Recall} = \frac{\text{True Positive}}{\text{True Positive} + \text{False Negative}} \quad (5)$$

HASIL DAN PEMBAHASAN

Dalam penelitian ini, proses data mining dilakukan terhadap 263 data yang telah melewati tahap *preprocessing*. Data tersebut kemudian dianalisis untuk menemukan fitur-fitur yang relevan dalam konteks topik tesis. Salah satu analisis yang dilakukan adalah menghitung nilai korelasi antar mata kuliah dari setiap topik skripsi, dengan tujuan untuk mencari mata kuliah yang akan dijadikan fitur dari topik tersebut. Data yang digunakan merupakan hasil integrasi antara data akademik mahasiswa Fasilkom UNEJ dan data topik skripsi, yang telah melalui tahap data *cleaning*, data *labelling*, transformasi data, serta *feature section* berdasarkan nilai korelasi antara mata kuliah dan topik skripsi. Hasil korelasi menunjukkan bahwa terdapat beberapa mata kuliah yang memiliki hubungan erat dengan topik skripsi tertentu, yang ditunjukkan pada Tabel 1.

Tabel 1. Hasil korelasi topik skripsi dengan mata kuliah yang berhubungan

Topik	Mata Kuliah
Data Mining	Algoritma dan Pemrograman II, Matematika Diskrit, Sistem Penunjang Pengambilan Keputusan
Enterprise Management System	Manajemen Keamanan Sistem Informasi, Bahasa Indonesia, Sistem Penunjang Pengambilan Keputusan
IT Adoption	Sistem Basis Data, Pengantar Kecerdasan Buatan, Teknik Rekayasa Perangkat Lunak
IT Evaluation	Sistem Basis Data, Pengantar Kecerdasan Buatan, Teknik Rekayasa Perangkat Lunak
IT Strategy	Kriptografi, Jaringan Komputer, Teknik Rekayasa Perangkat Lunak
Information Security	Manajemen Keamanan Sistem Informasi, Algoritma dan Pemrograman I, Dasar-dasar Sistem Informasi
Machine Learning	Data Mining, Teknik Rekayasa Perangkat Lunak, Profesional Issue
Software Construction	Sistem Penunjang Pengambilan Keputusan, Profesional Issue, Sistem Basis Data
Teori Graf	Arsitektur Komputer, Algoritma dan Pemrograman I, Interaksi Manusia dan Komputer
Text Mining	Sistem Penunjang Pengambilan Keputusan, Sistem Informasi Manajemen, Teknik Rekayasa Perangkat Lunak
UI/UX	Dasar-dasar Sistem Informasi, Arsitektur Komputer, Matematika Diskrit

Setelah hasil korelasi diperoleh, selanjutnya adalah melakukan skenario dengan metode SMOTE untuk menyeimbangkan jumlah sampel pada kelas minoritas agar tidak terjadi

ketimpangan dengan kelas mayoritas. Sebelum dilakukan SMOTE, terdapat perbedaan jumlah mahasiswa yang memilih berbagai topik skripsi, yang menyebabkan distribusi data tidak merata. Setelah dilakukan SMOTE, jumlah data pada setiap kategori menjadi lebih merata, seperti yang ditunjukkan dalam Tabel 2 yang menunjukkan bahwa SMOTE efektif dalam menyeimbangkan jumlah data di setiap kategori.

Tabel 2. Hasil korelasi topik skripsi dengan mata kuliah yang berhubungan

Skenario	Jumlah Data Sebelum SMOTE	Jumlah Data Setelah SMOTE
Korelasi keseluruhan mata kuliah	255	1155
Korelasi dengan Data Mining	143	671
Korelasi dengan Machine Learning	169	1166
Korelasi dengan IT Evaluation	202	759
Korelasi dengan Software Construction	145	693
Korelasi dengan UI/UX	202	979

Setelah dilakukan balancing data menggunakan SMOTE, dilakukan pemodelan menggunakan Support Vector Machine (SVM) dan Naïve Bayes, yang kemudian dievaluasi dengan K-Fold Cross Validation menggunakan K=5 dan K=10 dan F1-score. Pemilihan K=5 dan K=10 dilakukan karena kedua nilai ini merupakan standar umum dalam validasi silang yang memberikan keseimbangan antara bias dan variansi model [13]. Hasil evaluasi dengan K-Fold Cross Validation setelah penerapan SMOTE menunjukan peningkatan secara signifikan, terutama pada SVM dengan kernel RBF yang mencapai 96,67% pada K=10 (Tabel 3). Selain K-Fold Cross Validation, evaluasi juga dilakukan menggunakan F1-score. Hasil evaluasi menunjukkan bahwa model SVM dengan kernel RBF setelah SMOTE mencapai hasil terbaik, dengan F1-Score sebesar 96,81%, sementara Naïve Bayes tipe Categorical mencapai 88,33% (Tabel 4).

Tabel 3. Hasil akurasi terbaik dengan K-Fold

K-Fold Cross Validation	tanpa SMOTE		dengan SMOTE	
	SVM	Naïve Bayes	SVM	Naïve Bayes
K=5	52,04%	53,72%	95,31%	84,15%
K=10	53,00%	53,64%	96,67%	83,75%

Tabel 4. Hasil akurasi terbaik dengan F1-score

tanpa SMOTE		dengan SMOTE	
SVM	Naïve Bayes	SVM	Naïve Bayes
59,09%	56,82%	96,81%	88,33%

Berdasarkan tabel 3 dan tabel 4 terlihat bahwa penerapan SMOTE mampu meningkatkan performa model secara signifikan baik pada algoritma SVM maupun Naive bayes. Sebelum penerapan SMOTE, akurasi yang diperoleh dengan K-Fold Cross Validation berada pada kisaran 52-53% untuk SVM dan 53% untuk naive bayes. Setelah penerapan SMOTE, terjadi peningkatan akurasi yang cukup signifikan terutama pada algoritma SVM dengan akurasi mencapai 96,67% untuk K=10. Hal serupa juga terlihat pada evaluasi menggunakan F1-score dimana setelah menerapkan SMOTE akurasi terbaik pada algoritma SVM mencapai 96,81% dibandingkan dengan 59,09% tanpa menggunakan SMOTE. Peningkatan ini menunjukkan bahwa SMOTE berhasil menyeimbangkan data sehingga model mampu menangkap pola dari kelas minoritas dengan lebih baik. Sementara itu, naive bayes menunjukkan peningkatan yang lebih moderat dengan F1-score mencapai 88,33% dengan penerapan SMOTE.

KESIMPULAN

Hasil penelitian ini menunjukkan bahwa penerapan SMOTE berperan krusial dalam menangani data yang tidak seimbang sehingga mampu meningkatkan akurasi model klasifikasi

untuk rekomendasi topik skripsi. Kombinasi metode *feature selection*, SMOTE, dan algoritma SVM terbukti memberikan performa terbaik dibandingkan kombinasi lainnya. Hasil ini tidak hanya menunjukkan pkeunggulan Teknik oversampling dalam mengatasi data yang tidak seimbang, tetapi juga memberikan dasar yang kuat untuk mengembangkan system rekomendasi topik skripsi berbasis data mining. Penelitian ini dapat diperluas lebih lanjut dengan mempertimbangkan indikator lain yang mendukung kemampuan mahasiswa pada topik tertentu atau mengeksplorasi algoritma lain seperti metode ensemble untuk meningkatkan performa model. Selain itu, evaluasi terhadap data dari populasi mahasiswa yang lebih beragam dapat dilakukan untuk meningkatkan generalisasi model.

DAFTAR PUSTAKA

- [1] E. M. Sari Rochman *et al.*, “Classification of Thesis Topics Based on Informatics Science Using SVM,” *IOP Conf Ser Mater Sci Eng*, vol. 1125, no. 1, p. 012033, May 2021, doi: 10.1088/1757-899X/1125/1/012033.
- [2] F. Fitrah and A. Irianto, “Analisis Prokrastinasi Dalam Mengerjakan Skripsi Pada Mahasiswa Jurusan Pendidikan Ekonomi Angkatan 2015 Universitas Negeri Padang,” *Jurnal Ecogen*, vol. 2, no. 3, p. 412, Oct. 2019, doi: 10.24036/jmpe.v2i3.7412.
- [3] C. Slamet, F. M. Maliki, U. Syaripudin, A. S. Amin, and M. A. Ramdhani, “Thesis topic recommendation using simple multi attribute rating technique,” *J Phys Conf Ser*, vol. 1402, no. 6, p. 066105, Dec. 2019, doi: 10.1088/1742-6596/1402/6/066105.
- [4] H. Aldowah, H. Al-Samarraie, and W. M. Fauzy, “Educational data mining and learning analytics for 21st century higher education: A review and synthesis,” *Telematics and Informatics*, vol. 37, pp. 13–49, Apr. 2019, doi: 10.1016/j.tele.2019.01.007.
- [5] N. O. Adiwijaya and R. Sarno, “Specialty Coffees Classification Utilizes Feature Selection and Machine Learning,” 2024, pp. 94–101. doi: 10.2991/978-94-6463-445-7_11.
- [6] E. Haryatmi and S. Pramita Hervianti, “Penerapan Algoritma Support Vector Machine Untuk Model Prediksi Kelulusan Mahasiswa Tepat Waktu,” *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 5, no. 2, pp. 386–392, Apr. 2021, doi: 10.29207/resti.v5i2.3007.
- [7] L. B. Ilmawan and M. A. Mude, “Perbandingan Metode Klasifikasi Support Vector Machine dan Naïve Bayes untuk Analisis Sentimen pada Ulasan Tekstual di Google Play Store,” *ILKOM Jurnal Ilmiah*, vol. 12, no. 2, pp. 154–161, Aug. 2020, doi: 10.33096/ilkom.v12i2.597.154-161.
- [8] F. Farid, U. Enri, and Y. Umaidah, “Sistem Pendukung Keputusan Rekomendasi Topik Skripsi Menggunakan Naïve Bayes Classifier,” *JOINTECS (Journal of Information Technology and Computer Science)*, vol. 6, no. 1, p. 35, Jan. 2021, doi: 10.31328/jointecs.v6i1.2076.
- [9] Y.-Z. Li and F. Min, “Parallel–serial architecture with instance correlation label-specific features for multi-label learning,” *Knowl Based Syst*, vol. 304, p. 112568, Nov. 2024, doi: 10.1016/j.knosys.2024.112568.
- [10] G. Verhoeven, “Optimal performance with Random Forests: does feature selection beat tuning?,” github.io. Accessed: Dec. 26, 2024. [Online]. Available: https://gsverhoeven.github.io/post/random-forest-rfe_vs_tuning/
- [11] A. Pratiwi, N. Oktavia Adiwijaya, and Q. A’yuni Ar Ruhimat, “Pengukuran Efektivitas Endorsement Melalui Akun Kreator dan Akun Publik Instagram Menggunakan TEARS Model,” *JURNAL SOSIAL Jurnal Penelitian Ilmu-Ilmu Sosial*, vol. 25, no. 2, pp. 69–77, Nov. 2024, doi: 10.33319/sos.v25i2.211.
- [12] M. Stephanou and M. Varughese, “Sequential estimation of Spearman rank correlation using Hermite series estimators,” *J Multivar Anal*, vol. 186, p. 104783, Nov. 2021, doi: 10.1016/j.jmva.2021.104783.

-
- [13] S. Matharaarachchi, M. Domaratzki, and S. Muthukumarana, "Enhancing SMOTE for imbalanced data with abnormal minority instances," *Machine Learning with Applications*, vol. 18, p. 100597, Dec. 2024, doi: 10.1016/j.mlwa.2024.100597.
 - [14] P. Sun, Z. Wang, L. Jia, and Z. Xu, "SMOTE-kTLNN: A hybrid re-sampling method based on SMOTE and a two-layer nearest neighbor classifier," *Expert Syst Appl*, vol. 238, p. 121848, Mar. 2024, doi: 10.1016/j.eswa.2023.121848.
 - [15] D. Berrar, "Bayes' Theorem and Naive Bayes Classifier," in *Reference Module in Life Sciences*, Elsevier, 2024. doi: 10.1016/B978-0-323-95502-7.00118-4.
 - [16] H. R. Baghaee, D. Mlakic, S. Nikolovski, and T. Dragicevic, "Anti-Islanding Protection of PV-Based Microgrids Consisting of PHEVs Using SVMs," *IEEE Trans Smart Grid*, vol. 11, no. 1, pp. 483–500, Jan. 2020, doi: 10.1109/TSG.2019.2924290.
 - [17] M. N. Qureshi and N. A. Nor, "Optimisation and cross-validation of an e-hailing hybrid pricing algorithm using a supervised learning classification and regression tree model: a heuristic approach," *J King Saud Univ Sci*, vol. 36, no. 3, Mar. 2024, doi: 10.1016/j.jksus.2024.103107.
 - [18] H. S. Mansur, N. O. Adiwijaya, and T. Dharmawan, "Optimization of Machine Learning Algorithms with Bagging and AdaBoost Methods for Stroke Disease Prediction," 2023.

Halaman ini sengaja dikosongkan