

Perbandingan Algoritma Klasifikasi *Decision Tree*, *Random Forest*, dan *Naive Bayes* dalam Prediksi Kelulusan Mahasiswa

Sartika Lina Mulani Sitio ^{1*}, Susanna Dwi Yulianti Kusuma ², Hidayatullah Al Islami ³, Mukti Andrianto ⁴

^{1*,2,3,4} Program Studi Teknik Informatika, Fakultas Ilmu Komputer, Universitas Pamulang, Kota Tangerang Selatan, Provinsi Banten, Indonesia.

Email: dosen00847@unpam.ac.id ^{1*}, dosen00682@unpam.ac.id ², dosen02408@unpam.ac.id ³, muktiandrianto01@gmail.com ⁴

Histori Artikel:

Dikirim 10 Juni 2026; *Diterima dalam bentuk revisi* 15 Juni 2026; *Diterima* 30 Juli 2026; *Diterbitkan* 10 September 2026. Semua hak dilindungi oleh Lembaga Penelitian dan Pengabdian Masyarakat (LPPM) STMKI Indonesia Banda Aceh.

Abstrak

Kelulusan tepat waktu merupakan salah satu indikator penting dalam keberhasilan akademik perguruan tinggi. Namun, proses monitoring mahasiswa masih menghadapi kendala dalam mengidentifikasi mahasiswa yang berpotensi tidak lulus tepat waktu. Penelitian ini menawarkan solusi berupa penerapan machine learning untuk memprediksi status kelulusan sekaligus membandingkan algoritma yang paling efektif. Penelitian menggunakan 750 data mahasiswa Program Studi Informatika angkatan 2017 – 2021 yang bersumber dari Sistem Informasi Akademik. Metode penelitian meliputi preprocessing data, penanganan missing value dan outlier, encoding, normalisasi, pembentukan model Decision Tree, Random Forest, dan Naïve Bayes, serta evaluasi menggunakan accuracy, precision, recall, F1-score, dan AUC-ROC. Hasil penelitian menunjukkan bahwa Random Forest memberikan kinerja terbaik dengan accuracy 91,33%, precision 90,48%, recall 93,42%, F1-score 91,93%, dan AUC-ROC 0,952. Decision Tree memperoleh accuracy 84,67%, sedangkan Naïve Bayes sebesar 80,00%. Uji Friedman menghasilkan nilai $p = 0,015$ yang menunjukkan adanya perbedaan kinerja yang signifikan. Dengan demikian, Random Forest menjadi algoritma terbaik untuk prediksi kelulusan mahasiswa pada dataset penelitian ini dan berpotensi mendukung sistem peringatan dini akademik.

Kata Kunci: Prediksi Kelulusan Mahasiswa Machine Learning; Random Forest Decision Tree Student Graduation Prediction; Educational Data Mining Klasifikasi Akademik.

Abstract

Timely graduation is one of the important indicators of academic success of universities. However, the student monitoring process still faces obstacles in identifying students who have the potential to not graduate on time. This study offers a solution in the form of applying machine learning to predict graduation status while comparing the most effective algorithms. The research uses 750 data of students of the Informatics Study Program batch 2017 – 2021 sourced from the Academic Information System. The research methods include data preprocessing, handling missing values and outliers, encoding, normalization, forming Decision Tree, Random Forest, and Naïve Bayes models, as well as evaluation using accuracy, precision, recall, F1-score, and AUC-ROC. The results of the study show that Random Forest provides the best performance with an accuracy of 91.33%, precision of 90.48%, recall of 93.42%, F1-score of 91.93%, and an AUC-ROC of 0.952. Decision Tree obtained an accuracy of 84.67%, while Naïve Bayes was 80.00%. The Friedman test resulted in a value of $p = 0.015$ which showed a significant difference in performance. Thus, Random Forest is the best algorithm for predicting student graduation in this research dataset and has the potential to support academic early warning systems.

Keyword: Prediksi Kelulusan Mahasiswa Machine Learning; Random Forest Decision Tree Student Graduation Prediction; Educational Data Mining Klasifikasi Akademik.

1. Pendahuluan

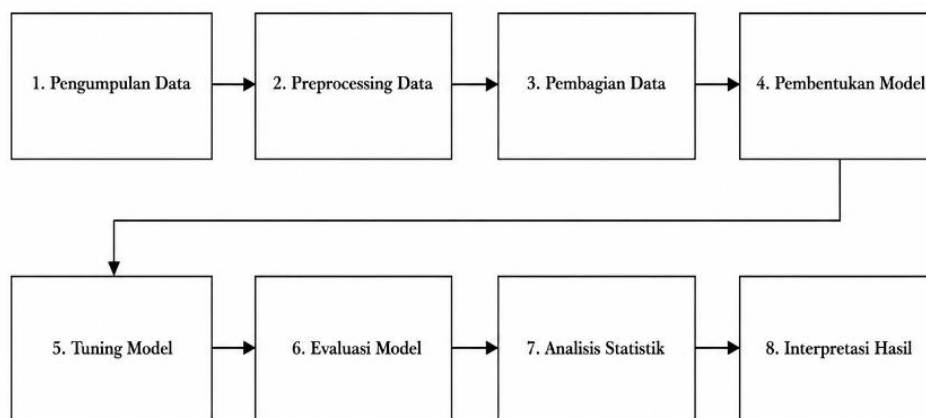
Kelulusan tepat waktu merupakan salah satu indikator penting dalam mengevaluasi keberhasilan mahasiswa dan efektivitas pengelolaan akademik perguruan tinggi. Keterlambatan penyelesaian studi dapat meningkatkan beban biaya bagi mahasiswa, memperpanjang masa studi, serta memengaruhi efisiensi pengelolaan sumber daya akademik. Oleh karena itu, kemampuan untuk mengidentifikasi sejak dini mahasiswa yang berpotensi tidak lulus tepat waktu menjadi penting bagi perguruan tinggi. Pemanfaatan data akademik melalui pendekatan *machine learning* dapat menjadi salah satu solusi karena mampu mempelajari pola dari data historis mahasiswa untuk menghasilkan prediksi yang dapat mendukung proses monitoring dan pengambilan keputusan akademik (Pelima *et al.*, 2024). Perkembangan *Educational Data Mining* dan *Learning Analytics* menunjukkan bahwa berbagai algoritma *machine learning* telah digunakan untuk memprediksi keberhasilan akademik, risiko *dropout*, dan kelulusan mahasiswa. Kajian sistematis Pelima *et al.* menunjukkan bahwa penelitian prediksi kelulusan mahasiswa semakin berkembang dan memanfaatkan data akademik maupun data dari sistem informasi pendidikan, dengan *Random Forest* dan *Support Vector Machine* termasuk algoritma yang banyak digunakan. Sementara itu, kajian sistematis mengenai prediksi *dropout* menunjukkan bahwa *Random Forest* merupakan salah satu algoritma yang paling sering digunakan dan memiliki kinerja prediksi yang kompetitif (Realinho *et al.*, 2022). Temuan tersebut menunjukkan bahwa pendekatan berbasis klasifikasi memiliki potensi besar untuk membantu perguruan tinggi dalam melakukan identifikasi mahasiswa berdasarkan kemungkinan keberhasilan studinya (Villar & de Andrade, 2024). Penelitian empiris terbaru juga memperlihatkan bahwa *Random Forest* memiliki kemampuan yang baik dalam prediksi kelulusan. Penelitian oleh (Pawitra *et al.*, 2024) menggunakan sembilan model *machine learning* pada 133 mahasiswa di perguruan tinggi Indonesia dan memperoleh hasil bahwa *Random Forest* memberikan performa terbaik dengan akurasi 85% dan AUC 0,875. Penelitian tersebut juga menunjukkan bahwa IPK kumulatif dan beberapa karakteristik akademik merupakan prediktor penting dalam menentukan kelulusan tepat waktu. Penelitian lain menggunakan *Random Forest* untuk memprediksi kelulusan pada program studi teknik komputer di Chile dan menemukan bahwa IPK menjadi salah satu faktor paling berpengaruh terhadap penyelesaian proses kelulusan (Quelopana *et al.*, 2024). Studi lain juga mulai mengembangkan pendekatan yang lebih kompleks, seperti *stacking*, SMOTE, dan optimasi PSO untuk meningkatkan performa prediksi kelulusan (Efendi *et al.*, 2025).

Meskipun berbagai penelitian tersebut telah menunjukkan efektivitas *machine learning*, masih terdapat kesenjangan penelitian. Sebagian penelitian berfokus pada pengembangan satu algoritma atau penggunaan model ensemble dan optimasi, sedangkan kajian empiris yang secara langsung membandingkan algoritma klasifikasi sederhana dengan karakteristik pembelajaran yang berbeda masih diperlukan untuk menentukan model yang paling sesuai pada dataset akademik tertentu. Selain itu, perbedaan karakteristik dataset, jumlah mahasiswa, distribusi kelas, atribut akademik, dan atribut sosial-ekonomi dapat menghasilkan performa algoritma yang berbeda (Johora *et al.*, 2025; Rahman *et al.*, 2025). Kajian sistematis terbaru juga menunjukkan bahwa sebagian besar penelitian prediksi kelulusan masih menggunakan dataset dari satu institusi dan memiliki keterbatasan dalam validasi eksternal serta generalisasi model. Hasil suatu algoritma pada penelitian sebelumnya tidak dapat secara langsung dianggap sebagai algoritma terbaik untuk seluruh konteks perguruan tinggi (Yáñez *et al.*, 2026). Berdasarkan kesenjangan tersebut, penelitian ini memfokuskan perbandingan tiga algoritma klasifikasi, yaitu *Decision Tree*, *Random Forest*, dan *Naïve Bayes*, menggunakan dataset akademik mahasiswa yang memiliki karakteristik akademik dan sosial-ekonomi. Dataset penelitian terdiri atas 750 mahasiswa Program Studi Informatika angkatan 2017–2021, dengan 513 mahasiswa atau 68,4% termasuk kategori lulus tepat waktu dan 237 mahasiswa atau 31,6% termasuk kategori tidak lulus tepat waktu. Variabel yang digunakan mencakup IPK semester, total SKS, jumlah mata kuliah gagal, status beasiswa, asal daerah, status pekerjaan, kehadiran, dan penghasilan orang tua. Pendekatan ini memungkinkan evaluasi yang lebih terukur terhadap kemampuan masing-masing algoritma dalam menangkap pola kelulusan mahasiswa. Kebaruan penelitian ini terletak pada evaluasi komparatif tiga algoritma klasifikasi dengan karakteristik berbeda pada dataset akademik mahasiswa yang

menggabungkan atribut akademik dan sosial-ekonomi, disertai evaluasi multidimensi dan pengujian signifikansi statistik. Berbeda dengan penelitian yang hanya berfokus pada pengembangan satu model atau peningkatan performa melalui optimasi, penelitian ini tidak hanya menentukan algoritma dengan performa prediksi tertinggi, tetapi juga menganalisis faktor yang berkontribusi terhadap prediksi kelulusan. Evaluasi dilakukan menggunakan *accuracy*, *precision*, *recall*, F1-score, dan AUC-ROC serta diperkuat dengan uji Friedman untuk mengetahui apakah perbedaan performa antaralgoritma signifikan secara statistik. Berdasarkan uraian tersebut, penelitian ini bertujuan untuk menerapkan dan membandingkan algoritma *Decision Tree*, *Random Forest*, dan *Naïve Bayes* dalam memprediksi kelulusan mahasiswa, mengevaluasi kinerja ketiga algoritma berdasarkan *accuracy*, *precision*, *recall*, F1-score, dan AUC-ROC, serta menentukan algoritma dengan performa terbaik dan faktor akademik yang paling berpengaruh terhadap prediksi kelulusan. Tujuan tersebut secara langsung menjawab permasalahan penelitian mengenai penerapan, perbandingan kinerja, dan pemilihan algoritma terbaik untuk prediksi kelulusan mahasiswa.

2. Metode Penelitian

Penelitian ini menggunakan pendekatan kuantitatif komparatif dengan desain eksperimental komputasional untuk membandingkan kinerja algoritma *Decision Tree*, *Random Forest*, dan *Naïve Bayes* dalam memprediksi kelulusan mahasiswa. Seluruh tahapan penelitian dilakukan secara sistematis, mulai dari pengumpulan dan persiapan data, pembentukan model, penyetelan parameter, evaluasi kinerja, hingga analisis statistik dan interpretasi hasil. Alur tahapan penelitian ditunjukkan pada Gambar 1.



Gambar 1. Model Proses Penelitian

Adapun penjelasan dari tahapan metode diatas adalah sebagai berikut:

1) Pengumpulan Data

Data penelitian merupakan data sekunder yang diperoleh dari Sistem Informasi Akademik (SIKAD) melalui koordinasi dengan Biro Administrasi Akademik dan Kemahasiswaan. Dataset terdiri atas 750 mahasiswa Program Studi Informatika angkatan 2017–2021. Sebanyak 513 mahasiswa (68,4%) dikategorikan lulus tepat waktu dan 237 mahasiswa (31,6%) dikategorikan tidak lulus tepat waktu. Data yang digunakan telah dianonimkan untuk menjaga privasi mahasiswa.

2) Preprocessing Data

Tahap pra-pemrosesan dilakukan setelah data dibagi menjadi data latih dan data uji untuk mencegah terjadinya data leakage. Dataset dibagi secara stratified menjadi 80% data latih dan 20% data uji dengan mempertahankan proporsi kelas pada kedua subset. Seluruh parameter transformasi pada tahap preprocessing dihitung berdasarkan data latih, kemudian diterapkan pada

data uji tanpa melakukan perhitungan ulang menggunakan data uji. Penanganan missing value dilakukan pada data latih dengan imputasi median untuk atribut numerik dan modus untuk atribut kategorikal (Cohausz, 2022). Nilai parameter imputasi yang diperoleh dari data latih kemudian diterapkan pada data uji. Deteksi outlier dilakukan menggunakan metode Interquartile Range (IQR), sedangkan penanganan outlier dilakukan dengan teknik capping berdasarkan batas bawah dan batas atas yang dihitung dari data latih. Atribut kategorikal selanjutnya dikonversi menggunakan One-Hot Encoding, sedangkan variabel target dikodekan secara biner, yaitu 1 untuk Lulus Tepat Waktu dan 0 untuk Tidak Lulus Tepat Waktu. Untuk atribut numerik digunakan Min-Max Scaling. Parameter transformasi scaling dihitung hanya dari data latih dan digunakan kembali untuk mentransformasi data uji (Atika, 2026; Silva *et al.*, 2025)

3) Pembagian Data

Dataset dibagi menjadi 80% data latih dan 20% data uji. Pembagian dilakukan secara *stratified* untuk mempertahankan proporsi masing-masing kelas pada data latih dan data uji. Dengan 750 data, pembagian tersebut menghasilkan sekitar 600 data untuk pelatihan dan 150 data untuk pengujian. Penggunaan *stratified split* sesuai dengan fungsi *train_test_split* yang memungkinkan distribusi kelas dipertahankan melalui parameter (Joseph, 2022; Melo *et al.*, 2022).

4) Pembentukan Model

Tiga algoritma klasifikasi dibangun menggunakan data latih, yaitu Decision Tree, Random Forest, dan Naïve Bayes. Decision Tree digunakan sebagai model berbasis pohon yang mudah diinterpretasikan, Random Forest menggunakan kombinasi beberapa pohon keputusan, sedangkan Naïve Bayes menggunakan pendekatan probabilistik. Untuk data numerik seperti IPK dan kehadiran, penelitian menggunakan varian Gaussian Naïve Bayes (Sitio & Nadiyanti, 2022)

5) Tuning Model

Optimasi parameter dilakukan menggunakan Grid Search Cross Validation untuk memperoleh konfigurasi model yang memberikan kinerja terbaik. Penelitian menggunakan *5-fold cross-validation*, dengan pemilihan konfigurasi berdasarkan nilai F1-score tertinggi pada data validasi (Bennemann *et al.*, 2022; Bischl *et al.*, 2023).

6) Evaluasi Model

Model terbaik dari masing-masing algoritma kemudian dievaluasi menggunakan data uji independen. Pengukuran kinerja dilakukan menggunakan accuracy, precision, recall, F1-score, dan AUC-ROC. Confusion matrix juga digunakan untuk melihat kesalahan klasifikasi antara mahasiswa yang diprediksi lulus dan tidak lulus. Metrik tersebut merupakan metrik klasifikasi yang tersedia dalam *scikit-learn*, termasuk accuracy, precision, recall, F1, dan ROC-AUC (Gorshkov *et al.*, 2025; Sitio *et al.*, 2024).

7) Analisis Statistik

Uji Friedman dilakukan menggunakan nilai accuracy yang diperoleh dari masing-masing lima fold pada Stratified 5-Fold Cross Validation. Lima fold diperlakukan sebagai lima pengukuran berulang, sedangkan accuracy digunakan sebagai satu-satunya metrik yang diranking pada setiap fold. Dengan demikian, perbandingan antaralgoritma tidak menggunakan kombinasi beberapa metrik, tetapi berdasarkan accuracy pada lima fold yang sama (Demšar, 2006).

8) Interpretasi Hasil

Tahap terakhir adalah menginterpretasikan hasil evaluasi untuk menentukan algoritma yang paling efektif dalam memprediksi kelulusan mahasiswa. Selain membandingkan performa ketiga algoritma, penelitian juga menganalisis *feature importance* untuk mengidentifikasi variabel yang paling berkontribusi terhadap prediksi. Hasil akhirnya digunakan untuk memberikan rekomendasi algoritma yang dapat mendukung pengembangan sistem prediksi atau peringatan dini akademik.

3. Hasil dan Pembahasan

3.1 Hasil

Bagian ini menyajikan hasil penerapan algoritma *Decision Tree*, *Random Forest*, dan *Naive Bayes* dalam memprediksi kelulusan mahasiswa. Hasil penelitian dianalisis berdasarkan tahapan pengolahan data, evaluasi kinerja model, *cross-validation*, *confusion matrix*, analisis *feature importance*, *learning curve*, serta uji signifikansi statistik. Selanjutnya, hasil dari masing-masing algoritma dibandingkan untuk menentukan model yang memiliki kinerja paling optimal dalam prediksi kelulusan mahasiswa.

3.1.1 Distribusi Data dan Pembagian Data

Dataset terdiri dari 750 mahasiswa, dengan 513 mahasiswa atau 68,4% berstatus Lulus Tepat Waktu dan 237 mahasiswa atau 31,6% berstatus Tidak Lulus Tepat Waktu. Data kemudian dibagi menjadi data pelatihan dan data pengujian secara stratifikasi. Data uji independen terdiri atas 150 mahasiswa, yaitu 103 mahasiswa Lulus dan 47 mahasiswa Tidak Lulus.

Tabel 1. Distribusi dan Pembagian Data

Status Kelulusan	Jumlah	Persentase
Lulus Tepat Waktu	513	68,4%
Tidak Lulus Tepat Waktu	237	31,6%
Total	750	100%

3.1.2 Hasil 5-Fold Cross Validation

Evaluasi awal dilakukan menggunakan 5-Fold Stratified Cross Validation. Setiap algoritma diuji pada lima fold sehingga dapat diketahui kestabilan performanya terhadap variasi data pelatihan dan validasi.

Tabel 2. Hasil 5-Fold Cross Validation

Fold	Decision Tree	Random Forest	Naive Bayes
Fold 1	83,33%	90,00%	78,33%
Fold 2	85,00%	91,67%	81,67%
Fold 3	84,17%	92,50%	79,17%
Fold 4	86,67%	90,83%	82,50%
Fold 5	84,17%	91,67%	79,17%
Rata - rata	84,67%	91,33%	80,17%
Standar Deviasi	$\pm 1,18\%$	$\pm 0,88\%$	$\pm 1,52\%$

3.1.3 Hasil Confusion Matrix

Pengujian menggunakan data uji independen dan menghasilkan confusion matrix untuk masing – masing algoritma:

Tabel 3. Hasil Confusion Matrix

Algoritma	TP	TN	FP	FN
Decision Tree	90	37	9	14
Random Forest	96	41	5	8
Naive Bayes	88	32	11	19

Random Forest menghasilkan jumlah *True Positive* (96) dan *True Negative* (41) paling tinggi. Pada saat yang sama, jumlah kesalahan klasifikasi juga paling rendah, yaitu hanya 5 *False Positive* dan 8 *False Negative*. *Decision Tree* menghasilkan 9 *False Positive* dan 14 *False Negative*, sedangkan *Naive Bayes* menghasilkan 11 *False Positive* dan 19 *False Negative*.

Temuan tersebut memperlihatkan bahwa *Random Forest* lebih mampu mengklasifikasikan mahasiswa ke dalam kelas yang tepat. Penurunan *False Negative* menjadi sangat penting dalam konteks prediksi akademik karena kesalahan tersebut berkaitan dengan mahasiswa yang tidak teridentifikasi secara tepat oleh sistem.

3.1.4 Perbandingan Matrix Evaluasi

Untuk memperoleh perbandingan yang lebih komprehensif, ketiga algoritma dievaluasi menggunakan *accuracy*, *precision*, recall, F1-score, dan AUC-ROC.

Tabel 4. Perbandingan Matrix Evaluasi

Matrix	Decision Tree	Random Forest	Naïve Bayes
Accuracy	84,67%	91,33%	80,00%
Precision	83,21%	90,48%	78,95%
Recall	87,10%	93,42%	85,71%
F1-Score	85,12%	91,93%	82,19%
AUC – ROC	0,881%	0,952	0,847
Waktu Pelatihan	0,23 detik	2,87 detik	0,18 detik

Berdasarkan tabel diatas, *Random Forest* menghasilkan performa terbaik pada seluruh metrik kualitas prediksi. *Accuracy* sebesar 91,33% menunjukkan bahwa model mampu melakukan klasifikasi dengan tingkat ketepatan yang tinggi. *Precision* sebesar 90,48% menunjukkan bahwa sebagian besar prediksi positif yang dihasilkan model merupakan prediksi yang benar. Sementara itu, recall sebesar 93,42% menunjukkan kemampuan model dalam mengenali kelas target dengan baik. F1-score sebesar 91,93% menunjukkan keseimbangan yang baik antara *precision* dan recall. Nilai AUC-ROC sebesar 0,952 juga menunjukkan bahwa *Random Forest* mempunyai kemampuan diskriminasi yang sangat baik dalam membedakan mahasiswa berdasarkan status kelulusan. Nilai tersebut lebih tinggi dibandingkan *Decision Tree* sebesar 0,881 dan *Naïve Bayes* sebesar 0,847. Namun, *Random Forest* membutuhkan waktu komputasi lebih tinggi. Waktu pelatihannya mencapai 2,87 detik, sedangkan *Decision Tree* hanya 0,23 detik dan *Naïve Bayes* 0,18 detik. Terdapat trade-off antara akurasi dan efisiensi komputasi. Dalam penelitian ini, peningkatan kualitas prediksi *Random Forest* lebih dominan dibandingkan tambahan waktu komputasinya.

Tabel 5. Analisis Learning Curve

Algoritma	Training Score	CV Score	Interpretasi
Decision Tree	≈95%	≈84 – 87%	High variance/overfitting
Random Forest	≈99%	≈91,3%	Generalisasi baik
Naïve Bayes	≈80 – 82%	≈80 – 82%	High bias

Learning curve Decision Tree menunjukkan adanya kesenjangan antara training score dan validation score. Kondisi tersebut mengindikasikan kecenderungan *high variance* atau *overfitting*. *Random Forest* menunjukkan pola yang lebih sehat karena validation score meningkat hingga sekitar 91,3% dan relatif stabil, sementara kesenjangan performanya masih terkendali. Sebaliknya, *Naïve Bayes* menunjukkan training dan validation score yang relatif rendah dan stagnan. Pola tersebut menunjukkan *high bias*, sehingga penambahan data diperkirakan tidak akan memberikan peningkatan performa yang besar tanpa perubahan pendekatan atau representasi fitur.

3.2 Pembahasan

Hasil penelitian ini menunjukkan bahwa algoritma *Random Forest* memiliki performa terbaik dalam memprediksi kelulusan mahasiswa, dengan nilai akurasi sebesar 91,33%, F1-score 91,93%, dan AUC-ROC 0,952. Keunggulan ini sejalan dengan temuan dari Pawitra *et al.* (2024), yang juga melaporkan bahwa *Random Forest* menunjukkan performa terbaik dibandingkan model lain dalam prediksi

kelulusan mahasiswa di perguruan tinggi Indonesia. Hal ini dapat dikarenakan kemampuan Random Forest dalam mengatasi data dengan fitur yang kompleks dan variabel yang saling berinteraksi, serta kemampuannya dalam mengurangi overfitting melalui proses ensemble learning. Di sisi lain, *Decision Tree* menunjukkan performa yang cukup baik namun cenderung mengalami overfitting, sebagaimana yang juga diungkapkan oleh Efendi *et al.* (2025), yang menyatakan bahwa model berbasis pohon seperti *Decision Tree* cenderung memiliki variansi yang tinggi jika tidak dilakukan pruning atau teknik regularisasi lainnya. Sementara itu, Naïve Bayes menunjukkan performa yang paling rendah di antara ketiga algoritma, dengan akurasi sebesar 80%, yang konsisten dengan studi sebelumnya oleh Sitio dan Nadiyanti (2022), yang menyebutkan bahwa Naïve Bayes mungkin kurang optimal untuk dataset yang memiliki fitur saling bergantung dan bersifat kompleks. Hasil uji Friedman menunjukkan adanya perbedaan signifikan secara statistik antara ketiga algoritma ini ($p = 0,015$), menegaskan bahwa pilihan algoritma sangat mempengaruhi hasil prediksi. Secara umum, temuan ini memperkuat hasil dari Villar dan de Andrade (2024), yang menyatakan bahwa algoritma ensemble seperti Random Forest cenderung memiliki kinerja prediksi yang lebih stabil dan akurat dalam konteks prediksi akademik dibandingkan algoritma sederhana seperti Naïve Bayes. Oleh karena itu, penggunaan Random Forest sebagai model prediksi kelulusan mahasiswa sangat relevan dan memiliki potensi besar untuk diterapkan dalam sistem peringatan dini akademik, meskipun harus diimbangi dengan pertimbangan efisiensi waktu komputasi yang lebih tinggi.

4. Kesimpulan

Berdasarkan hasil penelitian mengenai perbandingan algoritma *Decision Tree*, *Random Forest*, dan *Naïve Bayes* dalam prediksi kelulusan mahasiswa, dapat disimpulkan bahwa ketiga algoritma mampu melakukan klasifikasi status kelulusan berdasarkan data akademik dan faktor pendukung lainnya. Random Forest menunjukkan performa terbaik, dengan accuracy sebesar 91,33%, F1-score sebesar 91,93%, dan AUC-ROC sebesar 0,952. Analisis *feature importance* menunjukkan bahwa IPK rata-rata semester 1–6, persentase kehadiran, jumlah mata kuliah gagal, dan total SKS lulus merupakan fitur yang paling dominan pada model Random Forest. Keempat fitur tersebut secara kumulatif memberikan kontribusi sebesar 86,16% terhadap keseluruhan feature importance. Hasil ini menunjukkan bahwa karakteristik akademik merupakan faktor yang lebih dominan dalam prediksi kelulusan tepat waktu dibandingkan faktor sosial-ekonomi yang digunakan dalam penelitian. Berdasarkan performanya, *Random Forest* berpotensi digunakan sebagai model prediktif dalam pengembangan sistem peringatan dini akademik untuk membantu proses monitoring mahasiswa. Namun, penelitian ini belum mencakup implementasi sistem peringatan dini maupun pengujian penggunaan model oleh pengelola akademik. Oleh karena itu, penelitian selanjutnya dapat diarahkan pada integrasi model *Random Forest* ke dalam sistem informasi akademik, pengembangan dashboard monitoring, serta pengujian efektivitas sistem dalam mendukung intervensi akademik.

5. Daftar Pustaka

- Atika, P. D. (2026). A comparative study of machine learning-based student dropout risk prediction. *Piksel Journal*, 14(225), 167–174. <https://doi.org/10.33558/piksel.v14i1.12299>.
- Bennemann, B., Schwartz, B., Giesemann, J., & Lutz, W. (2022). Predicting patients who will drop out of out-patient psychotherapy using machine learning algorithms. *British Journal of Psychiatry*, 220(4), 192–201. <https://doi.org/10.1192/bjp.2022.17>.
- Bischl, B., Binder, M., Lang, M., Pielok, T., Richter, J., Coors, S., Thomas, J., Ullmann, T., Becker, M., Boulesteix, A. L., Deng, D., & Lindauer, M. (2023). Hyperparameter optimization:

Foundations, algorithms, best practices, and open challenges. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 13(2), 1–43. <https://doi.org/10.1002/widm.1484>.

Cohausz, L. (2022). When probabilities are not enough - A framework for causal explanations of student success models. *Journal of Educational Data Mining*, 14(3), 52–75. <https://doi.org/10.5281/zenodo.7304800>.

Demšar, J. (2006). Statistical comparisons of classifiers over multiple data sets. *Journal of Machine Learning Research*, 7, 1–30.

Efendi, A., Fitri, I., & Nurcahyo, G. W. (2025). Improving student graduation timeliness prediction using SMOTE and ensemble learning with stacking and GridSearchCV optimization. *Data and Metadata*, 4. <https://doi.org/10.56294/dm2025917>.

Gorshkov, S. S., Ignatov, D. I., & Chernysheva, A. Y. (2025). Predicting student dropout through text and media content analysis of VKontakte profiles. *IEEE Access*, 13(January), 46732–46746. <https://doi.org/10.1109/ACCESS.2025.3550817>.

Johora, F. T., Hasan, M. N., Rajbongshi, A., Ashrafuzzaman, M., & Akter, F. (2025). An explainable AI-based approach for predicting undergraduate students' academic performance. *Array*, 26(February). <https://doi.org/10.1016/j.array.2025.100384>.

Joseph, V. R. (2022). Optimal ratio for data splitting. *Statistical Analysis and Data Mining*, 15(4), 531–538. <https://doi.org/10.1002/sam.11583>.

Melo, E., Silva, I., Costa, D. G., Viegas, C. M. D., & Barros, T. M. (2022). On the use of explainable artificial intelligence to evaluate school dropout. *Education Sciences*, 12(12). <https://doi.org/10.3390/educsci12120845>.

Pawitra, M. A. S., Hung, H.-C., & Jati, H. (2024). A machine learning approach to predicting on-time graduation in Indonesian higher education. *Elinvo (Electronics, Informatics, and Vocational Education)*, 9(2), 294–308. <https://doi.org/10.21831/elinvo.v9i2.77052>.

Pelima, L. R., Sukmana, Y., & Rosmansyah, Y. (2024). Predicting university student graduation using academic performance and machine learning: A systematic literature review. *IEEE Access*, 12, 23451–23465. <https://doi.org/10.1109/ACCESS.2024.3361479>.

Quelopana, A., Keith, B., & Pizarro, R. (2024). Predictive modeling of on-time graduation in computing engineering programs: A case study from Northern Chile. *Computer Applications in Engineering Education*, 32(5), e22767. <https://doi.org/10.1002/cae.22767>.

Rahman, A., Mahdiana, D., & Fauzi, A. (2025). Predicting student on-time graduation using particle swarm optimization and random forest algorithms. *Indonesian Journal of Artificial Intelligence and Data Mining*, 8(1), 161. <https://doi.org/10.24014/ijaidm.v8i1.33577>.

Realinho, V., Machado, J., Baptista, L., & Martins, M. V. (2022). Predicting student dropout and academic success. *Data*, 7(11). <https://doi.org/10.3390/data7110146>.

Silva, F. da C., Santana, A. M., & Feitosa, R. M. (2025). An investigation into dropout indicators in secondary technical education using explainable artificial intelligence. *Revista Iberoamericana de Tecnologías del Aprendizaje*, 20, 105–114. <https://doi.org/10.1109/RITA.2025.3566095>.

- Sitio, S. L. M., & Nadiyanti, R. (2022). Analisis sentimen kenaikan harga BBM Pertamina pada media sosial menggunakan metode Naïve Bayes classifier. *Building of Informatics, Technology and Science (BITS)*, 4(3), 1224–1231. <https://doi.org/10.47065/bits.v4i3.2311>.
- Sitio, S. L. M., Thoyyibah, Maryani, Ardiansyah, M., Rosmawarni, N., Djaksana, Y. M., & Arianti, N. D. (2024). Comparison of the ensemble XGBoost and transformer models with machine learning for classification of Indonesian music moods of the 70's and 80's era. *Journal of Theoretical and Applied Information Technology*, 102(24), 9157–9165.
- Villar, A., & de Andrade, C. R. V. (2024). Supervised machine learning algorithms for predicting student dropout and academic success: A comparative study. *Discover Artificial Intelligence*, 4(1). <https://doi.org/10.1007/s44163-023-00079-z>.
- Yáñez, A., Crawford, B., Monfroy, E., Paz, Á., Barrera-García, J., Cisternas-Caneo, F., Peña Fritz, Á., & Soto, R. (2026). Machine learning for graduation prediction in higher education: A systematic review with a bio-inspired optimization perspective. *Biomimetics*, 11(7), 1–40. <https://doi.org/10.3390/biomimetics11070512>.