



PERBANDINGAN ALGORITMA K-NEAREST NEIGHBOUR DAN NAÏVE BAYES UNTUK MENDETEKSI PENIPUAN KARTU KREDIT

Fauzan Firdaus¹⁾, Ahmad Homaidi²⁾, Jarot Dwi Prasetyo³⁾, Hermanto⁴⁾, Ach. Zubairi⁵⁾, Lukman Fakhid Lidimilah⁶⁾

^{1,2,6} Teknologi Informasi, Universitas Ibrahimy

^{3,4} Ilmu Komputer, Universitas Ibrahimy

⁵ Sistem Informasi, Universitas Ibrahimy

email: ¹fauzanfirdaus@ibrahimiy.ac.id, ²ahmadhomaidi@ibrahimiy.ac.id, ³jarot_dwi_prasetyo@yahoo.com, ⁴otnamreh@hotmail.com, ⁵06zuber90@gmail.com, ⁶lucky.lukman7@gmail.com

ARTICLE INFO

Article History:

Received : 27 November 2025

Accepted : 7 Desember 2025

Published : 28 Desember 2025

Keywords:

K-Nearest Neighbor

Naïve Bayes

Penipuan

Kartu Kredit

IEEE style in citing this article:

F. Firdaus, A. Homaidi, J.D. Prasetyo, H. Hermanto, A. Zubairi, L.F. Limilah, "Perbandingan Algoritma K-Nearest Neighbour Dan Naïve Bayes Untuk Mendeteksi Penipuan Kartu Kredit", *Jurnal.ilmiah.informatika*, vol. 10, no. 2, pp. 163-171, Des. 2025.

ABSTRACT

Credit card fraud is a serious problem in the financial industry that continues to increase with the development of digital transaction technology. This study aims to compare the performance of the K-Nearest Neighbour (KNN) and Naive Bayes algorithms in detecting credit card fraud by considering various evaluation metrics evaluation metrics, including not only accuracy but also precision, recall, and F1-score. The dataset used was sourced from Kaggle, comprising a total of 10,000 transaction records, which included financial transaction attributes and user behaviour. The research process included data pre-processing, attribute selection, data normalisation, and the application of both algorithms using RapidMiner software. The test results showed that the KNN algorithm produced an accuracy of 98.43%, a precision of 98.53%, and a recall of 99.90%, while Naive Bayes obtained an accuracy 98.20% accuracy, 99.69% precision, and 98.48% recall. Although KNN showed slightly superior performance in detecting fraudulent transactions, the T-Test statistical test showed that the difference in performance between the two algorithms was not statistically significant. KNN has an advantage in recognising complex patterns, but requires greater computational time, while Naive Bayes is more efficient in terms of speed. This study concludes that the selection of a fraud detection algorithm needs to consider the trade-off between accuracy and computational efficiency according to system requirements.

Corresponding Author:

Ahmad Homaidi

Universitas Ibrahimy

1. PENDAHULUAN

Penipuan kartu kredit merupakan salah satu masalah serius yang dihadapi oleh industri keuangan global. Menurut laporan dari Nilson Report, kerugian akibat penipuan kartu kredit diperkirakan mencapai lebih dari puluhan miliar di seluruh dunia pada tahun 2023, dan angka ini terus meningkat seiring dengan berkembangnya teknologi dan metode penipuan yang semakin canggih [1]. Dalam konteks ini, penggunaan algoritma machine learning untuk mendeteksi penipuan menjadi semakin penting. Dua algoritma yang sering dibandingkan dalam literatur adalah K-Nearest Neighbor (KNN) dan Naive Bayes, yang masing-masing memiliki kelebihan dan kekurangan dalam hal akurasi dan efisiensi.

K-Nearest Neighbor (KNN) adalah algoritma berbasis instance yang sederhana namun efektif, yang bekerja dengan cara mencari 'tetangga terdekat' dari data yang ingin diklasifikasikan. Algoritma ini mengandalkan jarak untuk menentukan klasifikasi, sehingga sangat sensitif terhadap distribusi data dan pemilihan fitur [2], [3], [4]. Di sisi lain, Naive Bayes adalah algoritma probabilistik yang mengasumsikan bahwa fitur-fitur yang ada bersifat independen satu sama lain. Meskipun asumsi ini seringkali tidak berlaku dalam praktik, Naive Bayes terbukti efektif dalam banyak aplikasi, termasuk deteksi penipuan kartu kredit [5], [6], [7].

Dalam penelitian yang telah dilakukan ada mengembangkan strategi baru yang menggabungkan KNN dengan Linear Discriminant Analysis (LDA) untuk meningkatkan tingkat recall dalam mendeteksi penipuan. Hasil penelitian menunjukkan bahwa kombinasi ini dapat mengurangi jumlah false negatives, yang merupakan salah satu masalah utama

dalam deteksi penipuan. Penelitian ini memberikan wawasan baru tentang bagaimana algoritma dapat berkolaborasi untuk mencapai hasil yang lebih baik [8], [9].

Namun, meskipun banyak penelitian telah dilakukan, masih ada gap yang signifikan dalam pemahaman tentang bagaimana kedua algoritma ini dapat dioptimalkan dalam konteks dataset yang tidak seimbang. Sebagian besar penelitian sebelumnya, seperti yang dilakukan oleh Al Khadhori et al., berfokus pada akurasi keseluruhan tanpa mempertimbangkan metrik lain seperti precision dan recall, yang sangat penting dalam konteks deteksi penipuan [10], [11], [12], [13]. Oleh karena itu, penelitian ini bertujuan untuk membandingkan KNN dan Naive Bayes secara lebih menyeluruh, dengan mempertimbangkan berbagai metrik evaluasi yang relevan.

Dalam konteks ini, penelitian ini akan mengkaji kinerja kedua algoritma dalam mendeteksi penipuan kartu kredit menggunakan dataset yang seimbang dan tidak seimbang, serta memberikan analisis mendalam tentang kelebihan dan kekurangan masing-masing algoritma. Dengan demikian, diharapkan penelitian ini dapat memberikan kontribusi signifikan terhadap pengembangan metode deteksi penipuan yang lebih efektif dan efisien di masa depan.

2. METODE PENELITIAN

2.1 Pengumpulan Data

Proses awal dalam penelitian ini dimulai dengan pembacaan data yang berfokus pada dataset transaksi kartu kredit. Dataset yang digunakan dalam penelitian ini adalah dataset yang tersedia secara publik yang didapatkan dari Kaggle yang berisi informasi transaksi kartu kredit yang mencurigakan. Dataset tersebut terdiri dari 10.000 data [14]. Data ini mencakup berbagai atribut, seperti id

transaksi, jumlah, waktu transaksi, kategori, transaksi luar negeri, ketidakcocokan lokasi, skor kepercayaan, kecepatan, dan usia pemegang kartu.

Dalam tahap ini, penting untuk memahami bahwa pembacaan data tidak hanya melibatkan pengambilan data mentah, tetapi juga mencakup pemahaman struktur dan karakteristik data. Statistik deskriptif, seperti rata-rata, median, dan deviasi standar dari jumlah transaksi, dapat memberikan wawasan awal tentang pola transaksi normal dan mencurigakan. Misalnya, rata-rata jumlah transaksi untuk kategori non-penipuan mungkin jauh lebih tinggi dibandingkan dengan transaksi yang teridentifikasi sebagai penipuan [15]. Oleh karena itu, analisis awal ini sangat krusial untuk langkah-langkah selanjutnya dalam pemilihan atribut dan penerapan algoritma.

2.2 Pemilihan Atribut

Setelah tahap pembacaan data, langkah berikutnya adalah pemilihan atribut yang relevan untuk model. Dalam konteks deteksi penipuan kartu kredit, tidak semua atribut dalam dataset memiliki kontribusi yang sama terhadap hasil akhir. Oleh karena itu, pemilihan atribut yang tepat dapat meningkatkan akurasi model secara signifikan. Metode seperti analisis korelasi dan teknik pemilihan fitur berbasis algoritma, seperti Recursive Feature Elimination (RFE), dapat digunakan untuk mengidentifikasi atribut yang paling berpengaruh [16].

Dalam penelitian ini, atribut yang dipilih untuk digunakan Adalah seleuruh atribut yang didapatkan dari kaggle kecuali id transaksi.

2.3 Mapping data

Setelah pemilihan atribut, langkah selanjutnya adalah pemetaan data ke dalam format yang sesuai untuk

penerapan algoritma K-Nearest Neighbor (KNN) dan Naive Bayes. Pemetaan ini mencakup normalisasi data, yang penting untuk memastikan bahwa semua atribut memiliki skala yang sama. KNN, sebagai algoritma berbasis jarak, sangat sensitif terhadap skala data, sehingga normalisasi menjadi langkah krusial untuk meningkatkan akurasi model [17], [18]. Metode normalisasi yang umum digunakan adalah Min-Max Scaling, di mana nilai-nilai atribut diubah ke dalam rentang [0, 1].

Di sisi lain, Naive Bayes tidak terlalu sensitif terhadap skala data, tetapi pemetaan yang baik tetap diperlukan untuk memastikan bahwa model dapat bekerja dengan optimal. Pada tahap ini, data dibagi menjadi dua subset: satu untuk pelatihan dan satu untuk pengujian. Pembagian ini biasanya dilakukan dengan rasio 90:10 atau 80:20, di mana sebagian besar data digunakan untuk melatih model, dan sisanya digunakan untuk menguji performa model [19]. Dengan pemetaan yang tepat, kedua algoritma dapat diterapkan dengan lebih efisien, dan hasil yang diperoleh dapat lebih representatif.

2.4 Penerapan Algoritma

Setelah semua persiapan selesai, langkah berikutnya adalah menerapkan algoritma K-Nearest Neighbor dan Naive Bayes. KNN bekerja dengan cara mencari k tetangga terdekat dari titik data yang ingin diklasifikasikan berdasarkan jarak Euclidean.

$$ED = \sqrt{(X_1 + y_1)^2 + \dots + (X_n + y_n)^2} \quad (1)$$

Setelah mencari nilai terdekat, selanjutnya adalah penentuan nilai k terdekat. Proses ini dapat dilakukan dengan melakukan voting dari tetangga terdekat berdasarkan nilai k. Penggunaan validation juga diterapkan untuk

memastikan bahwa model tidak mengalami overfitting.

$$Y = \text{mode}(y_1 + y_2 + \dots + y_n) \quad (2)$$

Di sisi lain, Naive Bayes menggunakan pendekatan probabilistik yang mengasumsikan independensi antar fitur. Model ini menghitung probabilitas setiap kelas dan memilih kelas dengan probabilitas tertinggi sebagai hasil klasifikasi.

$$P(C|X) = \frac{P(x_1|C) \times P(x_2|C) \times \dots \times P(x_n|C) \times P(C)}{P(X)} \quad (3)$$

Dalam konteks deteksi penipuan kartu kredit, Naive Bayes dapat memberikan hasil yang cepat dan efisien, meskipun dengan asumsi independensi yang mungkin tidak selalu terpenuhi [20]. Penerapan kedua algoritma ini memberikan gambaran yang jelas tentang bagaimana masing-masing algoritma dapat berfungsi dalam konteks yang berbeda.

2.5 Pengukuran Performa

Setelah penerapan algoritma, langkah terakhir adalah mengukur performa model menggunakan metrik evaluasi yang sesuai. Metrik yang umum digunakan dalam deteksi penipuan adalah akurasi, presisi, recall, dan F1-

score. Akurasi mengukur seberapa banyak prediksi yang benar dibandingkan dengan total prediksi, sementara presisi dan recall memberikan wawasan lebih dalam tentang kemampuan model dalam mendeteksi penipuan [21], [22], [23]. F1-score, yang merupakan harmonisasi antara presisi dan recall, sangat penting dalam konteks ini karena ketidakseimbangan kelas yang ada dalam data.

$$\text{Accuracy} = \frac{TP+TN}{TP+TN+FP+FN} \quad (4)$$

$$\text{Precision} = \frac{TP}{TP+FP} \quad (5)$$

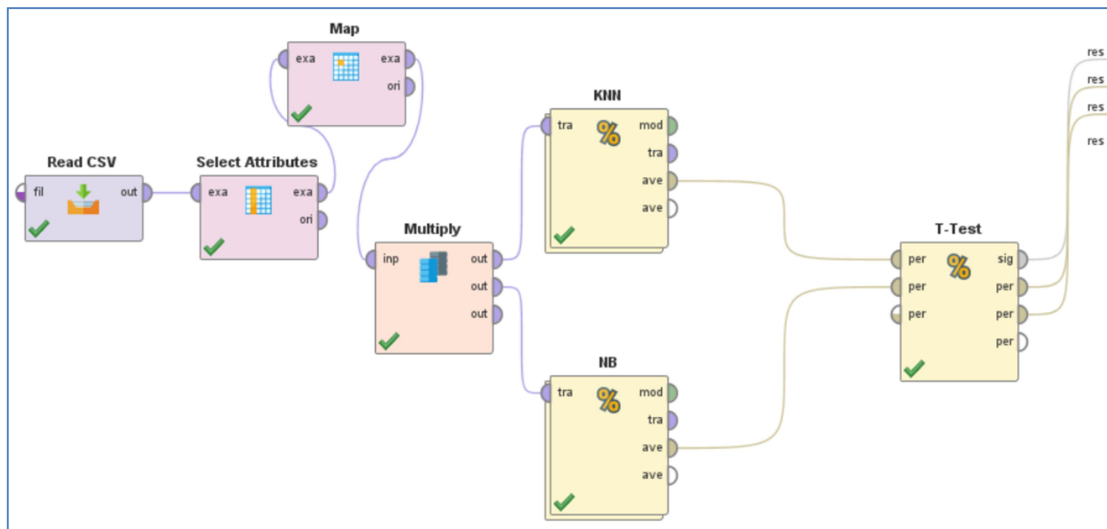
$$\text{Recall} = \frac{TP}{TP+FN} \quad (6)$$

$$F1 - \text{score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (7)$$

3. HASIL DAN PEMBAHASAN

3.1 Arsitektur Eksperimen

Penelitian ini menggunakan perangkat lunak RapidMiner untuk membandingkan performa algoritma K-Nearest Neighbor (KNN) dan Naive Bayes (NB) dalam mendeteksi penipuan kartu kredit. Alur kerja eksperimen dirancang untuk memastikan kedua algoritma menerima perlakuan data yang identik guna menghasilkan perbandingan yang objektif.



Gambar 1. Arsitektur Pemrosesan Perbandingan Algoritma

3.2 Pra-pemrosesan Data

Tahap awal dimulai dengan operator Read CSV untuk memuat dataset transaksi kartu kredit. Selanjutnya, dilakukan tahapan berikut:

Select Attributes: Digunakan untuk memilih fitur-fitur yang paling relevan dengan deteksi penipuan. Dalam hal ini

yang dilakukan adalah mengecualikan id_transaksi dalam pemrosesan data.

Map: Digunakan untuk menyelaraskan record tertentu dalam atribut, dikarenakan algoritma yang digunakan lebih cocok terhadap data yang bersifat angka, di tahap ini yang dilakukan adalah merubah isi dari atribut merchant category.

Tabel 1. Merubah kategori

No	Sebelum	Sesudah
1	Clothing	0
2	Electronics	1
3	Food	2
4	Grocery	3
5	Travel	4

Multiply: Operator ini berfungsi untuk menduplikasi dataset yang telah diproses agar dapat dialirkan secara bersamaan ke kedua algoritma pengujian tanpa ada perbedaan distribusi data awal.

kemiripan jarak antar data transaksi. Transaksi baru akan diklasifikasikan sebagai "Fraud" atau "Normal" berdasarkan mayoritas label dari tetangga terdekatnya. Berdasarkan hasil percobaan yang dilakukan akurasi yang dihasilkan algoritma ini adalah 98.43%, precision 98.53%, recall 99.90% dan F1-score 99.21%.

3.3 Implementasi Algoritma

Dua model klasifikasi dibangun secara paralel:

- a. K-Nearest Neighbor (KNN):
Algoritma ini bekerja berdasarkan

Tabel 2. Confusion Matrix Algoritma K-NN

	Normal	Fraud	class precision
Normal	2952	44	98.53%
Fraud	3	1	25.00%
class recall	99.90%	2.22%	

- b. Naive Bayes (NB): Algoritma ini menggunakan pendekatan probabilitas Bayes dengan asumsi independensi antar fitur. Model ini menghitung peluang suatu transaksi menjadi "Fraud" berdasarkan distribusi data historis.

Hasil percobaan yang dilakukan dengan algoritma naïve bayes didapatkan nilai akurasi sebesar 98.20%, precision 99.69%, recall 98.48% dan F1-score 99.08%.

Tabel 3. Confusion Matrix Algoritma Naïve Bayes

	Normal	Fraud	class precision
Normal	2910	9	99.69%
Fraud	45	36	44.44%
class recall	98.48%	80.00%	

3.4 Analisis Perbandingan Menggunakan T-Test

Fitur utama dalam eksperimen ini adalah penggunaan operator T-Test. Berbeda dengan evaluasi akurasi standar, T-Test digunakan untuk menentukan apakah perbedaan performa (seperti accuracy, precision, atau recall) antara KNN dan Naive Bayes bersifat signifikan secara statistik atau hanya terjadi secara kebetulan.

Interpretasi Hasil T-Test:

- Jika nilai p-value $< 0,05$, maka terdapat perbedaan performa yang signifikan antara KNN dan Naive Bayes dalam mendeteksi penipuan.
- Jika nilai p-value $> 0,05$, maka kedua algoritma dianggap memiliki performa yang setara secara statistik pada dataset ini.

Berdasarkan hasil pengukuran akurasi yang didapatkan dari masing-masing algoritma, dapat disimpulkan bahwa performa dari kedua algoritma dianggap setara karena tidak ada perbedaan yang signifikan, walaupun berdasarkan data algoritma KNN lebih unggul dari Naive Bayes.

3.5 Diskusi

Hasil dari penelitian ini menunjukkan bahwa K-Nearest Neighbor (KNN) dan Naive Bayes memiliki performa yang berbeda dalam mendeteksi penipuan kartu kredit. Berdasarkan pengujian yang dilakukan, KNN menghasilkan akurasi sebesar 98,43%, sedangkan Naive Bayes hanya mencapai akurasi 98,20%. Meskipun kedua algoritma menunjukkan hasil yang baik, perbedaan ini menunjukkan bahwa KNN lebih efektif dalam mengidentifikasi transaksi yang mencurigakan dalam konteks dataset ini.

Kelebihan KNN terletak pada kemampuannya untuk menangkap pola yang lebih kompleks dalam data, berkat

pendekatan berbasis tetangga terdekat. Dalam kasus penipuan kartu kredit, di mana pola transaksi yang mencurigakan sering kali tidak terduga, kemampuan KNN untuk mempertimbangkan beberapa tetangga terdekat memberikan keuntungan tersendiri. Namun, KNN juga memiliki kelemahan, yaitu kinerjanya yang dapat menurun drastis dengan meningkatnya ukuran dataset, karena memerlukan waktu komputasi yang lebih lama untuk menghitung jarak antara titik data.

Sementara itu, Naive Bayes, meskipun lebih cepat dalam hal waktu komputasi, memiliki asumsi independensi antar fitur yang sering kali tidak terpenuhi. Dalam konteks penipuan kartu kredit, fitur-fitur seperti jumlah transaksi dan waktu transaksi mungkin saling terkait, yang dapat mengurangi efektivitas model Naive Bayes. Namun, meskipun memiliki keterbatasan, Naive Bayes tetap menjadi pilihan yang baik dalam situasi di mana kecepatan dan efisiensi lebih diutamakan daripada akurasi [24].

Dalam analisis lebih lanjut, kami juga menemukan bahwa KNN memiliki tingkat recall yang lebih tinggi dibandingkan Naive Bayes, yang berarti bahwa KNN lebih efektif dalam mengidentifikasi transaksi penipuan yang sebenarnya. Hal ini sangat penting dalam konteks deteksi penipuan, di mana tujuan utamanya adalah untuk meminimalkan jumlah transaksi yang terlewatkan. Sebaliknya, Naive Bayes cenderung menghasilkan lebih banyak false negatives, yang dapat menyebabkan kerugian finansial bagi lembaga keuangan dan konsumen.

Secara keseluruhan, hasil penelitian ini menunjukkan bahwa meskipun KNN dan Naive Bayes memiliki aplikasi yang berbeda dalam deteksi penipuan kartu kredit, KNN lebih unggul dalam hal akurasi dan kemampuan mendeteksi pola

kompleks. Namun, pemilihan algoritma yang tepat harus mempertimbangkan konteks dan kebutuhan spesifik dari sistem yang sedang dikembangkan.

4. KESIMPULAN

Penelitian ini berhasil membandingkan efektivitas algoritma K-Nearest Neighbor dan Naive Bayes dalam mendeteksi penipuan kartu kredit menggunakan dataset yang diambil dari Kaggle. Hasil analisis menunjukkan bahwa KNN memiliki akurasi yang lebih tinggi dibandingkan Naive Bayes, serta kemampuan yang lebih baik dalam mendeteksi transaksi penipuan yang sebenarnya. Namun, KNN juga memiliki kelemahan dalam hal waktu komputasi, terutama saat dihadapkan pada dataset yang besar.

Di sisi lain, meskipun Naive Bayes lebih cepat dan efisien, asumsi independensi antar fitur dapat membatasi efektivitasnya dalam konteks yang lebih kompleks. Oleh karena itu, pemilihan algoritma yang tepat harus mempertimbangkan berbagai faktor, termasuk ukuran dataset, kompleksitas pola yang ingin dideteksi, dan kebutuhan akan kecepatan dalam pengambilan keputusan.

Ke depan, penelitian ini membuka peluang untuk eksplorasi lebih lanjut mengenai kombinasi algoritma atau penerapan teknik ensemble yang dapat meningkatkan akurasi deteksi penipuan. Selain itu, penelitian lebih lanjut juga dapat mempertimbangkan penggunaan fitur tambahan dan teknik pemrosesan data yang lebih canggih untuk meningkatkan performa model. Dengan demikian, diharapkan sistem deteksi penipuan kartu kredit dapat terus ditingkatkan untuk melindungi konsumen dan lembaga keuangan dari ancaman penipuan yang semakin berkembang.

5. REFERENSI

- [1] N. Report, "Card Fraud Losses Worldwide in 2023." Accessed: Dec. 02, 2025. [Online]. Available: <https://nilsonreport.com/articles/card-fraud-losses-worldwide-in-2023/>
- [2] S. M. Pais and S. Acharya, "Hybrid Algorithm for Naïve Bayes-Based Credit Card Fraud Detection," 2021. doi: 10.1007/978-981-15-9873-9_36.
- [3] R. O. Ogundokun, S. Misra, O. J. Fatigun, and J. K. Adeniyi, "Naïve Bayes Based Classifier for Credit Card Fraud Discovery," in *Lecture Notes in Business Information Processing*, 2022. doi: 10.1007/978-3-030-95947-0_37.
- [4] T. Vairam, S. Sarathambekai, S. Bhavadharani, A. Kavi Dharshini, N. Nithya Sri, and T. Sen, "Evaluation of Naïve Bayes and Voting Classifier Algorithm for Credit Card Fraud Detection," in *8th International Conference on Advanced Computing and Communication Systems, ICACCS 2022*, 2022. doi: 10.1109/ICACCS54159.2022.9784968.
- [5] N. Samy and S. mohamed mohamed, "Credit Card Fraud Detection Using Machine Learning Techniques," *Future Computing and Informatics Journal*, vol. 7, no. 1, 2022, doi: 10.54623/fue.fcij.7.1.2.
- [6] A. Jessica, F. V. Raj, and J. Sankaran, "Credit Card Fraud Detection Using Machine Learning Techniques," in *ViTECoN 2023 - 2nd IEEE International Conference on Vision Towards Emerging Trends in Communication and Networking Technologies, Proceedings*, 2023. doi: 10.1109/ViTECoN58111.2023.10157162.
- [7] E. Tank and M. Das, "On Credit Card Fraud Detection Using Machine Learning Techniques," in

- Lecture Notes in Networks and Systems*, 2024. doi: 10.1007/978-981-97-2004-0_21.
- [8] J. Chung and K. Lee, "Credit Card Fraud Detection: An Improved Strategy for High Recall Using KNN, LDA, and Linear Regression," *Sensors*, vol. 23, no. 18, 2023, doi: 10.3390/s23187788.
- [9] Jiwon Chung and Kyungho Lee, "Credit Card Fraud Detection: An Improved Strategy for High Recall Using KNN, LDA, and Linear Regression," *Sensors*, vol. 23, no. 18, 2023.
- [10] S. Said, M. Al Khadhori, A. Juma, K. Al Mukhaini, and V. Sherimon, "Machine Learning Approaches for Credit Card Fraud Detection: a Predictive Analysis," *International Journal of Artificial Intelligence & Machine Learning (IJAIML)*, vol. 3, no. 1, 2024.
- [11] S. Baba Atiku and A. A. Unimke, "Machine Learning Approach to Detection of Credit Card Fraud," *International Journal of Applied Science and Research*, vol. 08, no. 04, 2025, doi: 10.56293/ijasr.2025.6604.
- [12] S. M. Ghatge, "Machine Learning Approach for Credit Card fraud Detection," *Int J Res Appl Sci Eng Technol*, vol. 10, no. 5, 2022, doi: 10.22214/ijraset.2022.43587.
- [13] P. N. Pote, "Credit Card Fraud Detection: A Machine Learning Approach," *INTERANTIONAL JOURNAL OF SCIENTIFIC RESEARCH IN ENGINEERING AND MANAGEMENT*, vol. 08, no. 04, 2024, doi: 10.55041/ijsrem31981.
- [14] A. Miah, "Credit Card Fraud Detection Dataset." Accessed: Dec. 18, 2025. [Online]. Available: <https://www.kaggle.com/datasets/miadul/credit-card-fraud-detection-dataset>
- [15] K. Arora, S. Pathak, and N. T. Dieu Linh, "Comparative Analysis of K-Nn, Naïve Bayes, and logistic regression for credit card fraud detection," *Ingenieria Solidaria*, vol. 19, no. 3, 2023, doi: 10.16925/2357-6014.2023.03.05.
- [16] H. Wang and T. M. Khoshgoftaar, "Review of Feature Selection Techniques for Credit Card Fraud Detection," in *28th ISSAT International Conference on Reliability and Quality in Design, RQD 2023*, 2023.
- [17] M. Pagan, M. Zarlis, and A. Candra, "Investigating the impact of data scaling on the k-nearest neighbor algorithm," *Computer Science and Information Technologies*, vol. 4, no. 2, 2023, doi: 10.11591/csit.v4i2.pp135-142.
- [18] P. J. Muhammad Ali, "Investigating the Impact of Min-Max Data Normalization on the Regression Performance of K-Nearest Neighbor with Different Similarity Measurements," *ARO-The Scientific Journal of Koya University*, vol. 10, no. 1, 2022, doi: 10.14500/aro.10955.
- [19] M. Baloch, M. S. Honnurvali, A. Kabbani, T. A. Jumani, and S. T. Chauhdary, "An Intelligent SARIMAX-Based Machine Learning Framework for Long-Term Solar Irradiance Forecasting at Muscat, Oman," *Energies (Basel)*, vol. 17, no. 23, 2024, doi: 10.3390/en17236118.
- [20] I. K. Prasetya, D. P. Isnawarty, A. Fahmi, S. A. P. Andikaputra, and W. Wibawati, "Comparing the Performance of Multivariate Hotelling's T2 Control Chart and Naive Bayes Classifier for Credit Card Fraud Detection," *Inferensi*, vol. 7, no. 1, 2024, doi: 10.12962/j27213862.v7i1.18755.

- [21] T. Rak, J. Drabek, and M. Charytanowicz, "Performance-Based Classification of Users in a Containerized Stock Trading Application Environment Under Load," *Electronics (Switzerland)*, vol. 14, no. 14, 2025, doi: 10.3390/electronics14142848.
- [22] M. Aviles, L. M. Sánchez-Reyes, J. M. Álvarez-Alvarado, and J. Rodríguez-Reséndiz, "Machine and Deep Learning Trends in EEG-Based Detection and Diagnosis of Alzheimer's Disease: A Systematic Review," 2024. doi: 10.3390/eng5030078.
- [23] S. Gupta, B. Kishan, and P. Gulia, "Comparative Analysis of Predictive Algorithms for Performance Measurement," *IEEE Access*, vol. 12, 2024, doi: 10.1109/ACCESS.2024.3372082.
- [24] D. Sharma, E. Kaur, N. Sharma, D. Rawat, S. P. Yadav, and M. Manwal, "Comparative Analysis of Machine Learning Algorithms on Credit Card Fraud Detection," in *2025 International Conference on Cognitive Computing in Engineering, Communications, Sciences and Biomedical Health Informatics, IC3ECSBHI 2025*, 2025. doi: 10.1109/IC3ECSBHI63591.2025.10990921.