

KLASIFIKASI RISIKO GAGAL BAYAR KREDIT MENGGUNAKAN XGBOOST DENGAN MODEL EXPLAINABLE BERBASIS SHAP

Aditia Ananda Sutrisno, Ahmad Abdul Chamid, Rizkysari Meimaharani

Teknik Informatika, Universitas Muria Kudus

Jl. Lkr. Utara, Kayuapu Kulon, Gondangmanis, Kec. Bae, Kabupaten Kudus, Jawa Tengah

adityaananda883@gmail.com

ABSTRAK

Risiko gagal bayar kredit merupakan permasalahan penting dalam sektor keuangan yang berpotensi mengganggu stabilitas lembaga keuangan apabila tidak dikelola secara efektif. Permasalahan utama dalam prediksi risiko kredit moderen adalah kebutuhan akan model yang tidak hanya memiliki akurasi tinggi, tetapi juga mampu memberikan penjelasan yang transparan terhadap hasil prediksi. Penelitian ini bertujuan untuk mengembangkan model klasifikasi risiko gagal bayar kredit menggunakan algoritma *Extreme Gradient Boosting* (XGBoost) serta menganalisis interpretabilitas model melalui pendekatan *Explainable Artificial Intelligence* berbasis *Shapley Additive Explanations* (SHAP). Metode penelitian meliputi tahap *preprocessing* data berupa pembersihan data, imputasi nilai hilang menggunakan *KNN Imputer*, *encoding* variabel kategorikal, normalisasi data, pelatihan model XGBoost, serta evaluasi kinerja menggunakan metrik akurasi, *precision*, *recall*, dan *F1-score*. Dataset yang digunakan berasal dari Kaggle dengan jumlah 32.586 data nasabah dan 13 fitur. Hasil penelitian menunjukkan bahwa model XGBoost mampu mencapai akurasi sebesar 0,93, dengan nilai *recall* sebesar 0,98 untuk kelas *no default* dan 0,74 untuk kelas *default*. Analisis SHAP mengidentifikasi *customer_income*, *loan_grade*, *loan_amnt*, dan *loan_int_rate* sebagai fitur yang paling berpengaruh terhadap prediksi risiko gagal bayar. Hasil ini menunjukkan bahwa integrasi XGBoost dan SHAP mampu menghasilkan model prediksi risiko kredit yang akurat sekaligus transparan, sehingga dapat mendukung pengambilan keputusan dalam sistem evaluasi risiko kredit.

Kata kunci : XGBoost, Credit Default Risk, SHAP, Explainable AI, Machine Learning

1. PENDAHULUAN

Sektor finansial, sebagai pilar penunjang utama perekonomian, secara inheren dihadapkan pada tantangan sentral terkait *credit default risk* [1]. Risiko kredit, yang didefinisikan sebagai ketidakmampuan debitur dalam memenuhi kewajiban yang telah disepakati, merupakan perhatian signifikan bagi institusi keuangan [2]. Berdasarkan laporan KOMPAS tahun 2025 yang merujuk pada data Otoritas Jasa Keuangan (OJK) rasio kredit macet di industri pinjaman daring mencapai 2,93 persen pada April 2025, meningkat dari 2,77 persen pada Maret 2025 dan 2,6 persen pada akhir 2024 [3]. Jika tidak dikelola secara efektif, risiko ini dapat memicu kerugian finansial yang besar dan mengganggu stabilitas sistem perbankan secara keseluruhan, sebagaimana yang pernah terlihat dalam peristiwa krisis keuangan global [4]. Oleh karena itu, pengembangan model prediksi yang akurat, andal, dan transparan merupakan prasyarat mutlak dalam manajemen risiko kredit modern untuk menghindari potensi kerugian dan menjaga profitabilitas institusi [5], [6].

Seiring dengan pesatnya digitalisasi dan peningkatan volume data pengajuan pinjaman, metode tradisional dalam penilaian kelayakan kredit, yang sering mengandalkan variabel tetap dan bersifat manual, menjadi kurang adaptif dalam menangkap kompleksitas perilaku finansial moderen [2], [4], [6], [7]. Kekurangan ini mendorong adopsi teknologi pembelajaran mesin (*Machine Learning*) sebagai

solusi berbasis data untuk mengevaluasi risiko secara lebih dinamis, akurat, dan efisien [8], [9], [10], [11]. Dalam konteks klasifikasi risiko, algoritma *ensemble learning* berbasis *boosting*, seperti *Extreme Gradient Boosting* (XGBoost), telah diakui karena kinerja prediktifnya yang unggul, efektivitasnya dalam menangani kasus klasifikasi berskala besar, dan kemampuannya mencapai tingkat akurasi klasifikasi yang tinggi.

Meskipun XGBoost menunjukkan performa yang superior, kompleksitas internal arsitektur pohon-basisnya menjadikan model ini rentan diklasifikasikan sebagai "kotak hitam" (*black-box*) [12]. Kurangnya penjelasan yang memadai mengenai mengapa model mengambil keputusan tertentu, terutama dalam domain sensitif seperti keuangan yang menuntut akuntabilitas tinggi, menghambat kepercayaan pengguna, menyulitkan identifikasi bias, dan menghambat adopsi teknologi AI secara bertanggung jawab [12], [13], [14].

Untuk mengatasi masalah krusial antara akurasi dan transparansi ini, penelitian ini mengintegrasikan kerangka kerja *Explainable AI*, khususnya metode *Shapley Additive Explanations* (SHAP) [15]. SHAP, yang didasarkan pada teori permainan kooperatif, memberikan solusi interpretasi dengan menghitung secara aditif kontribusi setiap fitur terhadap hasil prediksi, baik secara global (keseluruhan model) maupun lokal (per individu) [15]. Integrasi SHAP memungkinkan model XGBoost yang akurat untuk

tidak hanya memberikan prediksi kelayakan kredit, tetapi juga menyajikan penjelasan yang kuantitatif dan intuitif mengenai faktor-faktor dominan yang mendorong keputusan tersebut, sebuah pendekatan yang terbukti efektif dalam menganalisis dinamika risiko gagal bayar [12]. Dengan demikian, penelitian ini bertujuan untuk meningkatkan keandalan dan akurasi klasifikasi risiko gagal bayar menggunakan XGBoost, sekaligus menyediakan transparansi dan interpretabilitas yang tinggi melalui analisis SHAP, sehingga dapat mendukung pengambilan keputusan finansial yang lebih terinformasi dan membangun kepercayaan yang lebih baik terhadap sistem prediksi berbasis pembelajaran mesin.

2. TINJAUAN PUSTAKA.

2.1. Penelitian Terdahulu

Penelitian dari Rahmi Anadra, dkk. Tahun 2024 tentang klasifikasi persetujuan pinjaman (*Loan Approval Classification*). Metode yang digunakan adalah *ensemble learning*, membandingkan Random Forest dan XGBoost pada dataset 4.296 data, di mana data tidak seimbang ditangani menggunakan *Synthetic Minority Over-sampling Technique* (SMOTE). Hasilnya menunjukkan XGBoost dengan penanganan data tidak seimbang mencapai tingkat akurasi tertinggi sebesar 98,52%.. Studi ini memvalidasi pemilihan XGBoost sebagai algoritma yang unggul dan sangat akurat dalam konteks risiko kredit [5].

Tahun 2024, Muhammad Iqbal, dkk. Melakukan analisis kinerja *ensemble learning* untuk mendapatkan nilai akurasi terbaik dari dataset prediksi persetujuan pinjaman dari Kaggle. Penelitian ini membandingkan XGBoost dengan berbagai model Decision Tree *ensemble* lain, termasuk LightGBM, Gradient Boosting, Random Forest, Adaboost. Hasil perbandingan menunjukkan bahwa XGBoost adalah model terbaik di antara semua algoritma yang diuji, mencapai Akurasi 0.9974, dengan AUC 0.9998, dan *F1-Score* 0.9966. Penelitian ini memperkuat justifikasi bahwa XGBoost adalah pilihan optimal untuk mencapai performa prediksi tertinggi dalam klasifikasi kelayakan pinjaman [7].

2.2. Machine Learning

Machine Learning merupakan cabang dari kecerdasan buatan (*Artificial Intelligence*) yang memungkinkan sistem komputer belajar secara otomatis dari data tanpa harus diprogram secara eksplisit untuk setiap tugas tertentu. Proses pembelajaran ini dilakukan melalui identifikasi pola dan hubungan dalam data sehingga sistem dapat melakukan prediksi, klasifikasi, maupun pengambilan keputusan berdasarkan pengalaman sebelumnya [16].

2.3. Algoritma XGBOOST

Extreme Gradient Boosting (XGBoost) merupakan algoritma *ensemble learning* canggih berbasis pohon keputusan, yang dikembangkan sebagai sistem tree boosting yang sangat terukur dan

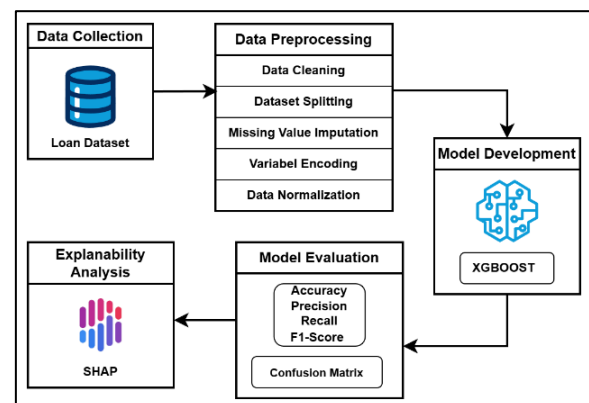
efisien. XGBoost adalah pengembangan dari *Gradient Boosting Machine* (GBM), dirancang untuk kinerja dan kecepatan superior saat memproses data dalam skala besar. Inti dari XGBoost terletak pada kemampuannya menggabungkan serangkaian model lemah secara iteratif untuk menghasilkan model prediksi yang jauh lebih kuat dan akurat [2], [5].

2.4. SHAP

Shapley Additive Explanations (SHAP) merupakan metodologi unggulan dalam ranah *Explainable AI* (XAI) yang berfungsi sebagai alat atribusi fitur pasca-model (*post-hoc*) untuk menjelaskan prediksi model *black-box* dengan tetap menjaga kinerja. Landasan teoretis SHAP berasal dari konsep *Shapley values* dari teori permainan kooperatif, yang menawarkan cara yang adil untuk mengalokasikan kontribusi setiap fitur (dianggap sebagai pemain dalam permainan koalisi) terhadap prediksi model [17].

3. METODE PENELITIAN

Metodologi penelitian dilakukan secara terstruktur menggunakan pendekatan berbasis data mining dan pembelajaran mesin. Tahapan divisualisasikan dalam gambar 1.



Gambar 1. Metodologi Penelitian

3.1. Data Collection

Data dikumpulkan atau didapatkan secara daring melalui Kaggle dengan judul *Loan Dataset* dari Prakash Raushan, 2023 yang terdiri dari 32.586 data nasabah dan 13 atribut.

3.2. Data Preprocessing

Mengingat dataset bersumber dari publik dan memiliki nilai kosong serta format yang beragam, maka dilakukan serangkaian tahap pembersihan dan transformasi data secara sistematis. Pertama adalah *data cleaning*, yaitu proses menghapus *noise* agar data dapat dianalisis dengan baik. Proses ini memperbaiki data yang salah, rusak, tidak akurat, atau berformat keliru untuk meningkatkan kualitasnya [18].

Tahap kedua *data splitting*, dataset dipisah menjadi dua *subset*, di mana 80% data pelatihan sementara 20% data pengujian. Proses pembagian ini

dilakukan menggunakan teknik *stratified sampling* berdasarkan fitur target guna menjaga konsistensi proporsi antara data pelatihan dan pengujian [19]. Selanjutnya yaitu *missing value imputation* dimana nilai kosong pada fitur numerik diisi menggunakan metode *K-Nearest Neighbors* (KNN) *Imputer* untuk mempertahankan konsistensi pola data [20].

Keempat adalah proses *variable encoding*. Variabel kategorikal yang bersifat ordinal ditransformasikan menjadi representasi numerik melalui teknik *Label Encoding*. Langkah ini diperlukan agar algoritma pembelajaran mesin dapat mengenali hubungan antar kategori dalam bentuk numerik tanpa kehilangan makna urutannya [21]. Variabel kategorikal *non-ordinal* dikonversi menjadi *dummy variables* menggunakan teknik *One-Hot Encoding* untuk menghindari asumsi urutan yang tidak berlaku [22]. Terakhir adalah *data normalization*. Data numerik diproses menggunakan teknik normalisasi *StandardScaler*, biar seluruh variabel memiliki skala yang setara dan mendukung kinerja model secara optimal [23].

3.3. Model Development

Model dibangun menggunakan algoritma XGBoost Classifier yang bekerja berdasarkan prinsip gradient boosting. Algoritma ini bekerja dengan menyesuaikan parameter pembelajaran secara berulang untuk meminimalkan loss function dan membangun model pohon regresi yang lebih teratur [24]. Pada tahap implementasi, model dijalankan tanpa proses penyesuaian parameter untuk menilai performanya dalam konfigurasi default. Inisialisasi model dilakukan dengan beberapa pengaturan, antara lain `use_label_encoder= False` guna menonaktifkan label encoder bawaan dari XGBoost, `eval_metric= 'logloss'` yang berfungsi menghitung tingkat kesalahan prediksi secara logaritmik, serta `random_state= 42` untuk menjamin konsistensi hasil pada setiap proses pelatihan model.

3.4. Model Evaluation

Evaluasi model dilakukan untuk menilai sejauh mana algoritma XGBoost Classifier mampu memprediksi risiko gagal bayar secara akurat dan konsisten. Dalam penelitian ini, proses evaluasi dilakukan dengan memakai *confusion matrix*. *Confusion matrix* digunakan untuk menampilkan hubungan perbandingan antara hasil prediksi model dan data sebenarnya [25].

Confusion matrix berisi metrik evaluasi, *accuracy* untuk mengukur seberapa banyak prediksi yang benar dari keseluruhan data [26]. *Precision* untuk mengukur tingkat ketepatan model saat memprediksi sebuah data sebagai positif [27]. *Recall* untuk mengukur tingkat ketepatan model dalam mendeteksi kembali semua data yang seharusnya positif [27]. *F1-Score* untuk nilai rata-rata harmonis antara *Precision* dan *Recall*, yang memberikan gambaran performa model secara seimbang [28].

3.5. Explainability Analysis

Tahap akhir adalah analisis *explainability* model menggunakan SHAP (*Shapley Additive Explanations*). Metode ini digunakan untuk mengukur kontribusi setiap fitur terhadap hasil prediksi. Analisis dilakukan secara global untuk melihat fitur yang paling berpengaruh terhadap keseluruhan model.

4. HASIL DAN PEMBAHASAN

4.1. Data Overview

Dataset penelitian diperoleh dari Kaggle (Prakash Raushan, 2023) yang berisi 32.586 data nasabah dengan 13 fitur. Beberapa fitur yang seharusnya numerik seperti *customer_income* dan *loan_amnt* masih bertipe *object* karena format data asli berupa *string*, sehingga memerlukan tahap *preprocessing*. Fitur target adalah *Current_loan_status* yang mengindikasikan status pinjaman dengan kategori *DEFAULT* dan *NO DEFAULT*. Deskripsi detail setiap atribut disajikan pada Tabel 1.

Tabel 1. Deskripsi Atribut Dataset

Fitur	Tipe Data	Keterangan
customer_id	Int	Nomer unik pelanggan
customer_age	Int	Usia pelanggan
customer_income	Object	Pendapatan tahunan
home_ownership	Object	Status kepemilikan rumah
employment_duration	Float	Lama bekerja
loan_intent	Object	Tujuan pinjaman
loan_grade	Object	Grade pinjaman
loan_amnt	Float	Jumlah pinjaman yang diajukan
loan_int_rate	Float	Suku bunga pinjaman
term_years	Int	Lama Pinjaman
historical_default	Object	Riwayat gagal bayar
cred_hist_length	Int	Lama riwayat kredit
Current_loan_status	Object	Status pinjaman (Target)

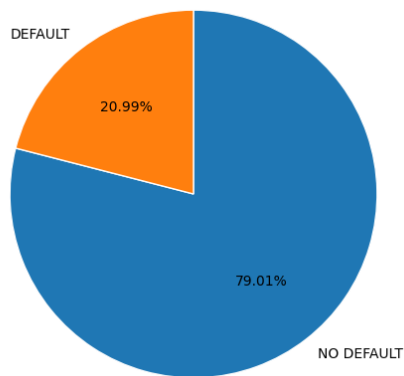
Pada dataset ini ditemukan beberapa fitur dengan nilai kosong, fitur *historical_default* memiliki 63.64% nilai hilang, diikuti oleh *loan_int_rate* 9.56% dan *employment_duration* 2.75%. Informasi ini menunjukkan perlunya tahap penanganan *missing value* pada bagian *preprocessing* selanjutnya. penjelasan lengkap terdapat pada Table 2.

Tabel 2. Missing Value pada Setiap Fitur

Fitur	Missing	Presentase
historical_default	20.737	63.64 %
loan_int_rate	3.116	9.56 %
employment_duration	895	2.75 %
Current_loan_status	4	0.01 %
loan_amnt	1	0.00 %

Distribusi pada fitur target menunjukkan bahwa kelas *NO DEFAULT* mendominasi sebesar 79.01%,

sedangkan kelas *DEFAULT* hanya 20.99%. Untuk lebih jelasnya, dapat dilihat pada Gambar 2.



Gambar 2. Distribusi Kelas Target

4.2. Data Preprocessing

Tahapan data *preprocessing* meliputi lima langkah utama, yaitu *data cleaning*, *data splitting*, *missing value imputation*, *encoding variabel*, dan *data normalization*. Urutan ini dipilih untuk mencegah *data leakage* serta menjaga keaslian distribusi data pada proses pelatihan dan pengujian model.

Langkah awal dilakukan dengan menghapus fitur *historical_default* yang memiliki proporsi *missing values* tinggi. Selain itu, atribut *customer_income* dan *loan_amnt* dibersihkan dari simbol mata uang (£) dan tanda koma (,) sebelum dikonversi menjadi tipe numerik. Nilai *customer_age* juga dibatasi maksimal 100 tahun untuk menghindari anomali data. Selanjutnya, dataset dibagi menjadi data latih (80%) dan data uji (20%) dengan *stratified sampling*.

Nilai kosong pada fitur *employment_duration*, *loan_amnt*, dan *loan_int_rate* diimputasi menggunakan *K-Nearest Neighbors Imputer* ($k = 5$). Fitur kategorikal *home_ownership* dan *loan_intent* ditransformasikan menggunakan *One-Hot Encoding*. Sementara itu, fitur *loan_grade* menggunakan *Label Encoding* karena bersifat ordinal. Terakhir, seluruh fitur numerik dinormalisasi menggunakan *StandardScaler*. Hasil akhir dari tahap ini berupa data latih dan data uji yang siap digunakan.

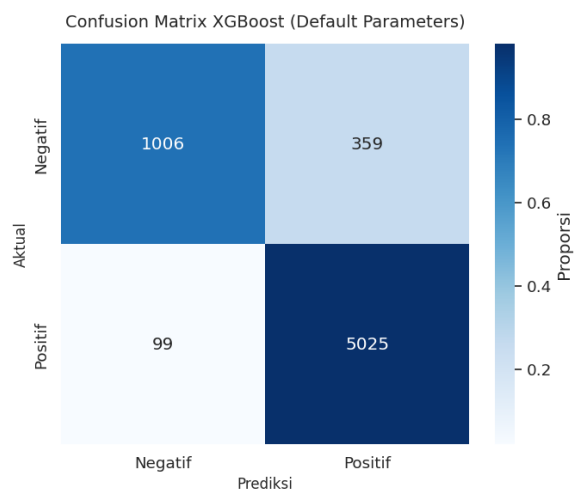
4.3. Model Evaluation

Berdasarkan hasil *classification report* pada gambar 3, model XGBoost menunjukkan performa klasifikasi yang sangat baik dengan nilai akurasi keseluruhan sebesar 0.93. Kelas positif (1) memperoleh nilai *precision* sebesar 0.93, *recall* sebesar 0.98, dan *F1-score* sebesar 0.96, menandakan model sangat efektif dalam mengidentifikasi sampel positif dengan tingkat kesalahan yang rendah. Sementara itu, kelas negatif (0) memiliki *precision* sebesar 0.91, *recall* sebesar 0.74, dan *F1-score* sebesar 0.81, yang menunjukkan adanya sebagian kecil data negatif yang masih salah terklasifikasi sebagai positif.

Class	Precision	Recall	F1-Score	Support
0	0.91	0.74	0.81	1365
1	0.93	0.98	0.96	5124
Accuracy	0.93			

Gambar 3. Classification Report

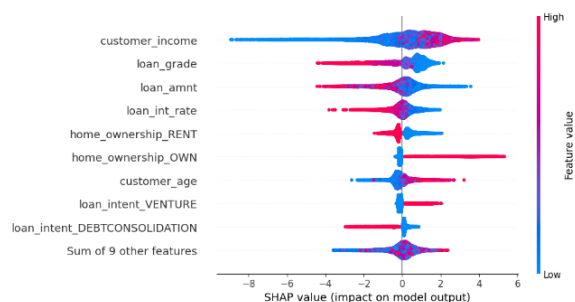
Confusion matrix menunjukkan bahwa model XGBoost berhasil mengklasifikasikan mayoritas data dengan benar. Sebanyak 5.025 data positif terprediksi dengan benar (*true positive*) dan 1.006 data negatif teridentifikasi sesuai label aslinya (*true negative*). Hanya terdapat 99 kasus *false negative* dan 359 kasus *false positive*, yang menunjukkan model memiliki kemampuan deteksi positif yang sangat baik. Lebih jelasnya ditampilkan pada gambar 4 dibawah.



Gambar 4. Confusion Matrix

4.4. Explainability Analysis Dengan SHAP

Memuat hasil, Pengujian dan pembahasan yang telah dilakukan Pada konteks penelitian ini, kelas positif (1) merepresentasikan nasabah dengan status *no default*, sedangkan kelas negatif (0) menunjukkan *default*. Untuk memahami bagaimana setiap fitur memengaruhi hasil prediksi model, dilakukan analisis menggunakan metode SHAP yang divisualisasikan pada Gambar 5.



Gambar 5. SHAP Value

Fitur yang paling berpengaruh terhadap keputusan model XGBoost adalah *customer_income*, diikuti oleh *loan_grade*, *loan_amnt*, dan *loan_int_rate*. Nilai SHAP positif pada pendapatan tinggi menunjukkan bahwa semakin besar penghasilan nasabah, semakin tinggi pula kecenderungan model memprediksi kelas positif, sedangkan pinjaman dengan jumlah besar atau tingkat bunga tinggi cenderung menurunkan peluang tersebut. Fitur *loan_grade* berperan positif, di mana grade pinjaman yang baik meningkatkan probabilitas prediksi positif. Sementara itu, faktor demografis seperti *home_ownership* dan *customer_age* berpengaruh lebih kecil, namun tetap memberikan kontribusi terhadap stabilitas prediksi model. Secara keseluruhan, hasil ini menunjukkan bahwa model XGBoost secara konsisten menilai faktor finansial utama sebagai indikator dominan dalam proses klasifikasi, selaras dengan prinsip evaluasi risiko kredit pada sektor keuangan.

5. KESIMPULAN DAN SARAN

Penelitian ini berhasil membuktikan bahwa model XGBoost sangat efektif dalam memprediksi status kelayakan pinjaman nasabah, mencapai akurasi keseluruhan sebesar 0.93. Kinerja model menunjukkan keunggulan yang signifikan dalam mengidentifikasi nasabah dengan pembayaran lancar (kelas *no default*), dibuktikan dengan nilai *precision* 0.93, *recall* 0.98, dan *F1-score* 0.96. Meskipun demikian, model menunjukkan sedikit keterbatasan dalam mendeteksi kasus gagal bayar (kelas *default*) dengan nilai *recall* 0.74, menunjukkan masih ada potensi perbaikan. Analisis SHAP mengidentifikasi *customer_income*, *loan_grade*, *loan_amnt*, dan *loan_int_rate* sebagai fitur paling berpengaruh dalam menentukan risiko gagal bayar. Secara keseluruhan, model XGBoost terbukti menjadi metode yang akurat dan terinterpretasi dengan baik, menjadikannya relevan untuk aplikasi dalam analisis risiko kredit.

Untuk penelitian berikutnya, disarankan untuk fokus pada penanganan ketidakseimbangan kelas (*class imbalance*) dalam dataset guna meningkatkan kemampuan model dalam mendeteksi kasus *default* secara lebih akurat. Penerapan teknik seperti SMOTE, *undersampling*, atau penyesuaian bobot kelas direkomendasikan untuk memperbaiki distribusi data dan performa model pada kelas minoritas. Selain itu, eksplorasi lebih lanjut terhadap penyempurnaan *hyperparameter* dan pengujian dengan *model gradient boosting* lainnya, seperti LightGBM atau CatBoost, dianjurkan untuk memperoleh hasil prediksi yang lebih optimal dan stabil dalam analisis risiko kredit.

DAFTAR PUSTAKA

- [1] A. Haque and M. M. Hassan, "BANK LOAN PREDICTION USING MACHINE LEARNING TECHNIQUES," 2024. doi: 10.48550/arXiv.2410.08886.
- [2] R. A. Pora, D. Maulina, and N. Trihartanti, "Komparasi XGBoost dan C4.5 dalam Klasifikasi Risiko Kredit untuk Kelayakan Pinjaman di India," *IJCSR: The Indonesian Journal of Computer Science Research*, vol. 4, no. 2, 2025, [Online]. Available: <https://subset.id/index.php/IJCSR>
- [3] KOMPAS, "Rekening Macet Pinjaman Daring Melonjak hingga 51 Persen," KOMPAS. Accessed: Nov. 05, 2025. [Online]. Available: <https://www.kompas.id/artikel/rekening-macet-pinjaman-daring-melonjak-hingga-51-persen>
- [4] S. Prayudani, Y. Sibarani, A. Salam, and A. R. Lubis, "Perbandingan Kinerja Model Pembelajaran Mesin Random Forest dan K-Nearest Neighbor (KNN) untuk Prediksi Risiko Kredit pada Layanan Pinjaman Online," *Journal Software, Hardware and Information Technology*, vol. 5, no. 2, pp. 118–127, Jun. 2025, doi: 10.24252/shift.v5i2.204.
- [5] R. Anadra, K. Sadik, A. M. Soleh, and R. A. Astari, "Loan Approval Classification Using Ensemble Learning on Imbalanced Data," *ENTHUSIASTIC INTERNATIONAL JOURNAL OF APPLIED STATISTICS AND DATA SCIENCE*, 2024, [Online]. Available: <https://journal.uir.ac.id/ENTHUSIASTIC>
- [6] V. Sinap, "A Comparative Study of Loan Approval Prediction Using Machine Learning Methods," *Gazi Üniversitesi Fen Bilimleri Dergisi Part C: Tasarım ve Teknoloji*, vol. 12, no. 2, pp. 644–663, Jun. 2024, doi: 10.29109/gujsc.1455978.
- [7] M. Iqbal, R. Dahlia, M. Ifan Rifani Ihsan, and R. Sa, "Performance Analysis of Ensemble Learning and Feature Selection Methods in Loan Approval Prediction at Banks," *Journal of Artificial Intelligence and Engineering Applications*, vol. 3, no. 2, pp. 2808–4519, 2024, [Online]. Available: <https://ioinformatic.org/>
- [8] A. A. Chamid, R. Nindyasari, N. Azizah, and A. Hariyadi, "Analysis of Public Opinion on The Governor Candidate Debate Using LDA and IndoBERT," *Kinetik: Game Technology, Information System, Computer Network, Computing, Electronics, and Control*, Jun. 2025, doi: 10.22219/kinetik.v10i3.2221.
- [9] A. A. Chamid, Widowati, and R. Kusumaningrum, "Graph-Based Semi-Supervised Deep Learning for Indonesian Aspect-Based Sentiment Analysis," *Big Data and Cognitive Computing*, vol. 7, no. 1, Mar. 2023, doi: 10.3390/bdcc7010005.
- [10] A. P. Hendrawan, E. Wijayanti, and A. A. Chamid, "Design of an Exam Cheating Detection System Application Based on Machine Learning with the Computer Vision Method," *Jurnal Teknologi Informatika dan Komputer*, vol. 11, no. 2, pp. 509–521, Jul. 2025, doi: 10.37012/jtik.v11i2.2704.
- [11] A. A. Chamid, Widowati, and R. Kusumaningrum, "Labeling Consistency Test of

- Multi-Label Data for Aspect and Sentiment Classification Using the Cohen Kappa Method,” *Ingenierie des Systemes d’Information*, vol. 29, no. 1, pp. 161–167, Feb. 2024, doi: 10.18280/isi.290118.
- [12] A. Gramegna and P. Giudici, “SHAP and LIME: An Evaluation of Discriminative Power in Credit Risk,” *Front Artif Intell*, vol. 4, 2021, doi: 10.3389/frai.2021.752558.
- [13] M. Mersha, K. Lam, J. Wood, A. AlShami, and J. Kalita, “Explainable Artificial Intelligence: A Survey of Needs, Techniques, Applications, and Future Direction,” Jan. 2025, doi: 10.1016/j.neucom.2024.128111.
- [14] I. H. Ariyanto, A. A. Chamid, and R. Fiati, “Demand Prediction and Apparel Production Management Using AI-Based Decision Tree,” *bit-Tech*, vol. 8, no. 1, pp. 59–67, Aug. 2025, doi: 10.32877/bt.v8i1.2325.
- [15] M. Kristanaya, Melinda Putri Azzahra, Trimono, and Mohammad Idhom, “Klasifikasi Status Rujukan Pasien Poliklinik Bandara Berbasis Random Forest dan Interpretabilitas Model Menggunakan SHAP,” *Data Sciences Indonesia (DSI)*, vol. 5, no. 1, pp. 60–74, Jul. 2025, doi: 10.47709/dsi.v5i1.6078.
- [16] O. A. Bello, “Citation: Bello O.A. (2023) Machine Learning Algorithms for Credit Risk Assessment: An Economic and Financial Analysis,” *International Journal of Management Technology*, vol. 10, no. 1, pp. 109–133, 2023, doi: 10.37745/ijmt.2013/vol10n1109133.
- [17] B. Babolcsay, “Shapley Values and Explainable Artificial Intelligence,” Master’s Thesis, Eötvös Loránd University, 2025.
- [18] N. B. Putri and A. W. Wijayanto, “Analisis Komparasi Algoritma Klasifikasi Data Mining Dalam Klasifikasi Website Phishing,” *Komputika : Jurnal Sistem Komputer*, vol. 11, no. 1, pp. 59–66, Jan. 2022, doi: 10.34010/komputika.v11i1.4350.
- [19] M. T. Hossain, J. Lee, D. Besenski, B. Dimitrijevic, and L. Spasovic, “Developing a Prediction Model for Real-Time Incident Detection Leveraging User-Oriented Participatory Sensing Data,” *Information (Switzerland)*, vol. 16, no. 6, Jun. 2025, doi: 10.3390/info16060423.
- [20] A. Fadlil, Herman, and M. Dikky Praseptian, “Single Imputation Using Statistics-Based and K Nearest Neighbor Methods for Numerical Datasets,” *Ingenierie des Systemes d’Information*, vol. 28, no. 2, pp. 451–459, Apr. 2023, doi: 10.18280/isi.280221.
- [21] A. Seveso, A. Campagner, D. Ciucci, and F. Cabitza, “Ordinal labels in machine learning: A user-centered approach to improve data validity in medical settings,” *BMC Med Inform Decis Mak*, vol. 20, Aug. 2020, doi: 10.1186/s12911-020-01152-8.
- [22] F. Bolikulov, R. Nasimov, A. Rashidov, F. Akhmedov, and Y. I. Cho, “Effective Methods of Categorical Data Encoding for Artificial Intelligence Algorithms,” *Mathematics*, vol. 12, no. 16, Aug. 2024, doi: 10.3390/math12162553.
- [23] L. B. V. de Amorim, G. D. C. Cavalcanti, and R. M. O. Cruz, “The choice of scaling technique matters for classification performance,” Dec. 2022, doi: 10.1016/j.asoc.2022.109924.
- [24] E. H. Yulianti, O. Soesanto, and Y. Sukmawaty, “Penerapan Metode Extreme Gradient Boosting (XGBOOST) pada Klasifikasi Nasabah Kartu Kredit,” *JOMTA Journal of Mathematics: Theory and Applications*, vol. 4, no. 1, 2022.
- [25] M. T. Syamkalla, S. Khomsah, and Y. S. R. Nur, “Implementasi Algoritma Catboost Dan Shapley Additive Explanations (SHAP) Dalam Memprediksi Popularitas Game Indie Pada Platform Steam,” *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 11, no. 4, pp. 777–786, Aug. 2024, doi: 10.25126/jtiik.1148503.
- [26] A. Damuri, U. Riyanto, H. Rusdianto, and M. Aminudin, “Implementasi Data Mining dengan Algoritma Naïve Bayes Untuk Klasifikasi Kelayakan Penerima Bantuan Sembako,” *JURIKOM (Jurnal Riset Komputer)*, vol. 8, no. 6, p. 219, Dec. 2021, doi: 10.30865/jurikom.v8i6.3655.
- [27] F. Damayanti, Arief Budiman, Siti Sundari, and Theodora MV Nainggolan, “Classification of Customer Credit Risk Levels Using the Random Forest Method: A Case Study on Microfinance Institutions,” *Journal of Computer Science Artificial Intelligence and Communications*, vol. 1, no. 2, pp. 52–56, Nov. 2024, doi: 10.62712/jocsaic.v1i2.20.
- [28] D. Galih Pradana, M. L. Alghifari, M. Farhan Juna, and S. Dwisiwi Palaguna, “Klasifikasi Penyakit Jantung Menggunakan Metode Artificial Neural Network,” *Indonesian Journal of Data and Science (IJODAS)*, vol. 3, no. 2, pp. 55–60, 2022.