

Explainable Machine Learning untuk Prediksi Risiko Kredit sebagai Pendukung Pengambilan Keputusan Manajemen Risiko

Jamilatul Badriyah^{1,*}, Agung Muliawan²

¹ Informatika; Universitas Madura; Jl Raya Panglegur Km. 3,5, Pamekasan, Jawa Timur, fax(0324)327418; e-mail: mila@unira.ac.id.

² Manajemen Informatika; Politeknik Negeri Jember; Jember, Jawa Timur, fax (0331) 333532; e-mail: agung.muliawan@polije.ac.id

* Korespondensi: e-mail: mia@unira.ac.id

Abstrak: Risiko kredit merupakan salah satu risiko utama yang dihadapi lembaga keuangan akibat potensi gagal bayar (default) nasabah. Meskipun berbagai metode machine learning telah digunakan untuk meningkatkan akurasi prediksi risiko kredit, sebagian besar penelitian masih berfokus pada performa model sehingga aspek interpretabilitas yang dibutuhkan dalam pengambilan keputusan manajemen risiko belum banyak diperhatikan. Penelitian ini bertujuan mengembangkan model Explainable Machine Learning dengan mengombinasikan algoritma Extreme Gradient Boosting (XGBoost) dan SHapley Additive exPlanations (SHAP). XGBoost dipilih karena memiliki kemampuan yang baik dalam memodelkan hubungan nonlinier dan menghasilkan akurasi prediksi yang tinggi, sedangkan SHAP digunakan untuk memberikan interpretasi terhadap kontribusi setiap fitur sehingga keputusan model menjadi lebih transparan. Penelitian menggunakan Taiwan Credit Card Clients Dataset dengan tahapan data preprocessing, penanganan ketidakseimbangan data menggunakan Synthetic Minority Over-sampling Technique (SMOTE), optimasi hiperparameter menggunakan RandomizedSearchCV dengan 3-fold cross-validation, serta evaluasi menggunakan metrik Accuracy, Precision, Recall, F1-Score, dan ROC-AUC. Hasil penelitian menunjukkan bahwa model memperoleh Accuracy sebesar 80,65%, Precision 59,39%, Recall 39,56%, F1-Score 47,49%, dan ROC-AUC 75,54%, yang menunjukkan kemampuan klasifikasi dan diskriminasi yang cukup baik. Interpretasi menggunakan SHAP mengidentifikasi bahwa PAY_0, LIMIT_BAL, AGE, serta riwayat dan jumlah pembayaran merupakan faktor-faktor yang paling berpengaruh terhadap prediksi risiko kredit. Kontribusi penelitian ini adalah menunjukkan bahwa integrasi XGBoost dan SHAP tidak hanya menghasilkan model dengan performa prediksi yang baik, tetapi juga meningkatkan transparansi dan kemudahan interpretasi, sehingga lebih mendukung proses pengambilan keputusan manajemen risiko kredit yang berbasis data..

Kata kunci: Risiko Kredit, Explainable Machine Learning, XGBoost, SHAP, Manajemen Risiko.

Abstract: Credit risk is one of the primary risks faced by financial institutions due to the potential for customer default. Although various machine learning methods have been employed to improve the accuracy of credit risk prediction, most existing studies primarily focus on predictive performance while paying limited attention to model interpretability, which is essential for risk management decision-making. This study proposes an Explainable Machine Learning approach by integrating Extreme Gradient Boosting (XGBoost) and SHapley Additive exPlanations (SHAP). XGBoost was selected for its capability to model complex nonlinear relationships and achieve high predictive performance, while SHAP was employed to provide transparent explanations of each feature's contribution to the prediction outcomes. The proposed approach was evaluated using the Taiwan Credit Card Clients Dataset through data preprocessing, class imbalance handling with the Synthetic Minority Over-sampling Technique (SMOTE), hyperparameter optimization using RandomizedSearchCV with 3-fold cross-validation, and model performance assessment based on Accuracy, Precision, Recall, F1-Score, and ROC-AUC. Experimental results demonstrate that the proposed model achieved an

Accuracy of 80.65%, Precision of 59.39%, Recall of 39.56%, F1-Score of 47.49%, and ROC-AUC of 75.54%, indicating satisfactory classification and discrimination capabilities in distinguishing default and non-default customers. SHAP analysis revealed that PAY_0, LIMIT_BAL, AGE, together with customers' payment history and payment amounts, were the most influential factors in credit risk prediction. The main contribution of this study is demonstrating that the integration of XGBoost and SHAP not only provides accurate credit risk prediction but also enhances model transparency and interpretability, thereby supporting more reliable and data-driven credit risk management decisions..

Keywords: Credit Risk, Explainable Machine Learning, XGBoost, SHAP, Risk Management.

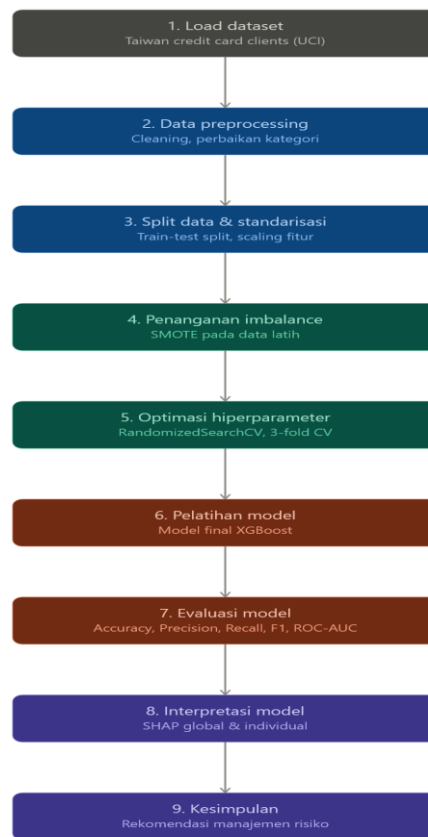
1. Pendahuluan

Risiko kredit merupakan salah satu risiko utama yang dihadapi lembaga keuangan akibat potensi gagal bayar (default) nasabah [1]. Kemampuan memprediksi risiko kredit secara akurat sangat penting untuk mendukung pengambilan keputusan dalam proses pemberian kredit, meminimalkan kerugian, serta menjaga stabilitas keuangan [2], [3]. Seiring meningkatnya volume dan kompleksitas data, pendekatan berbasis machine learning semakin banyak dimanfaatkan karena mampu menghasilkan prediksi yang lebih akurat dibandingkan metode konvensional [4]. Berbagai algoritma machine learning, seperti Logistic Regression, Random Forest, Support Vector Machine, dan Extreme Gradient Boosting (XGBoost), telah banyak digunakan dalam prediksi risiko kredit [5]. XGBoost dikenal memiliki performa yang baik dalam menangani data berdimensi tinggi dan hubungan nonlinier antarvariabel [6]. Namun, sebagian besar penelitian masih berfokus pada peningkatan akurasi model tanpa memberikan penjelasan mengenai faktor-faktor yang memengaruhi hasil prediksi [7], [8]. Kondisi tersebut menyebabkan model sulit diinterpretasikan sehingga mengurangi tingkat kepercayaan pengguna dalam pengambilan keputusan [9], [10]. Untuk mengatasi permasalahan tersebut, penelitian ini menerapkan pendekatan Explainable Machine Learning dengan mengombinasikan algoritma XGBoost dan metode SHapley Additive exPlanations (SHAP) [11]. Sebelum proses pelatihan model, ketidakseimbangan data ditangani menggunakan Synthetic Minority Over-sampling Technique (SMOTE), sedangkan optimasi hiperparameter dilakukan menggunakan RandomizedSearchCV untuk memperoleh konfigurasi model yang optimal. Pendekatan ini tidak hanya bertujuan meningkatkan performa prediksi, tetapi juga memberikan interpretasi terhadap kontribusi setiap fitur dalam proses klasifikasi.

Penelitian ini menggunakan Taiwan Credit Card Clients Dataset dari UCI Machine Learning Repository untuk membangun model prediksi risiko kredit. Kinerja model dievaluasi menggunakan metrik Accuracy, Precision, Recall, F1-Score, dan ROC-AUC, sedangkan interpretasi model dilakukan menggunakan SHAP. Hasil penelitian diharapkan dapat menghasilkan model prediksi yang akurat sekaligus transparan sehingga dapat menjadi pendukung pengambilan keputusan dalam manajemen risiko kredit di lembaga keuangan [12].

2. Metode Penelitian

Penelitian ini menggunakan pendekatan kuantitatif dengan metode *machine learning* untuk membangun model prediksi risiko kredit yang dapat diinterpretasikan (*Explainable Machine Learning*). Tahapan penelitian dimulai dari pengumpulan data, *preprocessing*, pelatihan model, evaluasi performa, hingga interpretasi hasil prediksi menggunakan SHAP. Alur penelitian ditunjukkan pada Gambar 1.



Gambar 1. Alur Penelitian

2.1 Dataset

Dataset yang digunakan adalah Taiwan Credit Card Clients Dataset yang diperoleh dari UCI Machine Learning Repository. Dataset ini berisi data historis nasabah kartu kredit yang terdiri dari informasi demografis, limit kredit, histori pembayaran, tagihan bulanan, jumlah pembayaran, serta status gagal bayar (default payment) sebagai variabel target. Berikut merupakan tabel dataset pada penelitian ini :

Tabel 1. Dataset

No	Nama Kolom	Tipe Data	Keterangan
1	ID	Numerik (Integer)	Nomor identitas unik nasabah (dihapus pada tahap <i>preprocessing</i>).
2	LIMIT_BAL	Numerik (Float)	Jumlah limit kredit yang diberikan kepada nasabah (Dollar NT).
3	SEX	Kategorik (Integer)	Jenis kelamin nasabah (1 = Laki-laki, 2 = Perempuan).
4	EDUCATION	Kategorik (Integer)	Tingkat pendidikan (1 = Graduate School, 2 = University, 3 = High School, 4 = Others).
5	MARRIAGE	Kategorik (Integer)	Status pernikahan (1 = Menikah, 2 = Lajang, 3 = Lainnya).
6	AGE	Numerik (Integer)	Usia nasabah (tahun).
7	PAY_0	Kategorik (Integer)	Status pembayaran bulan September 2005 (-1 = tepat waktu, 1–9 = keterlambatan pembayaran dalam bulan).
8	PAY_2	Kategorik (Integer)	Status pembayaran bulan Agustus 2005.
9	PAY_3	Kategorik (Integer)	Status pembayaran bulan Juli 2005.
10	PAY_4	Kategorik (Integer)	Status pembayaran bulan Juni 2005.
11	PAY_5	Kategorik (Integer)	Status pembayaran bulan Mei 2005.
12	PAY_6	Kategorik (Integer)	Status pembayaran bulan April 2005.
13	BILL_AMT1	Numerik (Float)	Jumlah tagihan pada bulan September 2005.
14	BILL_AMT2	Numerik (Float)	Jumlah tagihan pada bulan Agustus 2005.
15	BILL_AMT3	Numerik (Float)	Jumlah tagihan pada bulan Juli 2005.
16	BILL_AMT4	Numerik (Float)	Jumlah tagihan pada bulan Juni 2005.
17	BILL_AMT5	Numerik (Float)	Jumlah tagihan pada bulan Mei 2005.
18	BILL_AMT6	Numerik (Float)	Jumlah tagihan pada bulan April 2005.
19	PAY_AMT1	Numerik (Float)	Jumlah pembayaran pada bulan September 2005.
20	PAY_AMT2	Numerik (Float)	Jumlah pembayaran pada bulan Agustus 2005.
21	PAY_AMT3	Numerik (Float)	Jumlah pembayaran pada bulan Juli 2005.

22	PAY_AMT4	Numerik (Float)	Jumlah pembayaran pada bulan Juni 2005.
23	PAY_AMT5	Numerik (Float)	Jumlah pembayaran pada bulan Mei 2005.
24	PAY_AMT6	Numerik (Float)	Jumlah pembayaran pada bulan April 2005.
25	DEFAULT_PAYMENT	Kategorik (Integer)	Variabel target (0 = Tidak default, 1 = Default).

Sumber : Taiwan Credit Card Clients

2.2 Data Preprocessing

Tahap preprocessing bertujuan meningkatkan kualitas data sebelum digunakan dalam proses pelatihan model. Proses yang dilakukan meliputi penghapusan atribut ID karena tidak memiliki pengaruh terhadap proses klasifikasi, perbaikan nilai kategori yang tidak valid pada atribut EDUCATION dan MARRIAGE, serta perubahan nama variabel target menjadi DEFAULT_PAYMENT agar lebih mudah diinterpretasikan [13]. Selanjutnya, data dipisahkan menjadi data latih (training set) sebesar 80% dan data uji (testing set) sebesar 20% menggunakan metode stratified sampling untuk mempertahankan proporsi kelas. Fitur numerik kemudian distandarisasi menggunakan StandardScaler .

2.3 Penanganan Ketidakseimbangan Data dan Pelatihan Model

Permasalahan ketidakseimbangan kelas ditangani menggunakan Synthetic Minority Over-sampling Technique (SMOTE) yang diterapkan hanya pada data latih untuk menghindari data leakage. Selanjutnya, optimasi hiperparameter dilakukan menggunakan Randomized SearchCV

dengan 3-fold cross-validation dan metrik evaluasi F1-Score. Model klasifikasi yang digunakan adalah Extreme Gradient Boosting (XGBoost) dengan parameter terbaik hasil optimasi sebagai model akhir.

a. SMOTE (Synthetic Minority Over-sampling Technique)

SMOTE bekerja dengan membuat data sintesis baru dari kelas minoritas melalui interpolasi antara data asli dengan tetangga terdekatnya (*k-nearest neighbors*).

Berikut merupakan formulanya :

$$x_{new} = x_i + \lambda \times (x_{zi} - x_i) \quad \dots\dots\dots (1)$$

Keterangan:

x_i = data kelas minoritas yang dipilih secara acak

x_{zi} = salah satu dari k tetangga terdekat (*k-nearest neighbor*) dari x_i (biasanya $k=5$)

λ = bilangan acak antara 0 dan 1, $\lambda \in [0,1]$

x_{new} = data sintesis baru hasil interpolasi

b. RandomizedSearchCV

RandomizedSearchCV tidak memiliki "rumus" tunggal, tapi mekanismenya berbasis sampling acak dari ruang hiperparameter dan evaluasi menggunakan cross-validation. Berikut merupakan formula Skor evaluasi (rata-rata cross-validation):

$$Score(\theta) = \frac{1}{K} \sum_{k=1}^K F1_k(\theta) \quad \dots\dots\dots (2)$$

Keterangan :

K = jumlah fold cross-validation (pada kode Anda, $K = 3$)

$F1_k(\theta)$ = nilai F1-score pada fold ke- k dengan kombinasi hiperparameter θ

c. XGBoost (Extreme Gradient Boosting)

XGBoost membangun model secara bertahap (*additive*), dimana setiap pohon baru berusaha memperbaiki kesalahan (residual) dari pohon-pohon sebelumnya [14]. Berikut merupakan formula Fungsi prediksi (additive model) [15]:

$$\hat{y}_i = \sum_{t=1}^T f_t(x_i), \quad f_t \in \mathcal{F} \quad \dots\dots\dots (3)$$

Keterangan:

$\hat{y}_i$ = prediksi akhir untuk sampel ke- i

T = jumlah pohon (n_estimators)

f_t = fungsi pohon ke- t

$\mathcal{F}$ = ruang seluruh kemungkinan pohon (CART)

2.4 Evaluasi dan Interpretasi Model

Performa model dievaluasi menggunakan metrik Accuracy, Precision, Recall, F1-Score, dan Receiver Operating Characteristic – Area Under Curve (ROC-AUC). Selain itu, hasil klasifikasi divisualisasikan menggunakan Confusion Matrix dan ROC Curve untuk mengetahui kemampuan model dalam membedakan kelas default dan non-default [16]. Untuk meningkatkan transparansi model, penelitian ini menerapkan SHapley Additive exPlanations (SHAP) yang digunakan untuk mengidentifikasi kontribusi setiap fitur terhadap hasil prediksi, baik secara global (global feature importance) maupun pada tingkat individu (local explanation) [11].

3. Hasil dan Pembahasan

3.1 Hasil Preprocessing Data

Tahap preprocessing dilakukan untuk meningkatkan kualitas data sebelum proses pelatihan model. Pada tahap ini, atribut ID dihapus karena hanya berfungsi sebagai identitas unik sehingga tidak memiliki kontribusi terhadap proses klasifikasi. Selanjutnya dilakukan perbaikan nilai kategori pada atribut EDUCATION dan MARRIAGE untuk menghilangkan kategori yang tidak valid. Nama variabel target default payment next month juga diubah menjadi DEFAULT_PAYMENT agar lebih mudah diinterpretasikan selama proses analisis. Dataset kemudian dipisahkan menjadi data latih sebesar 80% dan data uji sebesar 20% menggunakan metode stratified sampling sehingga proporsi kelas tetap terjaga. Selanjutnya, fitur numerik distandarisasi menggunakan StandardScaler, sedangkan ketidakseimbangan kelas pada data latih ditangani menggunakan SMOTE sebelum dilakukan proses pelatihan model.

3.2 Hasil Optimasi Hiperparameter

Optimasi hiperparameter dilakukan menggunakan RandomizedSearchCV dengan 3-fold cross-validation dan metrik evaluasi F1-Score. Pendekatan ini bertujuan memperoleh kombinasi parameter yang mampu meningkatkan performa model XGBoost tanpa melakukan pencarian secara menyeluruh terhadap seluruh ruang parameter. Parameter terbaik yang diperoleh kemudian digunakan pada proses pelatihan model akhir.

Tabel 2. Parameter Terbaik Hasil RandomizedSearchCV

No	Parameter	Nilai
1	max_depth	9
2	learning_rate	0.1
3	n_estimators	200
4	subsample	1.0
5	colsample_bytree	1.0

Sumber : Penelitian google colab

Optimasi hiperparameter dilakukan menggunakan RandomizedSearchCV dengan 3-fold cross-validation dan metrik evaluasi F1-Score. Proses optimasi difokuskan pada lima parameter utama XGBoost, yaitu max_depth, learning_rate, n_estimators, subsample, dan colsample_bytree. Berdasarkan hasil optimasi, kombinasi parameter terbaik diperoleh dengan max_depth = 9, learning_rate = 0.1, n_estimators = 200, subsample = 1.0, dan colsample_bytree = 1.0. Parameter tersebut kemudian digunakan untuk membangun model akhir yang selanjutnya dievaluasi menggunakan data uji.

3.3 Hasil Evaluasi Model

Kinerja model dievaluasi menggunakan lima metrik, yaitu Accuracy, Precision, Recall, F1-Score, dan ROC-AUC. Hasil evaluasi disajikan pada Tabel 3.

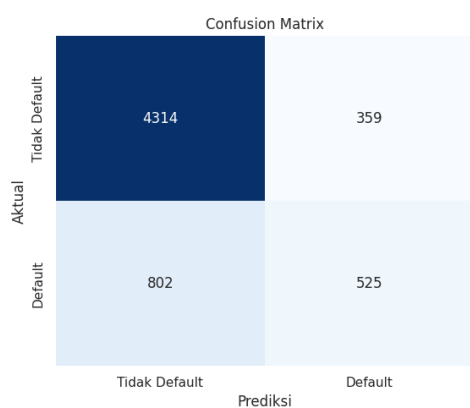
Tabel 3. Hasil Evaluasi Model XGBoost

No	Metrik	Nilai
1	Accuracy	80,65%
2	Precision	59,39%
3	Recall	39,56%
4	F1-Score	47,49%

5	ROC-AUC	75,54%
---	---------	--------

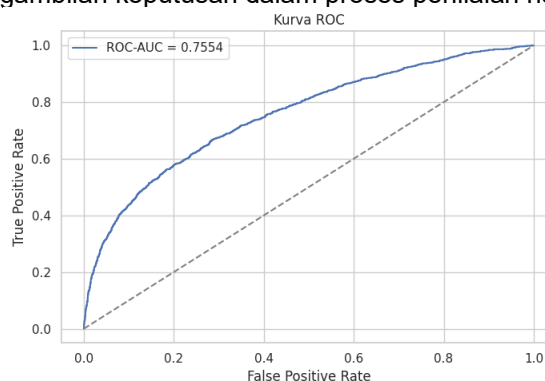
Sumber : Penelitian google colab

Berdasarkan hasil evaluasi, model menghasilkan Accuracy sebesar 80,65%, yang menunjukkan bahwa sebagian besar data berhasil diklasifikasikan dengan benar. Nilai Precision sebesar 59,39% mengindikasikan bahwa lebih dari separuh nasabah yang diprediksi mengalami default benar-benar termasuk dalam kategori tersebut. Sementara itu, Recall sebesar 39,56% menunjukkan bahwa model masih memiliki keterbatasan dalam mengidentifikasi seluruh nasabah yang benar-benar mengalami gagal bayar. Nilai F1-Score sebesar 47,49% mencerminkan keseimbangan antara Precision dan Recall, sedangkan ROC-AUC sebesar 75,54% menunjukkan bahwa model memiliki kemampuan diskriminasi yang cukup baik dalam membedakan nasabah default dan non-default. Selain metrik evaluasi, performa model divisualisasikan menggunakan Confusion Matrix dan ROC Curve.



Gambar 2. Confusion Matrix

Berdasarkan Confusion Matrix pada Gambar 2, model XGBoost berhasil mengklasifikasikan 4.314 data sebagai tidak default dengan benar (True Negative) dan 525 data sebagai default dengan benar (True Positive). Di sisi lain, terdapat 359 data yang sebenarnya tidak default tetapi diprediksi sebagai default (False Positive), serta 802 data yang sebenarnya default namun diprediksi sebagai tidak default (False Negative). Hasil ini menunjukkan bahwa model memiliki kemampuan yang baik dalam mengidentifikasi nasabah yang tidak mengalami gagal bayar, namun masih mengalami kesulitan dalam mendeteksi seluruh nasabah yang benar-benar berisiko default. Jumlah False Negative yang lebih tinggi dibandingkan False Positive sejalan dengan nilai Recall sebesar 39,56%, yang mengindikasikan bahwa masih terdapat cukup banyak kasus gagal bayar yang belum berhasil dikenali oleh model. Meskipun demikian, dominasi prediksi yang benar pada kedua kelas menunjukkan bahwa model memiliki performa klasifikasi yang cukup baik dan dapat digunakan sebagai pendukung pengambilan keputusan dalam proses penilaian risiko kredit.



Gambar 3. ROC Curve

Berdasarkan Kurva ROC pada Gambar 3, model XGBoost memperoleh nilai Area Under the Curve (ROC-AUC) sebesar 0,7554. Nilai tersebut menunjukkan bahwa model

memiliki kemampuan yang cukup baik dalam membedakan antara nasabah yang berpotensi mengalami default dan yang tidak mengalami default. Kurva ROC yang berada di atas garis diagonal (garis acak) mengindikasikan bahwa performa model jauh lebih baik dibandingkan klasifikasi secara acak. Semakin mendekati sudut kiri atas, semakin baik kemampuan model dalam mencapai True Positive Rate yang tinggi dengan False Positive Rate yang rendah. Dengan nilai ROC-AUC sebesar 75,54%, model memiliki kemampuan diskriminasi yang baik (good discrimination) dalam mengidentifikasi risiko kredit, sehingga dapat dijadikan sebagai alat pendukung dalam proses penilaian kelayakan kredit. Namun demikian, masih terdapat ruang untuk meningkatkan performa model, terutama dalam meningkatkan kemampuan mendeteksi nasabah yang benar-benar berisiko default sebagaimana tercermin dari nilai Recall yang relatif lebih rendah dibandingkan metrik evaluasi lainnya.

3.4 Interpretasi Model Menggunakan SHAP

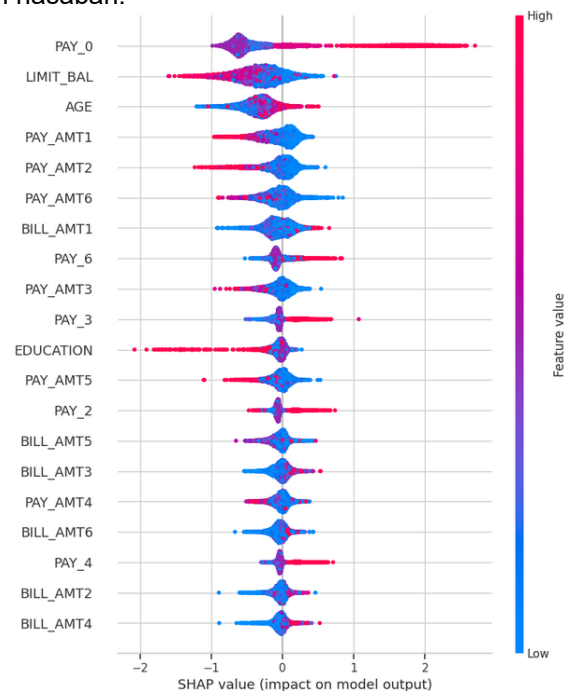
Untuk meningkatkan transparansi model, penelitian ini menerapkan SHapley Additive exPlanations (SHAP). Metode ini digunakan untuk mengidentifikasi kontribusi masing-masing fitur terhadap hasil prediksi yang dihasilkan oleh model XGBoost. Berdasarkan nilai rata-rata absolut SHAP, diperoleh sepuluh fitur yang memiliki pengaruh terbesar terhadap prediksi risiko kredit sebagaimana disajikan pada Tabel 4.

Tabel 4. Lima Fitur Paling Berpengaruh Berdasarkan SHAP

Peringkat	Fitur	Mean SHAP
1	PAY_0	0,639069
2	LIMIT_BAL	0,361456
3	AGE	0,332693
4	PAY_AMT1	0,168652
5	PAY_AMT2	0,146528

Sumber : Penelitian google colab

Hasil tersebut menunjukkan bahwa PAY_0 merupakan variabel yang paling berpengaruh dalam menentukan risiko kredit. Temuan ini mengindikasikan bahwa status pembayaran pada periode terakhir menjadi indikator utama dalam memprediksi kemungkinan gagal bayar. Selain itu, LIMIT_BAL dan AGE juga memberikan kontribusi yang signifikan terhadap keputusan model, diikuti oleh variabel yang berkaitan dengan riwayat pembayaran dan jumlah pembayaran nasabah.



Gambar 4. SHAP Summary Plot

Berdasarkan SHAP Summary Plot pada Gambar 4, diketahui bahwa PAY_0 merupakan fitur yang memberikan pengaruh paling besar terhadap prediksi risiko kredit, diikuti oleh LIMIT_BAL, AGE, PAY_AMT1, PAY_AMT2, dan PAY_AMT6. Urutan fitur pada sumbu vertikal menunjukkan tingkat kepentingan relatif berdasarkan nilai rata-rata absolut SHAP (mean absolute SHAP value), sedangkan sebaran titik pada sumbu horizontal menggambarkan besar dan arah kontribusi masing-masing fitur terhadap hasil prediksi model. Warna merah menunjukkan nilai fitur yang tinggi, sedangkan warna biru menunjukkan nilai fitur yang rendah. Pada fitur PAY_0, nilai yang lebih tinggi (status keterlambatan pembayaran) cenderung menghasilkan nilai SHAP positif sehingga meningkatkan probabilitas nasabah diklasifikasikan sebagai default. Sebaliknya, nilai LIMIT_BAL yang lebih tinggi cenderung menghasilkan nilai SHAP negatif, yang menunjukkan bahwa nasabah dengan limit kredit lebih besar memiliki kecenderungan risiko gagal bayar yang lebih rendah. Selain itu, variabel AGE serta beberapa variabel yang berkaitan dengan jumlah pembayaran (PAY_AMT1, PAY_AMT2, dan PAY_AMT6) juga memberikan kontribusi yang signifikan terhadap keputusan model. Hasil ini menunjukkan bahwa model tidak hanya mengandalkan satu variabel, tetapi memanfaatkan kombinasi riwayat pembayaran, kapasitas kredit, karakteristik nasabah, dan perilaku pembayaran untuk memprediksi risiko kredit secara lebih komprehensif. Interpretasi yang diberikan oleh SHAP meningkatkan transparansi model sehingga hasil prediksi dapat lebih mudah dipahami dan dijadikan dasar dalam pengambilan keputusan manajemen risiko kredit.

3.5 Pembahasan.

Hasil penelitian menunjukkan bahwa kombinasi SMOTE, RandomizedSearchCV, XGBoost, dan SHAP mampu menghasilkan model prediksi risiko kredit dengan performa yang baik serta interpretasi yang transparan. Nilai Accuracy sebesar 80,65% dan ROC-AUC sebesar 75,54% menunjukkan bahwa model memiliki kemampuan yang memadai dalam membedakan nasabah yang berpotensi mengalami default dan non-default. Di sisi lain, nilai Recall sebesar 39,56% mengindikasikan bahwa masih terdapat ruang untuk meningkatkan kemampuan model dalam mendeteksi seluruh kasus gagal bayar. Interpretasi menggunakan SHAP menunjukkan bahwa status pembayaran terakhir (PAY_0) merupakan faktor yang paling menentukan dalam prediksi risiko kredit. Selain itu, LIMIT_BAL, AGE, serta riwayat dan jumlah pembayaran juga memberikan kontribusi yang signifikan. Temuan ini sejalan dengan konsep manajemen risiko kredit, di mana perilaku pembayaran historis dan kapasitas finansial nasabah merupakan indikator utama dalam menilai tingkat risiko. Dengan demikian, penerapan Explainable Machine Learning tidak hanya menghasilkan model prediksi yang memiliki performa memadai, tetapi juga menyediakan informasi yang mudah dipahami sehingga dapat mendukung proses pengambilan keputusan yang lebih transparan dan berbasis data.

4. Kesimpulan

Penelitian ini berhasil mengembangkan model Explainable Machine Learning untuk prediksi risiko kredit menggunakan algoritma Extreme Gradient Boosting (XGBoost) yang dioptimasi dengan RandomizedSearchCV serta didukung oleh Synthetic Minority Over-sampling Technique (SMOTE) untuk mengatasi ketidakseimbangan kelas. Berdasarkan hasil pengujian, model memperoleh nilai Accuracy sebesar 80,65%, Precision sebesar 59,39%, Recall sebesar 39,56%, F1-Score sebesar 47,49%, dan ROC-AUC sebesar 75,54%, yang menunjukkan bahwa model memiliki kemampuan yang cukup baik dalam membedakan nasabah yang berpotensi mengalami default dan non-default. Penerapan metode SHapley Additive exPlanations (SHAP) berhasil meningkatkan transparansi model dengan mengidentifikasi faktor-faktor yang paling berpengaruh terhadap prediksi risiko kredit. Fitur PAY_0, LIMIT_BAL, dan AGE merupakan variabel yang memberikan kontribusi terbesar terhadap keputusan model, sehingga informasi tersebut dapat menjadi dasar dalam proses analisis dan pengambilan keputusan manajemen risiko kredit. Dengan demikian, pendekatan Explainable Machine Learning tidak hanya menghasilkan model prediksi yang memiliki performa yang baik, tetapi juga menyediakan interpretasi yang mudah dipahami sehingga dapat meningkatkan kepercayaan pengguna dalam penerapan model pada lembaga keuangan. Sebagai pengembangan pada penelitian selanjutnya, dapat dilakukan perbandingan dengan algoritma machine learning lainnya, penerapan teknik penanganan ketidakseimbangan data yang berbeda, maupun eksplorasi metode Explainable Artificial Intelligence (XAI) lainnya untuk meningkatkan kemampuan prediksi serta kualitas interpretasi model.

Referensi

- [1] S. Kaisar and S. Sifat, "Explainable Machine Learning Models for Credit Risk Analysis: A Survey," 2023, pp. 51–72. doi: 10.1007/978-3-031-36570-6_2.
- [2] S. Crone and S. Finlay, "Instance sampling in credit scoring: An empirical study of sample size and balancing," *International Journal of Forecasting*, vol. 28, pp. 224–238, Mar. 2012, doi: 10.1016/j.ijforecast.2011.07.006.
- [3] A. Wahid and A. Muliawan, "Strategi Retensi Pelanggan Berbasis Historis: Optimalisasi Model Prediksi Churn Menggunakan Machine Learning," *SemanTIK: Teknik Informasi*, vol. 11, no. 2, Sep. 2025, doi: 10.55679/semantik.v11i2.237.
- [4] M. Pebriadi, T. Fattah, and P. Salman, "Interpreting Credit Risk Prediction with XGBoost and SHAP," 2025, p. 582. doi: 10.1109/ISRIT168345.2025.11393290.
- [5] D. Kumar, "Explainable Machine Learning Models for Credit Risk Prediction in Retail Lending: A Comparative Study Using SHAP," *SSRN Electronic Journal*, Jan. 2025, doi: 10.2139/ssrn.5341125.
- [6] B. Hadji Misheva, J. Osterrieder, A. Hirsä, O. Kulkarni, and S. Lin, "Explainable AI in Credit Risk Management," 2021. doi: 10.48550/arXiv.2103.00949.
- [7] N. Bussmann, P. Giudici, D. Marinelli, and J. Papenbrock, "Explainable Machine Learning in Credit Risk Management," *Computational Economics*, vol. 57, Jan. 2021, doi: 10.1007/s10614-020-10042-0.
- [8] D. A. Fauziah, A. Muliawan, and M. Dimiyati, "Implementation Of Machine Learning On Employee Attrition Based On Performance Parameters Using Particle Swarm Optimization And Ensemble Classifier Methods," *Jurnal Teknik Informatika (Jutif)*, vol. 5, no. 6, Art. no. 6, Dec. 2024, doi: 10.52436/1.jutif.2024.5.6.3442.
- [9] L. Shu, C. Su, and J. Shu, "A Comparative Study of Explainable Machine Learning Models for Corporate Credit Scoring," 2025, p. 6. doi: 10.1109/ICECET63943.2025.11472005.
- [10] D. A. Fauziah, A. Muliawan, I. Sabilirasyad, and B. W. Maulana, "Smart Campaign for Smart Business Using Machine Learning on Improving Product Marketing Strategy," *Proceeding International Conference on Economics, Business and Information Technology*, vol. 6, no. 1, pp. 616–622, Sep. 2025, doi: 10.31967/icebit.v6i1.1681.
- [11] S. S and S. M., "Smart and Explainable Credit Card Fraud Detection Using XGBoost and SHAP," *J. IoT Soc. Mob. Anal. Cloud*, vol. 7, no. 2, pp. 155–169, Jul. 2025, doi: 10.36548/jismac.2025.2.004.
- [12] Shreya and H. Pathak, "Explainable Artificial Intelligence Credit Risk Assessment using Machine Learning," 2025. doi: 10.48550/arXiv.2506.19383.
- [13] H. Wicaksono, A. Riyanto, R. Darmawan, M. F. Hidayat, and A. Khumaidi, "Explainable Boosting Machine for Transparent Risk Assessment in BAZNAS Microfinance Desa," *ILKOM Jurnal Ilmiah*, vol. 17, no. 3, pp. 312–322, Dec. 2025, doi: 10.33096/ilkom.v17i3.3214.312-322.
- [14] H. Putra and R. Rumini, "Comparative Study of Logistic Regression, Random Forest, and XGBoost for Bank Loan Approval Classification," *Journal of Applied Informatics and Computing*, vol. 9, no. 5, pp. 2822–2835, Oct. 2025, doi: 10.30871/jaic.v9i5.10862.
- [15] S. Lin, D. Song, B. Cao, X. Gu, and J. Li, "Credit risk assessment of automobile loans using machine learning-based SHapley Additive exPlanations approach," *Engineering Applications of Artificial Intelligence*, vol. 147, p. 110236, May 2025, doi: 10.1016/j.engappai.2025.110236.
- [16] L. Lin and Y. Wang, "SHAP Stability in Credit Risk Management: A Case Study in Credit Card Default Model," 2025. doi: 10.48550/arXiv.2508.01851.