

BAB I

PENDAHULUAN

1.1 Latar Belakang

Penilaian kredit (*credit scoring*) menjadi salah satu alat utama yang digunakan lembaga keuangan untuk mengelola risiko kredit dengan mengklasifikasikan calon peminjam menjadi dua kelompok, yaitu peminjam yang baik (kemungkinan besar melunasi) dan peminjam yang buruk (berpotensi gagal bayar) [1], [2], [3]. Proses ini sangat penting karena kesalahan klasifikasi dapat menyebabkan kerugian finansial yang signifikan bagi bank dan perusahaan pembiayaan. Namun, dataset kredit pada umumnya bersifat tidak seimbang (*imbalanced*), di mana jumlah kasus gagal bayar jauh lebih sedikit dibandingkan kasus yang berhasil dilunasi. Ketidakseimbangan ini sering kali membuat model klasifikasi sulit mencapai kinerja yang baik, terutama pada portfolio risiko rendah (*low default portfolios*) seperti pinjaman korporasi berkualitas tinggi, pinjaman kepada bank, atau bahkan beberapa portofolio ritel (pinjaman pribadi seperti KPR atau kartu kredit) di negara tertentu yang memiliki tingkat gagal bayar sangat rendah [1]. Dalam kondisi tersebut, bahkan peningkatan akurasi sekecil apa pun dapat menghasilkan penghematan biaya yang besar bagi lembaga keuangan, sebagaimana juga didorong oleh Basel II Accord (standar internasional perbankan tahun 2004 yang mewajibkan bank menyimpan modal cukup sesuai risiko kredit yang dihitung lebih akurat) [1]. Oleh karena itu, pemilihan dan pengembangan metode klasifikasi yang mampu menangani ketidakseimbangan data menjadi tantangan utama dalam bidang ini.

Beberapa metode *ensemble* sudah banyak digunakan untuk mengatasi masalah klasifikasi *credit score*, terutama ketika data tidak seimbang (jumlah peminjam yang gagal bayar jauh lebih sedikit daripada yang melunasi). Metode yang paling sering dipakai adalah *random forest* (RF) dan *gradient boosting decision tree* (GBDT), termasuk pengembangannya seperti XGBoost. Cara kerja metode *ensemble* adalah menggabungkan banyak model sederhana sehingga hasil prediksi keseluruhan menjadi lebih tepat dibandingkan menggunakan satu model saja [4]. Penelitian Li dan Chen (2020) menunjukkan bahwa RF dan XGBoost merupakan dua algoritma teratas dalam menyelesaikan masalah *credit scoring*, dengan RF memberikan performa terbaik dan XGBoost berada di posisi kedua dalam hal akurasi, *Area Under the Curve* (AUC), *Kolmogorov-Smirnov statistic* (KS), serta *Brier score* pada tiga dataset kredit nyata [5].

Selain itu, Qin et al. (2021) membuktikan bahwa XGBoost yang diatur dengan optimasi khusus (adaptive particle swarm optimization) berhasil mengungguli metode pengaturan lain dalam hal akurasi, tingkat kesalahan, *Brier score*, dan *F1 score* pada empat dataset kredit [3]. Penelitian Andri et al. (2026) juga menunjukkan bahwa *tuning hyperparameter* pada *Random Forest & Catboost* memberikan peningkatan signifikan pada akurasi dan *F1-score* di dataset *real-world* yang kompleks dan tidak seimbang, termasuk data *credit scoring* [6]. Selain itu, penelitian tersebut menggunakan XGBoost sebagai model pembandingan yang juga dioptimalkan melalui proses *hyperparameter tuning*, secara konsisten menghasilkan performa yang kompetitif pada berbagai metrik evaluasi [6]. Mengingat kedua algoritma ini telah menunjukkan kinerja yang sangat baik dalam domain *credit scoring*, terdapat peluang yang signifikan untuk menggabungkan kelebihan mereka melalui pendekatan *hybrid*. Pendekatan *hybrid* RF-XGBoost telah diterapkan dengan hasil yang mengesankan di domain lain, seperti deteksi intrusi pada perangkat IoT oleh Al Faysal et al. (2022) dengan akurasi 99.94% dan klasifikasi sentimen ulasan produk e-commerce oleh Alghazzawi et al. (2023) dengan akurasi 98.7% [7], [8]. Keberhasilan tersebut mengindikasikan potensi pendekatan *hybrid* ini. Meskipun demikian, penerapan model *hybrid* ini pada klasifikasi *credit score* masih terbatas, padahal pendekatan semacam ini berpotensi memberikan solusi yang lebih baik dalam menangani data kredit yang tidak seimbang dan memiliki kompleksitas fitur tinggi.

Penelitian ini menawarkan solusi dengan mengimplementasikan model *hybrid* RF-XGBoost untuk klasifikasi *credit score* pada data yang tidak seimbang. Model *hybrid* ini menggabungkan kekuatan *random forest* yang unggul dalam menciptakan keragaman model sehingga lebih tahan terhadap variasi data [5], dengan XGBoost yang unggul dalam memperbaiki kesalahan secara bertahap sehingga akurasi semakin meningkat [3]. Gabungan ini diharapkan dapat menghasilkan performa yang lebih tangguh dalam menghadapi data kredit yang kompleks dan tidak seimbang. Dengan menggunakan metrik evaluasi standar (seperti akurasi, tingkat kesalahan, dan kemampuan membedakan peminjam baik dan buruk), penelitian ini akan mengevaluasi performa model tersebut dalam kondisi nyata. Hasil yang diperoleh diharapkan mengisi keterbatasan penelitian sebelumnya yang belum banyak menerapkan model *hybrid* RF-XGBoost pada domain *credit scoring*. Untuk itu, penelitian ini mengusulkan tugas akhir dengan judul **“Implementasi Model Hybrid Random Forest – XGBoost untuk Klasifikasi Skor Kredit pada Data Tidak Seimbang”**.

1.2 Rumusan Masalah

Penelitian sebelumnya telah menunjukkan bahwa metode *ensemble* seperti *random forest* (RF) dan XGBoost mampu meningkatkan kinerja klasifikasi *credit score*, terutama pada data yang tidak seimbang. Namun, ketidakseimbangan data pada dataset kredit masih seringkali membuat model klasifikasi tidak optimal. Selain itu, penelitian yang mengimplementasikan model *hybrid random forest* – XGBoost dalam klasifikasi *credit score* masih terbatas. Belum banyak studi yang menguji penerapan model *hybrid* tersebut pada dataset kredit yang sama, dengan mempertimbangkan aspek ketidakseimbangan data. Kondisi ini menyebabkan performa prediksi risiko kredit belum dapat dimaksimalkan secara optimal.

1.3 Tujuan

Penelitian ini bertujuan untuk mengimplementasikan dan mengevaluasi model *hybrid random forest* – XGBoost dalam klasifikasi *credit score* pada data yang tidak seimbang. Secara spesifik, tujuan penelitian ini adalah:

1. Mengimplementasikan model *hybrid random forest* – XGBoost yang menggabungkan kelebihan kedua algoritma tersebut untuk menangani data kredit yang tidak seimbang.
2. Melakukan evaluasi kinerja model *hybrid* menggunakan metrik standar dan membandingkannya dengan kinerja algoritma tunggalnya (*random forest* dan XGBoost)

1.4 Manfaat

Penelitian ini diharapkan memberikan manfaat sebagai berikut:

1. Bagi dunia akademik
Memberikan kontribusi teknis berupa pembuktian empiris efektivitas implementasi arsitektur *hybrid random forest* – XGBoost dengan pendekatan *stacking* menggunakan *Logistic Regression* sebagai *meta-learner* pada dataset kredit yang tidak seimbang dan kompleks. Penelitian ini diharapkan dapat menjadi referensi baru dalam pengembangan model *ensemble* untuk *credit scoring*, khususnya dalam menangani permasalahan *class imbalance* pada data keuangan yang memiliki karakteristik *high-dimensional* dan *low default portfolio*.
2. Bagi lembaga keuangan atau praktisi

Menyajikan hasil evaluasi kinerja model *hybrid* ini sebagai gambaran nyata untuk dipertimbangkan dalam implementasi sistem penilaian kredit yang lebih akurat dan efisien, terutama pada kondisi data yang tidak seimbang.

1.5 Ruang Lingkup

Ruang lingkup penelitian ini meliputi:

1. Dataset yang digunakan adalah *Credit Score Classification* dari Kaggle yang terdiri dari 100.000 sampel dengan distribusi kelas sebagai berikut:
 - A. *Good*: 17.800 sampel (17,8%)
 - B. *Standard*: 53.200 sampel (53,2%)
 - C. *Poor* : 29.000 sampel (29%)Dataset dapat diakses pada:
<https://www.kaggle.com/datasets/parisrohan/creditscoreclassification?resource=download&select=train.csv>
2. *Exploratory Data Analysis* (EDA) menggunakan teknik visualisasi dasar yaitu histogram, *boxplot*, dan *heatmap* korelasi untuk memahami pola distribusi data serta hubungan antar fitur.
3. Untuk keperluan implementasi dan evaluasi model, data dibagi menjadi *training set* (80%) dan *test set* tetap (20%) menggunakan *stratified split*. Model *hybrid* RF-XGBoost diimplementasikan dengan pendekatan *stacking* menggunakan *Logistic Regression* sebagai *meta-learner* dan *5-fold cross-validation* secara internal untuk pelatihan dan optimasi. Evaluasi dilakukan pada data uji yang sama. Penelitian ini menggunakan konfigurasi *hyperparameter* secara statis/manual (tidak dilakukan *hyperparameter tuning* otomatis).
4. Pengukuran performa model dinilai menggunakan metrik standar yaitu *accuracy*, *precision*, *recall*, *F1-score*, serta *Area Under the ROC Curve* (AUC-ROC).