

**ANALISIS OVERFITTING PADA PREDIKSI RISIKO GAGAL
BAYAR PINJAMAN MENGGUNAKAN ALGORITMA
DECISION TREE**

SKRIPSI

Diajukan untuk memenuhi salah satu syarat mencapai derajat Sarjana
Program Studi S1-Informatika



disusun oleh

OXA DEFRIZAL KHASAY

21.11.4388

Kepada

**FAKULTAS ILMU KOMPUTER
UNIVERSITAS AMIKOM YOGYAKARTA
YOGYAKARTA**

2025

**ANALISIS OVERFITTING PADA PREDIKSI RISIKO GAGAL
BAYAR PINJAMAN MENGGUNAKAN ALGORITMA
DECISION TREE**

SKRIPSI

untuk memenuhi salah satu syarat mencapai derajat Sarjana
Program Studi S1-Informatika



disusun oleh

OXA DEFRIZAL KHASAY

21.11.4388

Kepada

**FAKULTAS ILMU KOMPUTER
UNIVERSITAS AMIKOM YOGYAKARTA
YOGYAKARTA**

2025

HALAMAN PERSETUJUAN

HALAMAN PERSETUJUAN

SKRIPSI

ANALISIS OVERFITTING PADA PREDIKSI RISIKO GAGAL BAYAR PINJAMAN MENGGUNAKAN ALGORITMA DECISION TREE

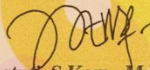
yang disusun dan diajukan oleh

Oxa Defrizal Khasay

21.11.4388

telah disetujui oleh Dosen Pembimbing Skripsi
pada tanggal 24 Januari 2025

Dosen Pembimbing,



Yuli Astuti, S.Kom., M.Kom.
NIK. 190302146

HALAMAN PENGESAHAN

HALAMAN PENGESAHAN

SKRIPSI

ANALISIS OVERFITTING PADA PREDIKSI RISIKO GAGAL BAYAR PINJAMAN MENGGUNAKAN ALGORITMA DECISION TREE

yang disusun dan diajukan oleh

Oxa Defrizal Khasay

21.11.4388

Telah dipertahankan di depan Dewan Penguji
pada tanggal 24 Januari 2025

Susunan Dewan Penguji

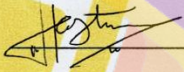
Nama Penguji

Tanda Tangan

Dr. Ferry Wahyu Wibowo, S.Si., M.Cs.
NIK. 190302235



Hastari Utama, S.Kom., M.Cs.
NIK. 190302230



Yuli Astuti, S.Kom., M.Cs.
NIK. 190302146



Skrripsi ini telah diterima sebagai salah satu persyaratan
untuk memperoleh gelar Sarjana Komputer
Tanggal 24 Januari 2025

DEKAN FAKULTAS ILMU KOMPUTER



Hanif Al Fatta, S.Kom., M.Kom., Ph.D.
NIK. 190302096

HALAMAN PERNYATAAN KEASLIAN SKRIPSI

HALAMAN PERNYATAAN KEASLIAN SKRIPSI

Yang bertandatangan di bawah ini,

Nama mahasiswa : Oxa Defrizal Khasay
NIM : 21.11.4388

Menyatakan bahwa Skripsi dengan judul berikut:

**Analisis Overfitting pada Prediksi Risiko Gagal Bayar Pinjaman
Menggunakan Algoritma Decision Tree**

Dosen Pembimbing : Yuli Astuti, S.Kom., M.Kom.

1. Karya tulis ini adalah benar-benar ASLI dan BELUM PERNAH diajukan untuk mendapatkan gelar akademik, baik di Universitas AMIKOM Yogyakarta maupun di Perguruan Tinggi lainnya.
2. Karya tulis ini merupakan gagasan, rumusan dan penelitian SAYA sendiri, tanpa bantuan pihak lain kecuali arahan dari Dosen Pembimbing.
3. Dalam karya tulis ini tidak terdapat karya atau pendapat orang lain, kecuali secara tertulis dengan jelas dicantumkan sebagai acuan dalam naskah dengan disebutkan nama pengarang dan disebutkan dalam Daftar Pustaka pada karya tulis ini.
4. Perangkat lunak yang digunakan dalam penelitian ini sepenuhnya menjadi tanggung jawab SAYA, bukan tanggung jawab Universitas AMIKOM Yogyakarta.
5. Pernyataan ini SAYA buat dengan sesungguhnya, apabila di kemudian hari terdapat penyimpangan dan ketidakbenaran dalam pernyataan ini, maka SAYA bersedia menerima SANKSI AKADEMIK dengan pencabutan gelar yang sudah diperoleh, serta sanksi lainnya sesuai dengan norma yang berlaku di Perguruan Tinggi.

Yogyakarta, 24 Januari 2025

Yang Menyatakan,



Oxa Defrizal Khasay

KATA PENGANTAR

Puji syukur penulis panjatkan ke hadirat Allah SWT atas limpahan rahmat, nikmat, dan karunia-Nya, sehingga penulis dapat menyelesaikan skripsi yang berjudul *"Analisis Overfitting pada Prediksi Risiko Gagal Bayar Pinjaman Menggunakan Algoritma Decision Tree"* ini dengan baik. Skripsi ini disusun sebagai salah satu syarat untuk memperoleh gelar Sarjana Komputer pada Program Studi Informatika, Universitas Amikom Yogyakarta.

1. Allah SWT, atas segala rahmat, kemudahan, dan kekuatan yang diberikan selama proses pengerjaan skripsi ini.
2. Ibu Dosen Pembimbing, Yuli Astuti M.Kom, yang telah memberikan bimbingan, arahan, dan masukan yang sangat berarti selama proses penyusunan skripsi ini.
3. Kedua orang tua, yang selalu mendoakan, memberikan dukungan moral maupun material, serta menjadi sumber motivasi utama dalam setiap langkah perjalanan penulis.
4. Rekan-rekan mahasiswa angkatan 2021, yang senantiasa memberikan semangat, inspirasi, dan kebersamaan selama masa perkuliahan maupun proses penyelesaian skripsi ini.
5. Komunitas Kaggle dan Nikhil, yang menyediakan data serta sumber daya yang sangat membantu dalam proses penelitian ini.

Penulis menyadari bahwa skripsi ini masih memiliki keterbatasan dan jauh dari sempurna. Oleh karena itu, penulis sangat mengharapkan saran dan kritik yang membangun untuk perbaikan karya ini di masa mendatang.

Akhir kata, penulis berharap semoga skripsi ini dapat memberikan manfaat, baik bagi perkembangan ilmu pengetahuan maupun bagi pembaca yang tertarik dengan bidang analisis data dan pembelajaran mesin.

Yogyakarta, 24 Januari 2025

Penulis

DAFTAR ISI

HALAMAN JUDUL	i
HALAMAN PERSETUJUAN.....	ii
HALAMAN PENGESAHAN	iii
HALAMAN PERNYATAAN KEASLIAN SKRIPSI	iv
KATA PENGANTAR	v
DAFTAR ISI.....	vi
DAFTAR TABEL.....	ix
DAFTAR GAMBAR.....	x
DAFTAR LAMBANG DAN SINGKATAN	xi
DAFTAR ISTILAH.....	xiii
INTISARI	xix
<i>ABSTRACT</i>	xx
BAB I PENDAHULUAN.....	1
1.1 Latar Belakang	1
1.2 Rumusan Masalah.....	2
1.3 Batasan Masalah	3
1.4 Tujuan Penelitian	4
1.5 Manfaat Penelitian	4
1.6 Sistematika Penulisan	5
BAB II TINJAUAN PUSTAKA	6
2.1 Studi Literatur	6
2.2 Dasar Teori.....	12
2.2.1 Ekonomi.....	12

2.2.2	Risiko	12
2.2.3	Pinjaman	12
2.2.4	Overfitting.....	12
2.2.5	Prediksi	12
2.2.6	Machine Learning	13
2.2.7	SMOTE	13
2.2.8	K-Fold Cross Validation	13
2.2.9	Hyperparameter tuning	14
2.2.10	Decision Tree	14
2.2.11	Evaluasi Model	15
BAB III METODE PENELITIAN		18
3.1	Alur Penelitian	18
3.1.1	Pengumpulan Data	20
3.1.2	Exploratory Data Analysis (EDA)	20
3.1.3	Data Preprocessing.....	21
3.1.4	Splitting Data	23
3.1.5	Implementasi Model Decision Tree.....	23
3.1.6	Tunning Hyperparameter Model.....	24
3.1.7	Evaluasi.....	25
3.2	Alat dan Bahan.....	26
3.2.1	Alat.....	26
3.2.2	Bahan	28
BAB IV HASIL DAN PEMBAHASAN		29
4.1	Pengumpulan Data	29
4.2	Exploratory Data Analysis (EDA)	29

4.3	Data Preprocessing.....	33
4.3.1	Drop Kolom	33
4.3.2	Label Encoding	33
4.3.3	Penanganan Duplikat	34
4.3.4	Pemisahan Label dan fitur.....	34
4.3.5	Penanganan Data yang Tidak Seimbang.....	35
4.4	Modeling	36
4.4.1	Splitting Data	36
4.4.2	Pelatihan Model	37
4.5	Evaluasi.....	42
BAB V PENUTUP		48
5.1	Kesimpulan	48
5.2	Saran	48
REFERENSI		50

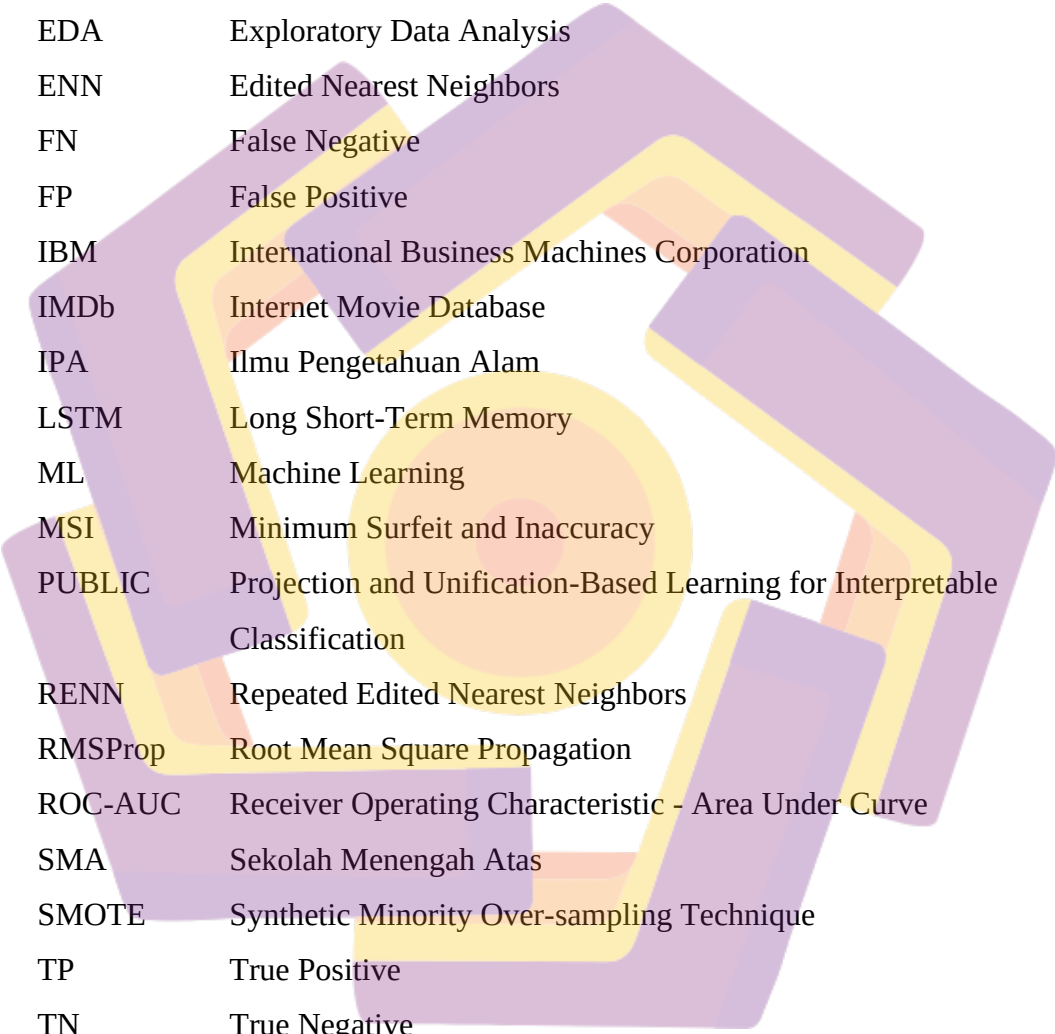
DAFTAR TABEL

Tabel 2.1 Keaslian Penelitian.....	8
Tabel 3.1 Parameter pada model Decision tree	24
Tabel 4.1 Hasil parameter terbaik pada konfigurasi SMOTE + hyperparameter tunning + 80:20 split	41
Tabel 4.2 Hasil parameter terbaik pada konfigurasi SMOTE + hyperparameter tunning + K-fold Cross Validation	42
Tabel 4.3 Hasil evaluasi model Decision Tree sebelum penerapan teknik penanganan ketidakseimbangan data.	42
Tabel 4.4 Hasil evaluasi model Decision Tree setelah penerapan teknik penanganan ketidakseimbangan data.	43
Tabel 4.5 Hasil evaluasi model Decision Tree menggunakan cross-validation....	43
Tabel 4.6 Hasil evaluasi model Decision Tree menggunakan hyperparameter tunning dengan teknik grid search.	44
Tabel 4.7 Hasil evaluasi model Decision Tree menggunakan hyperparameter tunning dengan teknik grid search dan cross-validation.	44

DAFTAR GAMBAR

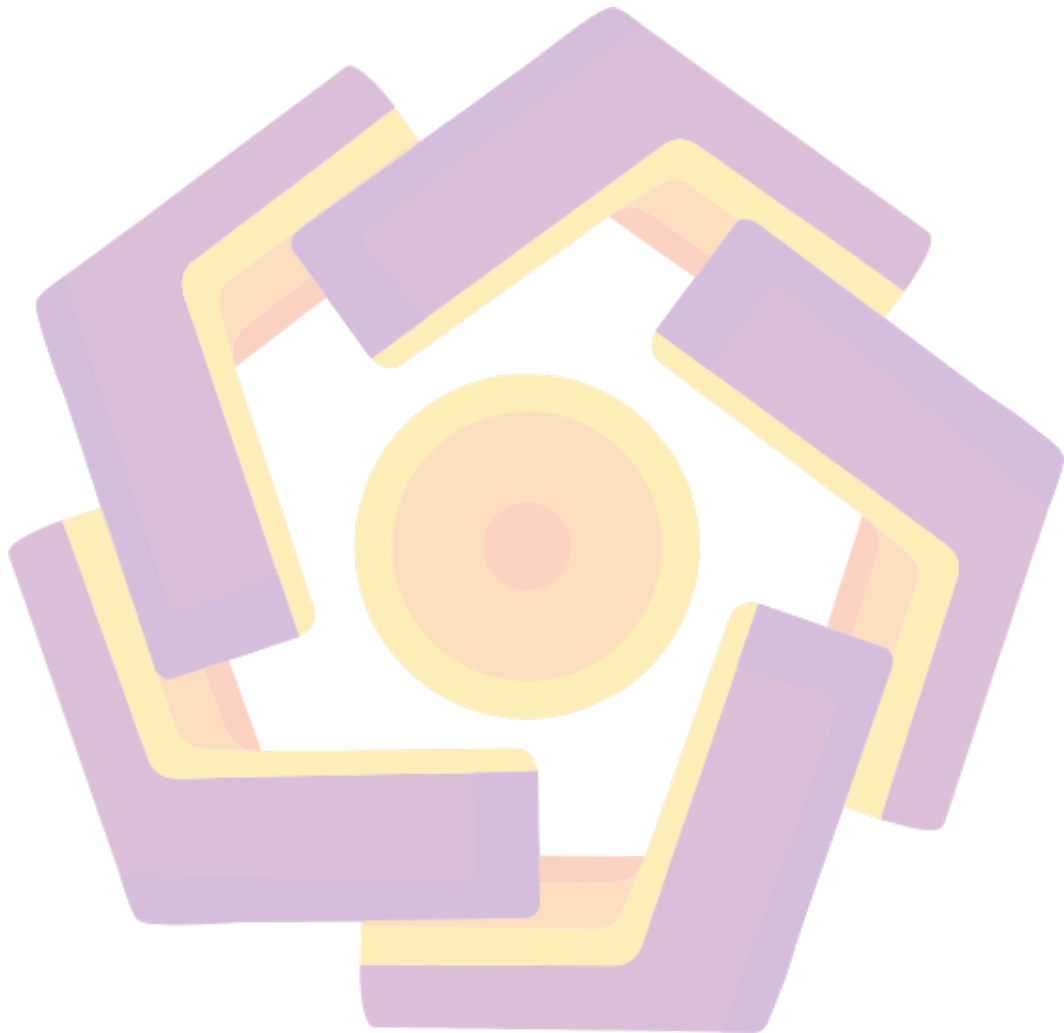
Gambar 2.1 Confusion Matrix.....	16
Gambar 3.1 Alur penelitian tanpa teknik SMOTE.....	18
Gambar 3.2 Alur penelitian dengan teknik SMOTE.....	19
Gambar 4.1 Hasil pengambilan data dari Kaggle.....	29
Gambar 4.2 Tabel statistik deskriptif dataset Prediksi Peminjaman.	29
Gambar 4.3 Tabel informasi dataset Prediksi Peminjaman.....	30
Gambar 4.4 Grafik Pie Chart Keputusan Peminjaman.....	31
Gambar 4.5 Grafik Box Plot kolom Income dan LoanAmount.	32
Gambar 4.6 Grafik Box Plot kolom age, CreditScore, MonthsEmployed, NumCreditLines, InterestRate, LoanTerm, dan DTRatio.....	32
Gambar 4.7 Kode program drop kolom.	33
Gambar 4.8 Hasil kolom yang telah encode.....	33
Gambar 4.9 Kode program penanganan duplikat.....	34
Gambar 4.10 Kode Program pemisahan label fitur dan target.....	35
Gambar 4.11 Grafik pie chart Keputusan peminjaman setelah data di seimbangkan menggunakan teknik SMOTE.....	36
Gambar 4.12 Kode Program Splitting data training dan test.	37
Gambar 4.13 Kode program Splitting data dengan teknik Cross Validation.	37
Gambar 4.14 Kode program pelatihan model Decision Tree.....	38
Gambar 4.15 Kode program pelatihan model Decision Tree dengan K-Fold Cross Validation.....	39
Gambar 4.16 Kode Program pelatihan model Decision Tree yang Tunning dengan teknik Grid Search.....	40
Gambar 4.17 Kode Program pelatihan model Decision Tree yang Tunning dengan teknik grid search dan di Cross Validation.....	41
Gambar 4.18 Grafik line chart perbandingan akurasi pelatihan dan pengujian berdasarkan kedalaman pohon di berbagai konfigurasi.	45
Gambar 4.19 Grafik bar chart perbandingan akurasi data pelatihan dan pengujian di berbagai konfigurasi.	46

DAFTAR LAMBANG DAN SINGKATAN



ALLKNN	All K-Nearest Neighbors
Bagging	Bootstrap Aggregating
CART	Classification and Regression Trees
CHAID	Chi-Square Automatic Interaction Detection
EDA	Exploratory Data Analysis
ENN	Edited Nearest Neighbors
FN	False Negative
FP	False Positive
IBM	International Business Machines Corporation
IMDb	Internet Movie Database
IPA	Ilmu Pengetahuan Alam
LSTM	Long Short-Term Memory
ML	Machine Learning
MSI	Minimum Surfeit and Inaccuracy
PUBLIC	Projection and Unification-Based Learning for Interpretable Classification
RENN	Repeated Edited Nearest Neighbors
RMSProp	Root Mean Square Propagation
ROC-AUC	Receiver Operating Characteristic - Area Under Curve
SMA	Sekolah Menengah Atas
SMOTE	Synthetic Minority Over-sampling Technique
TP	True Positive
TN	True Negative
SNMPTN	Seleksi Nasional Masuk Perguruan Tinggi Negeri
WEB	World Wide Web
K	Jumlah lipatan (Folds) dalam cross-validation
$Akurasi_i$	Akurasi model pada lipatan ke- i
S	Himpunan kasus
A	Atribut

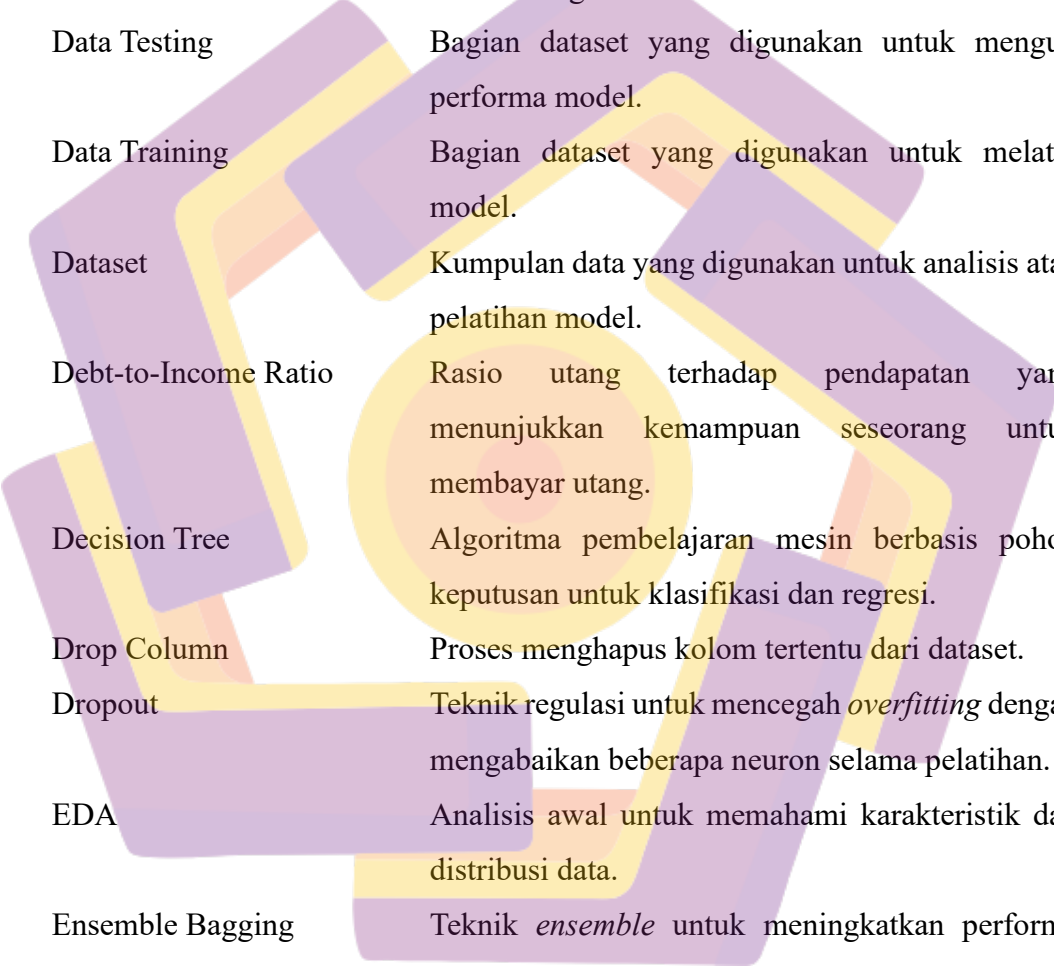
N	Jumlah Partisi atribut A
$ S_i $	Jumlah kasus pada partisi ke- i
$ S $	Jumlah kasus dalam S
N	Jumlah partisi S
P_i	Proporsi dari S_i terhadap S



DAFTAR ISTILAH



ALLKNN	Teknik <i>undersampling</i> untuk mengurangi <i>noise</i> dalam data.
Average Error	Kesalahan rata-rata yang dihitung selama proses evaluasi model.
Bagging	Teknik <i>ensemble</i> untuk meningkatkan stabilitas dan akurasi model.
Bayesian Optimization	Metode untuk mengoptimalkan <i>hyperparameter</i> model secara efisien.
Bayes Risk Post-Pruning	Teknik <i>post-pruning</i> untuk mengurangi <i>overfitting</i> dalam pohon keputusan.
Biner	Sistem numerik yang hanya menggunakan dua nilai, 0 dan 1.
Box Plot	Grafik yang digunakan untuk menggambarkan distribusi data dan mendeteksi <i>outlier</i> .
C4.5	Algoritma <i>Decision Tree</i> yang merupakan pengembangan dari algoritma ID3.
CART	Algoritma <i>Decision Tree</i> yang digunakan untuk klasifikasi dan regresi.
CHAID	Algoritma <i>Decision Tree</i> berbasis statistik Chi-squared.
Churn	Istilah yang mengacu pada pelanggan yang berhenti menggunakan layanan.
Confusion Matrix	Matriks untuk mengevaluasi performa model klasifikasi dengan membandingkan prediksi dan data aktual.
Criterion	Parameter dalam <i>Decision Tree</i> yang menentukan cara memisahkan data, seperti <i>Gini</i> atau <i>Entropy</i> .



Cross Validation	Teknik evaluasi model dengan membagi dataset menjadi beberapa bagian untuk melatih dan menguji model.
Data Mining	Proses menemukan pola atau informasi penting dari data besar.
Data Preprocessing	Proses membersihkan dan mempersiapkan data sebelum digunakan dalam model.
Data Testing	Bagian dataset yang digunakan untuk menguji performa model.
Data Training	Bagian dataset yang digunakan untuk melatih model.
Dataset	Kumpulan data yang digunakan untuk analisis atau pelatihan model.
Debt-to-Income Ratio	Rasio utang terhadap pendapatan yang menunjukkan kemampuan seseorang untuk membayar utang.
Decision Tree	Algoritma pembelajaran mesin berbasis pohon keputusan untuk klasifikasi dan regresi.
Drop Column	Proses menghapus kolom tertentu dari dataset.
Dropout	Teknik regulasi untuk mencegah <i>overfitting</i> dengan mengabaikan beberapa neuron selama pelatihan.
EDA	Analisis awal untuk memahami karakteristik dan distribusi data.
Ensemble Bagging	Teknik <i>ensemble</i> untuk meningkatkan performa model melalui <i>bagging</i> .
Entropy	Ukuran ketidakpastian dalam data yang digunakan untuk membangun <i>Decision Tree</i> .
ENN	Teknik <i>undersampling</i> untuk menghilangkan <i>noise</i> dari data.



F1-Score	Rata-rata harmonis antara presisi dan <i>recall</i> yang digunakan untuk mengevaluasi model pada data tidak seimbang.
Feature (X)	Variabel independen dalam dataset yang digunakan sebagai input untuk model.
Filter Noise	Proses menghapus data yang mengganggu atau tidak relevan dalam dataset.
Forecasting	Proses memperkirakan nilai atau tren masa depan berdasarkan data yang ada.
Gain	Ukuran yang digunakan untuk menentukan atribut terbaik dalam pembagian pada <i>Decision Tree</i> .
Gini	Ukuran ketidakmurnian dalam data yang digunakan dalam <i>Decision Tree</i>
Grid Search	Metode untuk menemukan kombinasi terbaik dari <i>hyperparameter</i> dengan mencoba setiap kemungkinan.
Hyperparameter	Parameter yang harus ditentukan sebelum proses pelatihan model.
Hyperparameter Tuning	Proses mengoptimalkan <i>hyperparameter</i> untuk meningkatkan performa model.
IBM	Perusahaan teknologi yang menyediakan berbagai dataset untuk analisis dan pengembangan model.
Imbalance Data	Ketidakseimbangan dalam distribusi data antar kelas.
Instance	Satu entri atau baris data dalam dataset.
Integer	Jenis data numerik yang tidak memiliki nilai desimal
K-Fold Cross-Validation	Teknik <i>cross-validation</i> yang membagi dataset menjadi K bagian untuk melatih dan menguji model secara bergantian.



Kaggle	Platform komunitas untuk berbagi dataset dan kompetisi <i>data science</i>
Label (y)	Variabel dependen dalam dataset yang menjadi target prediksi model.
Label Encoding	Teknik mengubah data kategorikal menjadi bentuk numerik
LSTM	Jenis jaringan saraf tiruan untuk data berurutan seperti teks dan waktu.
Maximum Error	Kesalahan terbesar yang ditemukan dalam prediksi model.
Min_Samples_Leaf	Parameter dalam <i>Decision Tree</i> yang menentukan jumlah minimum sampel di setiap daun.
Min_Samples_Split	Parameter dalam <i>Decision Tree</i> yang menentukan jumlah minimum sampel untuk memisahkan simpul.
Model	Representasi matematis yang digunakan untuk membuat prediksi berdasarkan data.
MSI	Teknik untuk menghindari <i>overfitting</i> tanpa menggunakan <i>hyperparameter</i> .
Naive Bayes	Algoritma klasifikasi berbasis teorema <i>Bayes</i> dengan asumsi independensi antar fitur.
Node	Titik dalam <i>Decision Tree</i> yang mewakili atribut atau kondisi
Node Daun	Titik akhir dalam <i>Decision Tree</i> yang menunjukkan kelas prediksi.
Noise	Data yang tidak relevan atau mengganggu dalam dataset.
Object	Jenis data yang berisi string atau nilai kategorikal.
Outlier	Data yang berbeda secara signifikan dari data lain dalam dataset.



Pie Chart	Grafik melingkar untuk menunjukkan proporsi data
Post-Pruning	Proses mengurangi ukuran <i>Decision Tree</i> setelah pohon selesai dibangun untuk mencegah overfitting.
Pre-Pruning	Proses menghentikan pembangunan <i>Decision Tree</i> lebih awal untuk mencegah overfitting.
Pruning	Teknik untuk mengurangi kompleksitas <i>Decision Tree</i> dengan menghapus cabang yang tidak relevan.
PUBLIC	Algoritma pohon keputusan yang sering digunakan untuk klasifikasi.
Random Forest	Algoritma ensemble yang terdiri dari beberapa <i>Decision Tree</i> untuk meningkatkan akurasi.
Random Search	Metode mencari <i>hyperparameter</i> terbaik dengan mencoba kombinasi secara acak.
Random_State	Parameter yang memastikan hasil pembagian data atau pengacakan konsisten.
Recall	Kemampuan model untuk mendeteksi semua data positif yang benar.
RENN	Teknik untuk menghapus <i>noise</i> dari data melalui beberapa iterasi.
RMSProp	Optimizer untuk jaringan saraf tiruan yang memperbaiki <i>learning rate</i> secara adaptif.
Scikit-Learn	Library Python untuk pembelajaran mesin dan analisis data.
SMOTE	Teknik untuk menyeimbangkan data dengan membuat sampel sintetis dari kelas minoritas.
Splitting	Proses membagi dataset menjadi bagian pelatihan dan pengujian.
Sqrt	Akar kuadrat, digunakan sebagai parameter dalam <i>Decision Tree</i> .

Standard Deviation	Ukuran seberapa jauh data tersebar dari nilai rata-rata.
Supervised Machine Learning	Jenis pembelajaran mesin di mana model dilatih menggunakan data berlabel.
Train-Test Split	Teknik membagi dataset menjadi data pelatihan dan pengujian dengan rasio tertentu.
Validation Loss	Kesalahan model pada data validasi selama pelatihan.
Variabel Dependen	Variabel target yang diprediksi oleh model.
Variabel Independen	Variabel yang digunakan sebagai input untuk model.
Web	Platform untuk mengakses informasi secara online.
Weight Decay	Teknik regulasi untuk mencegah overfitting dengan menambahkan penalti pada bobot besar dalam model.
X_test	Data fitur untuk pengujian model.
X_train	Data fitur untuk melatih model.
y_pred_test	Hasil prediksi model pada data pengujian.
y_pred_train	Hasil prediksi model pada data pelatihan.
y_test	Label sebenarnya pada data pengujian.
y_train	Label sebenarnya pada data pelatihan.

INTISARI

Pertumbuhan ekonomi suatu negara dapat terganggu oleh berbagai permasalahan, seperti pengangguran, inflasi, dan ketimpangan pendapatan. Dalam konteks perekonomian Indonesia, tingkat pengangguran yang tinggi sering kali mendorong terjadinya tindakan kriminal, termasuk penipuan dan penggelapan uang di sektor keuangan. Salah satu bentuk kejahatan yang sering terjadi adalah gagal bayar pinjaman, yang tidak hanya merugikan lembaga keuangan, tetapi juga memengaruhi nasabah baru yang beritikad baik. Teknologi *machine learning* telah diterapkan secara luas dalam mendukung pengambilan keputusan, khususnya dalam prediksi risiko gagal bayar pinjaman. Salah satu permasalahan utama dalam pengembangan model prediksi adalah *overfitting*, yaitu kondisi ketika model terlalu menyesuaikan data pelatihan sehingga kinerjanya menurun pada data baru. Penelitian ini bertujuan untuk menganalisis dan mengurangi *overfitting* pada model prediksi gagal bayar pinjaman menggunakan algoritma *Decision Tree*. Metode penelitian meliputi pengumpulan data dari platform *Kaggle*, eksplorasi data untuk memahami karakteristik dataset, *preprocessing* untuk membersihkan dan menyiapkan data, serta pembagian data menggunakan *train-test split* dan *k-fold cross-validation*. Selain itu, *tuning* parameter dilakukan melalui *Grid Search* untuk menemukan kombinasi parameter terbaik. Hasil penelitian menunjukkan bahwa kombinasi teknik SMOTE, dan *hyperparameter tuning* mampu menghasilkan model dengan akurasi 90,06%, presisi 90,52%, *recall* 90,06%, *f1-score* 90.03%, dan perbedaan kecil antara akurasi pelatihan dan uji. Penerapan teknik SMOTE terbukti efektif dalam menangani ketidakseimbangan data, yang berkontribusi pada peningkatan kinerja model. Penggunaan *k-fold cross-validation* juga berperan penting dalam memastikan stabilitas model saat memprediksi risiko gagal bayar pinjaman.

Kata kunci: *Overfitting, Decision Tree, Prediksi, SMOTE, Hyperparameter Tuning*

ABSTRACT

Economic growth in a country can be disrupted by various issues, such as unemployment, inflation, and income inequality. In the context of Indonesia's economy, the high unemployment rate often leads to criminal activities, including fraud and embezzlement in the financial sector. One of the most common crimes is loan default, which not only harms financial institutions but also affects new borrowers with good intentions. Machine learning technology has been widely applied to support decision-making, particularly in predicting loan default risk. One of the main challenges in developing predictive models is overfitting, a condition where the model overly fits the training data, resulting in poor performance on new data. This study aims to analyze and mitigate overfitting in predicting loan default risk using the Decision Tree algorithm. The research methodology includes data collection from the Kaggle platform, data exploration to understand dataset characteristics, preprocessing to clean and prepare the data, and data splitting using train-test split and k-fold cross-validation. Additionally, parameter tuning was conducted using Grid Search to find the best parameter combinations. The results show that the combination of SMOTE and hyperparameter tuning achieved a model with an accuracy of 90.06%, precision of 90.52%, recall of 90.06%, and an F1-score of 90.03%, with a minimal difference between training and testing accuracy. The application of SMOTE proved effective in handling imbalanced data, contributing to improved model. The use of k-fold cross-validation also played a crucial role in ensuring model stability when predicting loan default risk.

Keyword: Overfitting, Decision Tree, Prediction, SMOTE, Hyperparameter Tuning