

IMPLEMENTASI METODE SMOTE DAN RANDOM OVER-SAMPLING PADA ALGORITMA MACHINE LEARNING UNTUK PREDIKSI CUSTOMER CHURN DI SEKTOR PERBANKAN

Fannisa Salsabila Pratiwi¹, Mula Agung Barata², Aprillia Dwi Ardianti³
Program Studi Teknik Infomatika^{1,2}, Universitas Nahdlatul Ulama Sunan Giri
Program Studi Teknik Mesin³, Universitas Nahdlatul Ulama Sunan Giri
Jl. Ahmad Yani No.10, Jamban, Sukorejo, Kec. Bojonegoro, Kabupaten
Bojonegoro, Jawa Timur 62115
e-mail : *1fannisasabil2020@gmail.com, 2mula.ab26@gmail.com,
3aprilliadwia@unugiri.ac.id

Abstract

The ability to anticipate unsubscribed customers is a challenge in the competitive banking industry, where it is more efficient to retain customers than to attract new ones. The purpose of this study is to improve the effectiveness of churn prediction by overcoming data imbalances using SMOTE (Synthetic Minority Oversampling Technique) and Random Over-sampling. The data set used consists of 10. 000 bank customer data, with 12 important attributes, including churn indicators as targets. The machine learning algorithms used are Random Forest and Neive Bayes, evaluated based on accuracy, precision, recall, and F1 scores. The results of the experiment showed that the highest accuracy of 87.13% could be achieved with the Random Forest algorithm without using the oversampling method, but its effectiveness in detecting churn customers was slightly limited. The use of SMOTE and Random Over-sampling methods has improved the model's performance in identifying churn patterns, although it has led to a decrease in accuracy to 86.20% for Random Over-sampling and 81.47% for SMOTE. Nevertheless, the Neive Bayes algorithm showed the best accuracy rate of 79.20% without oversampling, although it was still slightly lacking in optimal churn handling. The study underscores the importance of using oversampling methods to improve prediction balance in minority classes, which is often overlooked in conventional models. It is hoped that the results of this research can be used as a guide in improving strategies to maintain customer trust that are more up-to-date and efficient.

Keywords: Customer Churn, Machine Learning Algorithms, Random Oversampling, SMOTE

PENDAHULUAN

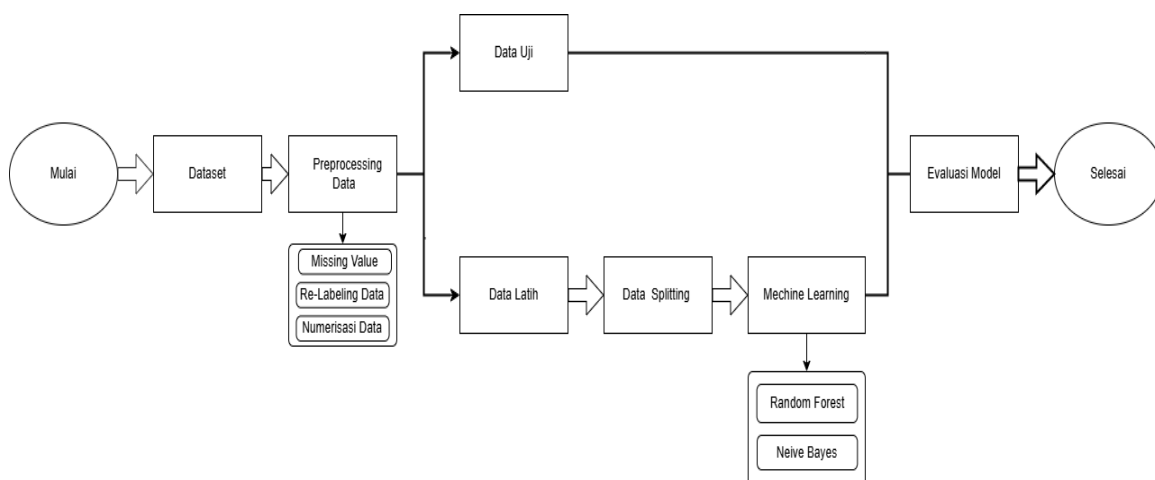
Kepuasan nasabah mempunyai pengaruh yang signifikan terhadap loyalitas nasabah yang merupakan ukuran kesetiaan nasabah terhadap suatu bank atau perusahaan. Retensi pelanggan adalah tujuan strategis yang bertujuan untuk menjaga hubungan pelanggan jangka panjang, dan retensi pelanggan mencerminkan pergantian pelanggan (Istiqomawati et al., 2022). Beralihnya nasabah dari suatu bank ke bank lain umumnya dikenal dengan istilah *Churn*. *Customer churn* atau berhentinya pelanggan dalam menggunakan layanan merupakan tantangan besar bagi sektor perbankan, terutama di tengah persaingan yang ketat dan perubahan ekspektasi pelanggan (Farid Naufal et al., 2023). *Customer churn* atau kehilangan pelanggan adalah fenomena dimana pelanggan berhenti menggunakan produk atau jasa suatu perusahaan, baik secara sukarela maupun tidak. Dalam industri perbankan, *churn* sering terjadi ketika nasabah menutup rekeningnya, beralih ke bank lain, atau berhenti menggunakan layanan yang diberikan (Istiqomawati et al., 2022). *Churn* dapat disebabkan oleh berbagai faktor, antara lain ketidakpuasan terhadap layanan, biaya yang tinggi, penawaran yang lebih menarik dari bank lain, dan pengalaman yang buruk terhadap layanan bank itu sendiri (Christanto & Wiwik, 2024). Dalam industri yang kompetitif seperti perbankan, biaya untuk memperoleh pelanggan baru biasanya jauh lebih tinggi dibandingkan biaya mempertahankan pelanggan yang sudah ada, sehingga memahami dan mencegah churn pelanggan sangatlah penting (Syarif & Nugraha, 2023). Dengan mengidentifikasi nasabah yang berisiko *churn*, bank dapat menerapkan strategi retensi seperti mempersonalisasi layanan dan menawarkan insentif untuk memperkuat loyalitas nasabah dan meningkatkan profitabilitas jangka panjang (Abdullah & Widarto, 2024). Klasifikasi banyak digunakan dalam berbagai sistem, seperti prediksi penjualan, sistem diagnosa medis, dan sebagainya. Sampel data yang digunakan dalam penelitian ini adalah data sekunder nasabah Bank ABC Multistate yang berjumlah 10.000 data yang diperoleh dari *website* Kaggle. Data ini berisi informasi tentang nasabah bank dan terdiri dari 12 atribut, termasuk atribut status *churn* pelanggan (Berhenti Berlangganan) (Namira et al., 2024).

Pada Studi mengenai prediksi customer *churn* telah diselidiki secara luas di berbagai sektor, terutama di industri yang bergantung pada retensi pelanggan sebagai faktor penting untuk keberlangsungan bisnisnya, seperti telekomunikasi dan perbankan (Hermawan, 2024). Dalam ranah sektor perbankan, satu dari utama ialah ketidakseimbangan data ketika meramal *churn*, dengan kebanyakan pelanggan tetap setia maupun cenderung berhenti. Hal ini dapat menyebabkan model prediksi menjadi berat sebelah dan kurang responsif terhadap pelanggan yang mungkin memilih untuk berhenti. Penelitian terkait prediksi *churn* di sektor perbankan dengan menggunakan algoritma *Random forest* dan pohon keputusan (Azmi, 2024). Dalam konteks ini menunjukkan hasil bahwa *Random Forest*, melebihi *Decision Tree* dengan mencapai tingkat *precision* yang lebih tinggi (79% vs 72%), *recall* (78% vs 72%), *F1-score* (78% vs 72%), dan *accuracy* (78% vs 72%) (Azmi, 2024). Namun Penelitian lain ,mennggunaakn algoritma *Random Forest* yang menghasilkan akurasi hingga 90,83% deengan data testing dan data training 90:10.

Penelitian ini memanfaatkan *Confusion Matrix* sebagai metode evaluasi utama untuk memperhitungkan akurasi, presisi, *recall*, dan *F1-score* dari model yang dibangun (Azmi, 2024). Tujuannya adalah agar hasil penelitian bisa dibandingkan langsung dengan penelitian sebelumnya yang juga memanfaatkan *Confusion Matrix* untuk mengevaluasi kinerja model. Dengan demikian, akan memberikan pemahaman yang lebih menyeluruh tentang kelebihan dan kekurangan dari metode yang digunakan (Anjani et al., 2023). Dengan memanfaatkan metode *SMOTE* dan *Random Over-sampling*, penelitian ini bertujuan untuk mengatasi tantangan ketidakseimbangan data dalam memprediksi *churn*, di mana kebanyakan jumlah nasabah tetap (*non-churn*) jauh lebih besar daripada nasabah yang berhenti (Novitasari et al., 2024). Diharapkan bahwa penerapan metode ini akan memberikan kontribusi positif bagi model dalam meningkatkan keefektifannya dalam mendeteksi pola *churn*.

METODE PENELITIAN

Penelitian ini mengembangkan model untuk memprediksi customer *churn* di sektor perbankan dengan menerapkan teknik *Synthetic Minority Oversampling (SMOTE)* dan teknik *Random Oversampling* untuk mengevaluasi kinerja algoritma *machine learning* untuk memperbaiki ketidakseimbangan kelas (Nugroho & Rilvani, 2024). Model ini dirancang untuk membantu bank memahami karakteristik nasabah yang cenderung memilih untuk tidak ikut serta dan mengembangkan strategi retensi nasabah yang efektif. Penelitian ini menggunakan pendekatan kuantitatif dengan desain eksperimental dan menguji setiap tahapannya untuk mengetahui efektivitas metode *oversampling* dalam meningkatkan akurasi prediksi (Anjani et al., 2023).



Gambar 1. Flowcart alur Penelitian

a. Dataset

Dataset yang digunakan pada penelitian ini menggunakan data sekunder (Namira et al., 2024). Data ini mencakup data *Bank Customer Churn Dataset* yang bersumber dari data public Kaggle <https://www.kaggle.com/datasets/gauravtopre/bank-customer-churn-dataset/data> sebanyak 10.000 data dengan 12 variabel yang di bagi menjadi 11 fitur dan 1 kelas Pengumpulan data dilakukan dengan memastikan bahwa data yang diperoleh cukup *representatif* untuk membentuk model prediksi churn yang akurat.

Tabel 1. Fitur Dataset

No	Atribut	Description
1	<i>customer_id</i>	Nomor unik yang mengidentifikasi setiap pelanggan dalam dataset.
2	<i>credit_score</i>	Skor kredit pelanggan, menggambarkan kelayakan kredit mereka.
3	<i>Country</i>	Negara tempat pelanggan tinggal (France, Spain, Germany)
4	<i>Gender</i>	Jenis kelamin pelanggan, seperti Male atau Female.
5	<i>Age</i>	Usia pelanggan dalam satuan tahun.
6	<i>Tenure</i>	Lamanya pelanggan menjadi nasabah bank (dalam satuan tahun).
7	<i>Balance</i>	Jumlah uang yang dimiliki pelanggan di rekening mereka.
8	<i>products_number</i>	Jumlah produk atau layanan bank yang digunakan oleh pelanggan.
9	<i>credit_card</i>	Indikator apakah pelanggan mempunyai kartu kredit (1 = Ya, 0 = Tidak).
10	<i>active_member</i>	Status aktif nasabah sebagai anggota bank (1 = Aktif, 0 = Tidak aktif).
11	<i>estimated_salary</i>	Perkiraan pendapatan pelanggan per tahun dalam satuan mata uang lokal

Tabel 2. Kelas Dataset

No	Atribut	Description
1	<i>Churn</i>	Indikator apakah pelanggan berhenti berlangganan (1 = Berhenti, 0 = Tidak Berhenti).

Kelas *churn* (berhenti berlangganan dalam bahasa Indonesia) merupakan atribut target dalam suatu kumpulan data yang digunakan untuk memprediksi apakah seorang nasabah akan tetap menggunakan layanan perbankan atau berhenti menjadi nasabah.

b. Preprocessing Data

Preprocessing Data awal data adalah fase penting dalam alur analisis data untuk memastikan kumpulan data bersih, terstruktur, dan siap digunakan oleh model *prediktif* (Farid Naufal et al., 2023). Proses ini mencakup beberapa langkah penting untuk menyelesaikan masalah kualitas data seperti nilai yang hilang, format yang salah, dan atribut yang tidak relevan.

1. Pemeriksaan dan Penanganan *Missing Value*

Dataset yang digunakan dalam penelitian ini melalui tahap pemeriksaan untuk memastikan bahwa setiap atribut tidak memiliki nilai yang hilang. Pemeriksaan dilakukan menggunakan fungsi dan menunjukkan bahwa semua kolom memiliki nilai lengkap tanpa data yang hilang, termasuk atribut utama seperti usia, jenis kelamin, keseimbangan, dan logout. Oleh karena itu, tidak diperlukan langkah imputasi atau penghapusan data untuk menangani nilai yang hilang. Kumpulan data ini dinyatakan bersih dan dapat digunakan untuk prapemrosesan data selanjutnya dapat di lihat pada gambar 1.

```
Data columns (total 12 columns):
#   Column              Non-Null Count  Dtype
---  -
0   customer_id          10000 non-null   int64
1   credit_score          10000 non-null   int64
2   country               10000 non-null   object
3   gender                10000 non-null   object
4   age                  10000 non-null   int64
5   tenure                10000 non-null   int64
6   balance               10000 non-null   float64
7   products_number       10000 non-null   int64
8   credit_card           10000 non-null   int64
9   active_member         10000 non-null   int64
10  estimated_salary       10000 non-null   float64
11  churn                 10000 non-null   int64
```

Gambar 2. Missing Value

2. Re-labeling Kolom

Re-labeling nama kolom dalam konteks penelitian ini berkaitan dengan perubahan bahasa melibatkan perubahan nama kolom atau atribut dalam suatu kumpulan dataset dari satu bahasa ke bahasa lain untuk menyesuaikan dengan kebutuhan pembaca atau pengguna kumpulan data tersebut (Taufik et al., 2024).

Dalam hal ini, nama kolom yang aslinya ditulis dalam bahasa Inggris telah diubah menjadi bahasa Indonesia untuk meningkatkan keterbacaan dan pemahaman, terutama bila penelitian ditujukan untuk publik.

Tabel 3. Transformation

NO	ARTIBUT (BAHASA INGGRIS)	ARTIBUT (BAHASA INDONESIA)
1	<i>customer_id</i>	Id_pelanggan
2	<i>credit_score</i>	Skor_kredit
3	<i>Country</i>	Negara
4	<i>Gender</i>	Jenis_kelamin
5	<i>Age</i>	Usia
6	<i>Tenure</i>	Masa aktif
7	<i>Balance</i>	Saldo
8	<i>products_number</i>	Jumlah_produk
9	<i>credit_card</i>	Kartu_kredit
10	<i>active_member</i>	Anggota_aktif
11	<i>estimated_salary</i>	Gaji_perkiraan
12	<i>Churn</i>	Berhenti_berlangganan

3. Numerisasi Data

Numerisasi atau istilahnya adalah Konversi Data Menjadi *Numerik* pada data nasabah bank, banyak fitur atau kolom yang bersifat kategoris atau teks (contoh ; "*Jenis_Gender*", atau "Negara"). Algoritme *mechine learning* biasanya hanya dapat memproses data numerik. Oleh karena itu, langkah ini penting untuk memastikan bahwa semua data berada dalam format yang dapat dipahami oleh model. Representasi numerik lebih efisien untuk perhitungan komputasi dibandingkan dengan data teks.

3.1. Label Encoding

Pada penelitian ini digunakan untuk mengubah kategori menjadi angka dan menetapkan nomor unik untuk setiap kategori. Untuk kolom jenis_kelamin, untuk "Laki-laki" di beri label 0 dan "Perempuan" di beri label 1. Di samping itu, di dalam kolom yang berkaitan dengan "Negara", terdapat kategori-kategori seperti *Spain*, *Germany*, dan *France*. Dalam hal ini, diganti dengan menggunakan metode Label Encoding dengan cara memberikan label numerik yang sesuai setiap kategori tersebut. Di *Spain* di beri label 0, di *Germany* di beri label 1, dan di *France* beri label 2. Dengan ni memungkinkan model *machine learning* untuk memahami dan memproses data tersebut dengan lebih baik.

c. Data Splitting

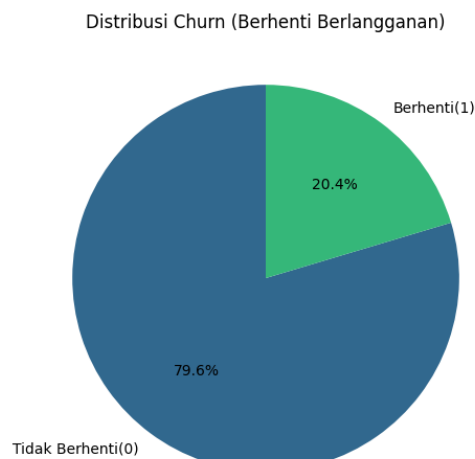
Pada tahap ini, data dibagi menjadi dua bagian: data latih untuk membangun algoritma berdasarkan data, dan data uji untuk mengevaluasi model berdasarkan data. Secara total, 70% dari total data digunakan untuk data pelatihan dan 30% dari data digunakan untuk data pengujian.

Tabel 4. Split Data

Set Data	Jumlah Fitur	Jumlah Target	Distribusi Target (0)	Distribusi Target (1)
Data Latih	70 %	70%	5574	1426
Data Uji	30%	30%	2389	611

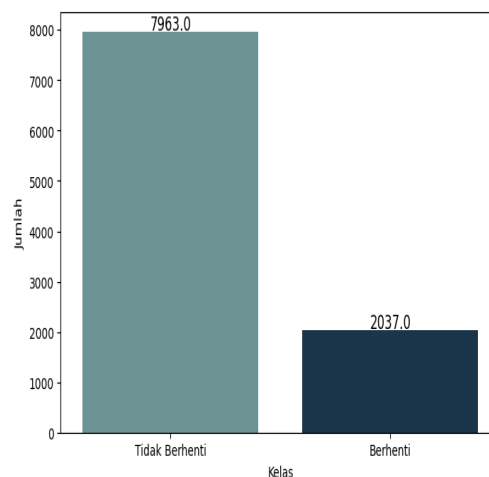
d. Metode Oversampling

Masalah ketidakseimbangan kelas sering terjadi dalam analisis data, khususnya dalam memprediksi *customer churn*, dimana jumlah pelanggan yang berhenti berlangganan (kelas minoritas) jauh lebih sedikit dari pada yang tetap berlangganan (kelas mayoritas) (Syukron et al., 2023). Ketidakseimbangan ini bisa membuat model machine learning cenderung memihak pada kelas mayoritas, sehingga mengakibatkan prediksi kurang akurat untuk kelas minoritas yang lebih penting (Nur Azizah et al., 2023). Karena itu, perlu adanya tindakan untuk mengatasi ketidakseimbangan kelas dengan menerapkan teknik *oversampling*.



Gambar 3. Distribusi *churn*

Oversampling adalah metode menambah jumlah sampel kelas minoritas agar lebih mendekati kelas mayoritas. Dua teknik *oversampling* yang umum digunakan adalah *SMOTE* (*Synthetic Minority Oversampling Technique*) dan *random oversampling*.

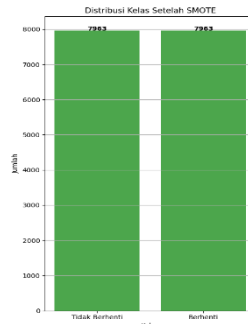


Gambar 4. Distribusi Kelas sebelum Over

1. SMOTE (Synthetic Minority Oversampling Technique)

Penggunaan *SMOTE* (Synthetic Minority Oversampling Technique) sangat penting dalam penelitian ini karena dataset yang digunakan untuk memprediksi customer churn di sektor perbankan sangat tidak seimbang. Kelas minoritas (pelanggan yang churn) seringkali jauh lebih kecil dibandingkan kelas mayoritas (pelanggan yang tidak churn) (Chieka & Kurniasih, 2023).

Jika ketidakseimbangan ini tidak ditangani dengan baik, model pembelajaran mesin akan cenderung mengabaikan kelas minoritas dan fokus pada prediksi kelas mayoritas, sehingga menghasilkan akurasi yang menipu. *SMOTE* memberikan solusi dengan membuat data sintetis baru untuk populasi yang kurang terwakili (pelanggan yang churn) daripada sekadar mereplikasi data yang sudah ada (Syukron et al., 2023). Setelah sudah di *SMOTE* akan memucul kesimbangan pada data kelas yang dapat di lihat pada tabel berikut.



Gambar 5. Distribusi Setelah SMOTE

2. Random Over-Sampling (ROS)

ROS (Random Over Sampling) adalah teknik untuk menangani masalah class imbalance dalam data, yaitu ketika salah satu kelas jauh lebih sedikit dibandingkan kelas lainnya. Ketidakseimbangan ini sering menyebabkan model pembelajaran mesin lebih cenderung memprediksi kelas mayoritas, sehingga mengurangi performa model pada kelas minoritas. Dalam ROS, data dari kelas minoritas diperbanyak dengan cara menyalin sampel yang sudah ada secara acak hingga jumlahnya menjadi seimbang dengan kelas mayoritas. Teknik ini tidak menghasilkan data baru, tetapi hanya memperbanyak data yang sudah ada.

e. Klasifikasi Algoritma Machine Learning

Beberapa algoritme pembelajaran *machine learning* yang paling umum digunakan untuk prediksi churn meliputi *Random Forest* dan *Naive Bayes* (Anggelia et al., 2024). Penelitian menunjukkan bahwa menggabungkan beberapa model atau model ansambel sering kali memberikan hasil prediksi yang lebih baik (Hartawan et al., 2023).

1. Random Forests

Random Forests sangat efektif karena dapat menangkap hubungan kompleks antar fitur dan mengurangi risiko *overfitting*, sehingga cocok untuk data yang tidak seimbang seperti prediksi churn. Random Forest bisa lebih berat dalam hal komputasi, terutama jika data yang digunakan sangat besar atau memiliki banyak fitur. Secara keseluruhan, Random Forest merupakan pilihan yang sangat baik untuk memprediksi churn karena kemampuannya menangani data yang kompleks dan memberikan hasil yang lebih akurat dan stabil.

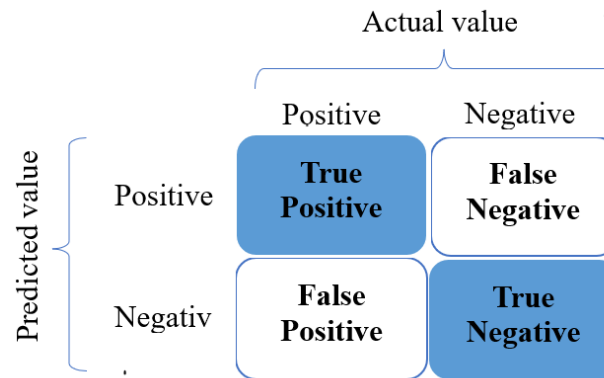
2. Naive Bayes

Naive Bayes Classifier (NBC) adalah proses klasifikasi probabilistik sederhana yang berkaitan dengan teori Bayesian. Naive Bayes adalah algoritma klasifikasi yang digunakan untuk memprediksi kategori atau kelas berdasarkan data yang ada. Algoritma ini bekerja dengan cara menghitung probabilitas dari setiap kelas berdasarkan fitur-fitur yang dimiliki oleh data. Naive Bayes menjadi model yang cepat dan efisien dalam memproses data. Teori

menyatakan bahwa meskipun di dunia nyata jarang terjadi, tetap saja Naive Bayes memberikan hasil yang memuaskan dalam berbagai situasi praktis. (Munir et al., 2024).

3. Evaluasi Model

Dalam penelitian data latih yang dibuat mengalami tahap pengujian yang mencakup skor akurasi dan nilai skor F1 (Silfana & Barata, 2024). Nilai akurasi diperoleh dari prediksi yang benar sebanyak data positif dan negatif dari total data. Skor F1 menunjukkan bahwa Model memiliki presisi dan *recall* yang baik.



Gambar 6. Confusion Matrix

True Positive (TP): Kasus churn yang diprediksi benar sebagai churn.

False Positive (FP): Kasus tidak churn yang diprediksi salah sebagai churn.

False Negative (FN): Kasus churn yang diprediksi salah sebagai tidak churn.

True Negative (TN): Kasus tidak churn yang diprediksi benar sebagai tidak churn.

a. Accuracy (Akurasi)

Akurasi dapat menggambarkan keakuratan model klasifikasi yang digunakan.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad [1]$$

b. Precision (Presisi):

Precision mengukur proporsi prediksi churn yang benar dari semua prediksi *churn*.

$$Precision = \frac{TP}{TP + FP} \quad [2]$$

c. Recall

Recall menggambarkan hasil dari model klasifikasi yang di gunakan dalam menemukan kembali sebuah informasi

$$Recall = \frac{TP}{TP + FN} \quad [3]$$

d. F1 Score

F1 Score adalah perhitungan kombinasi dari nilai *precision* dan nilai *recall* yang kemudian hasilnya disebut sebagai nilai pengukuran.

$$F1\ Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad [4]$$

e. Perbandingan Metriks Evaluasi Model

Ini adalah proses membandingkan kinerja model prediktif menggunakan metrik

evaluasi yang berbeda. Penting untuk memahami bagaimana model memproses data dan apakah model tersebut memenuhi tujuan analitis. Terkait dengan prediksi churn pelanggan.

HASIL DAN PEMBAHASAN

Pada bagian ini akan membahas tentang hasil eksperimen yang sudah dilakukan menggunakan aplikasi Google colab dengan membuat pemodelan tanpa menggunakan Metode Oversampling dan yang menggunakan metode *SMOTE (Synthetic Minority Over-sampling Technique)* serta Random Over-sampling. Pemodelan ini dilakukan dengan menggunakan berbagai algoritma klasifikasi untuk prediksi *customer churn* di sektor perbankan.

A. Dataset

	id_pelanggan	skor_kredit	negara	jenis_kelamin	usia	masa_aktif	saldo	jumlah_produk	kartu_kredit	anggota_aktif	gaji_perkiraan	berhenti_berlangganan
0	15634602	619	0	0	42	2	0.00	1	1	1	101348.88	1
1	15647311	608	2	0	41	1	83807.86	1	0	1	112542.58	0
2	15619304	502	0	0	42	8	159660.80	3	1	0	113931.57	1
3	15701354	699	0	0	39	1	0.00	2	0	0	93826.63	0
4	15737888	850	2	0	43	2	125510.82	1	1	1	79084.10	0
...	...	...	...	...	...	...	...	...	...	...	...	...
9995	15606229	771	0	1	39	5	0.00	2	1	0	96270.64	0
9996	15569892	516	0	1	35	10	57369.61	1	1	1	101699.77	0
9997	15584532	709	0	0	36	7	0.00	1	0	1	42085.58	1

Gambar 7. Dataset

B. Hasil Evaluasi algoritma Random Forest dan Neive Bayes Sebelum dan Setelah Metode Oversampling

Hasil evaluasi performa algoritma *Random Forest* dibandingkan sebelum dan setelah metode oversampling diterapkan disajikan dalam tabel berikut. Metode oversampling yang digunakan adalah *Random Over-Sampling* dan *SMOTE*

Tabel 5. Algoritma random Forest

Random Forest				
Label	Precision	Recall	F1 Score	Accuracy
0	88.22%	96.98%	92.39%	87.13%
1	78.78%	46.40%	58.41%	

Tabel 6. Algoritma Random Forest Dengan ROS

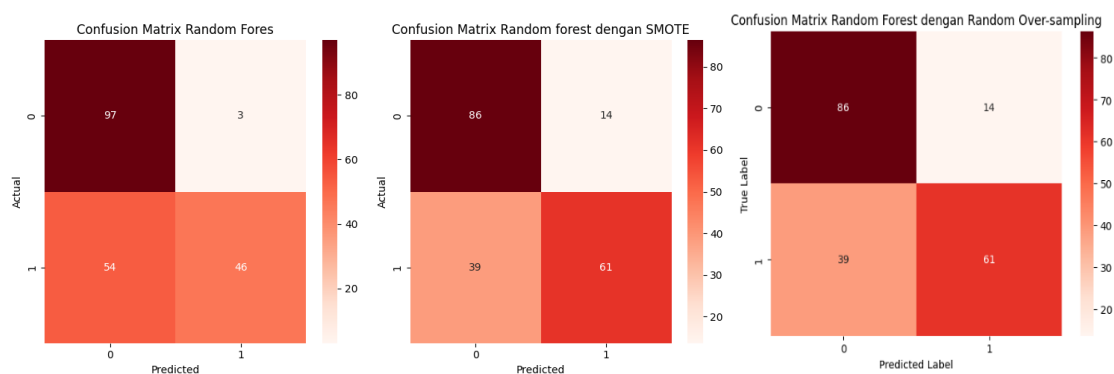
Random Forest dengan Random Over-sampling				
Label	Precision	Recall	F1 Score	Accuracy
0	89.25%	94.21%	91.66%	86.20%
1	68.89%	53.08%	59.96%	

Tabel 7. Algoritma Random Forest Dengan SMOTE

Random Forest dengan SMOTE				
Label	Precision	Recall	F1 Score	Accuracy
0	90.19%	86.38%	88.25%	81.47%
1	52.04%	61.13%	56.22%	

Berdasarkan hasil evaluasi model Random Forest yang diterapkan pada data churn

pelanggan dengan tiga metode oversampling Random Forest, Random Over-sampling, dan SMOTE dapat disimpulkan bahwa *Random Forest* memberikan akurasi tertinggi sebesar 87.13%. Meskipun akurasi ini cukup tinggi, model ini cenderung lebih baik dalam memprediksi kelas non-churn dan kurang efektif dalam mendeteksi churn. Sementara itu, *Random Over-sampling* mengalami sedikit penurunan dalam akurasi menjadi 86.20%, namun memberikan keseimbangan yang lebih baik dalam mendeteksi *churn*, meskipun akurasi secara keseluruhan sedikit menurun. *SMOTE*, di sisi lain, menunjukkan penurunan yang lebih signifikan dalam akurasi dengan nilai 81.47%, yang mencerminkan pengorbanan dalam akurasi demi meningkatkan kemampuan model dalam mendeteksi *churn*. Secara keseluruhan, jika hanya mengutamakan akurasi, *Random Forest* adalah yang paling unggul di antara ketiga metode tersebut.



Gambar 8. Confusion Matrix Algoritma

Hasil evaluasi performa algoritma *Neive Bayes* dibandingkan sebelum dan setelah *metode oversampling* diterapkan disajikan dalam tabel berikut. *Metode oversampling* yang digunakan adalah *Random Over-Sampling* dan *SMOTE*.

Tabel 8. Algoritma Neive Bayes

NEIVE BAYES				
Label	Precision	Recall	F1 Score	Accuracy
0	81.18%	96.56%,	88.20%	79.20%
1	34.13%	7.36%	12.11%	

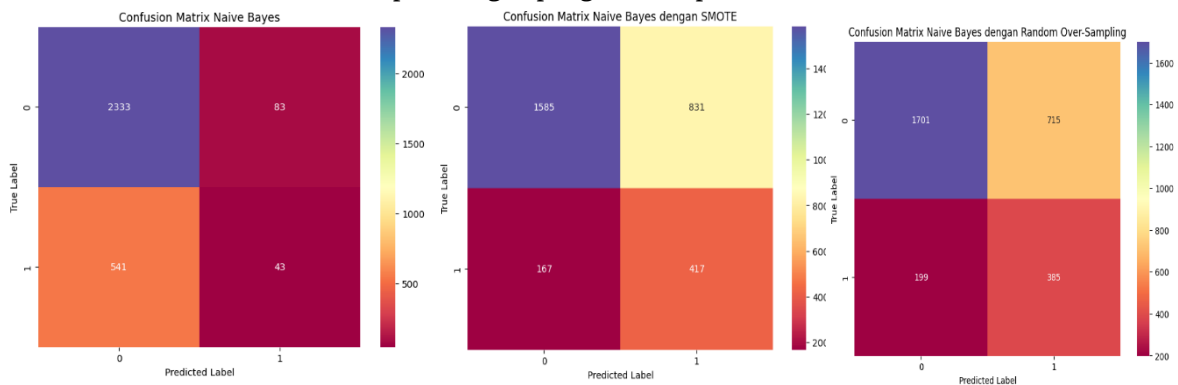
Tabel 9. Algoritma Neive Bayes Dengan SMOTE

Neive Bayes dengan SMOTE				
Label	Precision	Recall	F1 Score	Accuracy
0	90.47%	65.60%	76.06%	66.73%
1	33.41%	71.40%	45.52%	

Tabel 10. Algoritma Neive Bayes Dengan Random Over-sampling

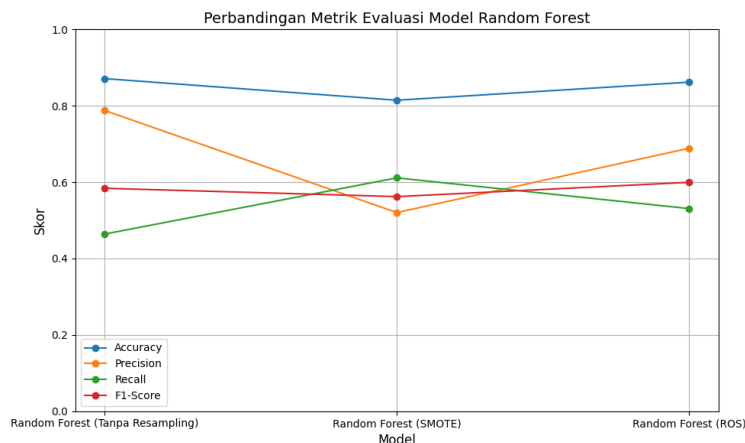
Neive Bayes dengan Random Over-sampling				
Label	Precision	Recall	F1 Score	Accuracy
0	89.53%	70.41%	78.82%	69.53%
1	35.00%	65.92%	45.72%	

Model *Neive Bayes Biasa (tanpa oversampling)* menunjukkan akurasi terbaik (79.20%), karena model ini sangat baik dalam memprediksi kelas mayoritas (*non-churn*). Namun, meskipun akurasi tinggi, model ini tidak efektif dalam mendeteksi churn (kelas minoritas), karena cenderung mengabaikan pelanggan churn. Sementara itu, setelah penerapan *Random Over-sampling*, akurasi model menurun menjadi 69.53%. Penurunan ini disebabkan oleh meningkatnya prediksi false positives, yaitu *non-churn* yang diprediksi sebagai churn, sehingga mengurangi akurasi keseluruhan. Akhirnya, dengan penggunaan *SMOTE*, akurasi model turun lebih jauh menjadi 66.73%, meskipun *recall* untuk churn meningkat. Penurunan akurasi ini terjadi karena jumlah *false positives* semakin meningkat, dan ketidakseimbangan antara *precision* dan *recall* semakin jelas. Secara keseluruhan, meskipun *Neive Bayes Biasa* memberikan akurasi tertinggi, penerapan *Random Over-sampling* dan *SMOTE* lebih efektif dalam meningkatkan kemampuan model untuk mendeteksi churn, meskipun dengan pengorbanan pada akurasi.

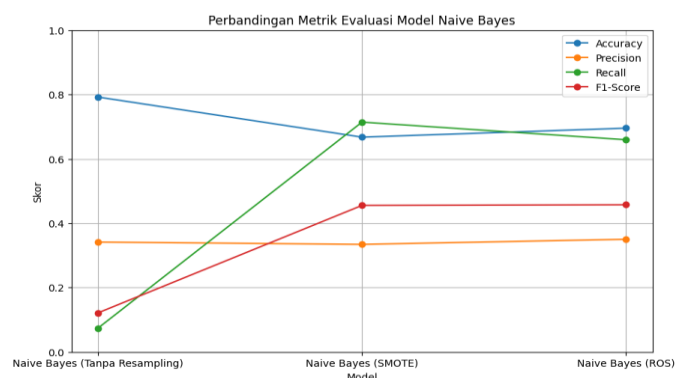


Gambar 9. Confusion Matrix Algoritma

C. Perbandingan Metrik Evaluasi Mode



Gambar 10. Perbandingan Matrix Algoritma



Gambar 11. Perbandingan Matrix Algoritma

Pada gambar 10 menyajikan valuasi kinerja model Random Forest berdasarkan nilai *precision*, *recall*, dan *F1-Score* menggunakan tiga metode pendekatan yang berbeda: tanpa *resampling*, *SMOTE*, serta *ROS*. Tanpa melakukan *resampling*, model memberikan hasil yang tidak seimbang dimana *precision* memiliki performa yang sangat baik, namun *recall* dan *F1-Score* rendah karena kelas minoritas tidak terdeteksi dengan baik. Dengan menggunakan metode *SMOTE*, tingkat *recall* mengalami peningkatan, namun sedikit terjadi penurunan pada tingkat *precision*. Pada sisi lain, *ROS* menunjukkan keseimbangan yang lebih baik dengan pertumbuhan yang terlihat pada semua ukuran.

Gambar 11 membandingkan empat metrik utama, yaitu *akurasi*, *presisi*, *recall*, dan *F1-Score*. Tanpa melakukan *resampling*, meskipun akurasi model mencapai tingkat tertinggi, namun kinerja pada kelas minoritas masih menunjukkan performa buruk karena tingkat *recall* yang rendah. *SMOTE* berkontribusi dalam peningkatan nilai *recall*, walaupun akurasi sedikit mengalami penurunan. *ROS* telah terbukti menjadi metode yang sangat optimal karena mampu mencapai keseimbangan yang sempurna antara semua metrik dengan tingkat performa yang stabil dan konsisten.

KESIMPULAN

Hasil penelitian ini menunjukkan bahwa *Random Forest* dapat mencapai akurasi tertinggi dan menangani kumpulan data yang tidak seimbang tanpa menerapkan teknik pengambilan sampel yang berlebihan. Namun, penerapan teknik oversampling meningkatkan kemampuan deteksi churn tetapi menurunkan keakuratan *Random Forest*. Sebaliknya, *Neive Bayes* menunjukkan akurasi yang tinggi dalam memprediksi kelas mayoritas, namun memiliki pengaruh yang kecil dalam mendeteksi churn, yang dapat mengakibatkan hilangnya nasabah yang berharga bagi bank. Keuntungan dari *Random Forest* terletak pada kemampuannya menangkap hubungan kompleks antar fitur, mengurangi risiko *overfitting*, dan meningkatkan keandalan dalam konteks data yang tidak seimbang.

Namun, kelemahan kedua algoritma ini adalah rendahnya efektivitasnya dalam mendeteksi kelas minoritas, yang penting untuk strategi retensi pelanggan. Hal ini menunjukkan bahwa meskipun keakuratan model secara keseluruhan tinggi, model tersebut masih gagal mengidentifikasi pelanggan yang berhenti.

SARAN

Untuk pengembangan selanjutnya, disarankan untuk mengeksplorasi algoritma lain seperti *Gradient Boosting* atau *XGBoost*, serta meningkatkan teknik preprocessing dan menerapkan metode oversampling alternatif. Analisis fitur yang lebih mendalam juga perlu dilakukan untuk mengidentifikasi faktor-faktor yang paling berpengaruh terhadap churn pelanggan.

DAFTAR PUSTAKA

- Abdullah, M., & Widarto, N. U. (2024). Analisis Strategi E-Commerce Pada PT. Spada Inovasi Digital Menggunakan Metode It Balanced Scorecard. *Jurnal Sistem Informasi dan Informatika (Simika)*, 7(1), 73–81.
- Anggelia, D., Riti, Y. F., & Siswanto, P. W. (2024). Analisis Perbandingan Metode Arima Dan Least Square Untuk Prediksi Harga Emas : Pendekatan Probabilistik Dan Statistik. *Jurnal Sistem Informasi Dan Informatika (Simika)*, 7(1), 95–103. <https://doi.org/10.47080/simika.v7i1.3197>
- Anjani, A. F., Anggraeni, D., & Tirta, I. M. (2023). Implementasi Random Forest Menggunakan SMOTE untuk Analisis Sentimen Ulasan Aplikasi Sister for Students UNEJ. *Jurnal Nasional Teknologi Dan Sistem Informasi*, 9(2), 163–172. <https://doi.org/10.25077/teknosi.v9i2.2023.163-172>
- Azmi, A. F. (2024). Prediksi Churn Nasabah Bank Menggunakan Klasifikasi Random Forest Dan Decision Tree Dengan Evaluasi Confusion Matrix. *Komputa : Jurnal Ilmiah Komputer Dan Informatika*, 13(1), 111–119. <https://ojs.unikom.ac.id/index.php/komputa/article/view/12639>
- Chieka, M., & Kurniasih, A. (2023). Teknik SMOTE untuk Mengatasi Imbalance Class dalam Klasifikasi Bank Customer Churn Menggunakan Algoritma Naïve bayes dan Logistic

- Regression. *Prosiding Seminar Nasional Mahasiswa Bidang Ilmu Komputer Dan Aplikasinya*, 4(2), 552–559.
- Christanto, R., & Wiwik, A. (2024). Analisis Churn Pelanggan Menggunakan Pembelajaran Mesin untuk Meningkatkan Retensi Pelanggan di Vissie Net. *Jurnal Internasional Penelitian dan Manajemen Ilmiah (IJSRM)*, September, 7379–7387. <https://doi.org/10.18535/ijsrm/v12i09.em05>
- Farid Naufal, M., Fernando Susanto, A., Nathaneil Kansil, C., Huda, S., & kunci, K. (2023). Analisis Perbandingan Algoritma Machine Learning untuk Prediksi Potensi Hilangnya Nasabah Bank Application of Machine Learning to Predict Potential Loss of Bank Customer. *Techno*, 22(1), 1–11.
- Hartawan, M. S., Erkamim, M., Rachmawati, S., Santi, N. C., Legito, L., & Sepriano, S. (2023). Penerapan Algoritma Supervised Learning untuk Klasifikasi Program Keluarga Harapan. *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 3(2), 83–91. <https://doi.org/10.57152/malcom.v3i2.873>
- Hermawan, A. (2024). Membangun Model Prediksi Churn Pelanggan yang Akurat. *Jurnal Riset Sistem Informasi dan Teknik Informatika*, 2(6), 68–80. <https://doi.org/https://doi.org/10.61132/mercurius.v2i6.398>
- Istiqomawati, R., Quraisy, M., Widiyastuti, A., & Yogyakarta, S. (2022). Pengaruh Switching Cost terhadap Customer Retention di Bank Syariah. *Jurnal Kajian Akutansi Dan Keuangan*, 2(2), 7–14. <https://doi.org/10.56393/pacioli.v2i2.1346>
- Munir, A. S., Saputra, A. B., Aziz, A., & Barata, M. A. (2024). Perbandingan Akurasi Algoritma Neive Bayes dan Algoritma Decision Tree dalam Pengklasifikasian Penyakit Kanker Payudara. *Jurnal Ilmiah Informatika Global*, 15(1), 23–29. <https://doi.org/10.36982/jiig.v15i1.3578>
- Namira, Slamet, I., & Susanto, I. (2024). Prediksi Nasabah Churn Dengan Algoritma Decision Tree, Random Forest Dan Support Vector Machine. *Escaf*, 27(2) 1045–1053. <https://help.sap.com>
- Novitasari, D. T., Barata, M. A., Rochmatin, N. N., Muzakka, M. A., & Andiyani, P. (2024). Analisis Penerapan Program Reward Kepada Customer Menggunakan Metode Clustering. *Jurnal Bisnis Kolega*, 10(1), 29–35. <https://doi.org/10.57249/jbk.v10i1.140>
- Nugroho, A., & Rilvani, E. (2024). Penerapan Metode Oversampling SMOTE Pada Algoritma Random Forest Untuk Prediksi Kebangkrutan Perusahaan Application of the SMOTE Oversampling Method to the Random Forest Algorithm for Predicting Company Bankruptcy. *Februari*, 22(1), 207–214.
- Nur Azizah, A., Falach Asy'ari, M., Wisma Dwi Prastya, I., & Purwitasari, D. (2023). Easy Data Augmentation untuk Data yang Imbalance pada Konsultasi Kesehatan Daring. *Jurnal Teknologi Informasi Dan Ilmu Komputer*, 10(5), 1095–1104. <https://doi.org/10.25126/jtiik.20231057082>
- Silfana, F. I., & Barata, M. A. (2024). Using K-NN Algorithm for Evaluating Feature Selection on High. *Jurnal Teknik Informatika*, 17(2), 191-202. <https://doi.org/10.15408/jti.v17i2.40866>
- Syarif, M., & Nugraha, W. (2023). Mwmote Dalam Mengatasi Ketidakseimbangan Kelas Pada Prediksi Churn Menggunakan Klasifikasi C4.5. *JATI (Jurnal Mahasiswa Teknik Informatika)* 7(1), 54–62.
- Syukron, A., Sardiarinto, S., Saputro, E., & Widodo, P. (2023). Penerapan Metode Smote Untuk Mengatasi Ketidakseimbangan Kelas Pada Prediksi Gagal Jantung. *Jurnal Teknologi Informasi Dan Terapan*, 10(1), 47–50. <https://doi.org/10.25047/jtit.v10i1.313>
- Taufik, R., Jimah, R., & Solichin, A. (2024). Implementasi dan Analisis Model Machine Learning Decision Tree untuk Deteksi Akun Palsu di Twitter. *Jurnal Media Informatika Budidarma*, 8(2), 797. <https://doi.org/10.30865/mib.v8i2.7548>