

## Prediksi Keputusan Reload Pelanggan Perusahaan Telekomunikasi X Menggunakan Machine Learning

### ARTICLE INFO

### Abstrak

#### Article history

*Received: June 5, 2026*

*Revised: June 30, 2026*

*Accepted: June 30, 2026*

#### Kata Kunci

Ensemble Learning Keputusan  
Reload SMOTE  
Telekomunikasi  
XGBoost

Percepatan kemajuan sektor telekomunikasi mengharuskan perusahaan untuk memahami perilaku konsumen demi mempertahankan loyalitas dan meningkatkan pendapatan. Kendala utama yang dihadapi adalah pengelolaan data dalam memprediksi keputusan reload pelanggan pada perusahaan telekomunikasi. Penelitian ini bertujuan untuk memprediksi keputusan reload pelanggan menggunakan pendekatan Ensemble Learning. Tahapan penelitian meliputi pengumpulan data transaksi historis, pra-pemrosesan, penanganan ketidakseimbangan kelas menggunakan teknik Synthetic Minority Over-sampling Technique (SMOTE), pemilihan fitur, serta evaluasi model. Algoritma yang digunakan mencakup Random Forest, Gradient Boosting, XGBoost, dan Stacking Ensemble. Kinerja model dievaluasi berdasarkan metrik Akurasi, Presisi, Recall, dan F1-Score. Hasil analisis pemilihan fitur menunjukkan bahwa masa aktif berlangganan (*tenure\_rgu*) merupakan faktor paling dominan yang memengaruhi keputusan reload pelanggan, diikuti oleh saldo saat ini dan riwayat total pengisian ulang dalam 30 hari terakhir. Pada tahap evaluasi model, XGBoost menghasilkan kinerja klasifikasi terbaik dengan tingkat akurasi tertinggi sebesar 74%, sedangkan Gradient Boosting memiliki recall tertinggi sebesar 50%. Temuan penelitian ini diharapkan dapat membantu manajemen perusahaan dalam merancang strategi pemasaran yang lebih akurat dan efisien untuk meningkatkan kepuasan pelanggan.

---

ABSTRACT

---

**Keywords**  
Ensemble Learning  
Reload Decision SMOTE  
Telecommunication  
XGBoost

*The accelerated progress of the telecommunications sector requires companies to understand consumer behavior in order to maintain loyalty and increase revenue. The main challenge faced is managing data to predict customer reload decisions at telecommunication companies. This study aims to predict these reload decisions using an Ensemble Learning approach. The research stages include historical transaction data collection, preprocessing, handling class imbalance using the SMOTE technique, feature selection, and model evaluation. The algorithms used include Random Forest, Gradient Boosting, XGBoost, and Stacking Ensemble. Model performance was evaluated based on Accuracy, Precision, Recall, and F1-Score metrics. The results of the feature selection analysis show that active subscription period (tenure\_rgu) is the most dominant factor influencing customer reload decisions, followed by the current balance and the history of total reloads in the last 30 days. In the model evaluation stage, XGBoost produced the best classification performance with the highest accuracy rate of 74%, while Gradient Boosting achieved the highest recall at 50%. The findings of this study are expected to assist company management in designing more accurate and efficient marketing strategies to improve customer satisfaction.*

---

## 1. Pendahuluan

Kemajuan teknologi komunikasi yang pesat telah meningkatkan penggunaan layanan seluler sehingga perusahaan telekomunikasi dituntut untuk memahami perilaku pelanggan guna mempertahankan loyalitas serta meningkatkan pendapatan. Salah satu perilaku pelanggan yang memiliki dampak langsung terhadap pendapatan perusahaan adalah keputusan reload, yaitu keputusan pelanggan untuk melakukan pengisian ulang pulsa atau paket layanan setelah masa penggunaan sebelumnya berakhir. Keputusan reload mencerminkan keberlanjutan penggunaan layanan sekaligus menjadi indikator loyalitas pelanggan karena menunjukkan bahwa pelanggan masih aktif memanfaatkan layanan yang disediakan perusahaan. Oleh karena itu, kemampuan memprediksi keputusan reload menjadi informasi yang penting bagi perusahaan dalam menyusun strategi pemasaran, pemberian promosi yang tepat sasaran, serta upaya mempertahankan pelanggan.

Perusahaan telekomunikasi menghasilkan jutaan data transaksi setiap hari yang berasal dari aktivitas penggunaan layanan, riwayat pengisian ulang, pola konsumsi data, durasi penggunaan, serta karakteristik pelanggan lainnya. Data tersebut memiliki karakteristik berupa volume yang sangat besar, terdiri atas berbagai jenis variabel numerik maupun kategorikal, serta mengandung hubungan antarvariabel yang kompleks dan cenderung nonlinier. Selain itu, perilaku pelanggan juga bersifat dinamis karena dipengaruhi oleh kebiasaan penggunaan layanan, kondisi ekonomi, hingga program promosi yang diberikan perusahaan. Karakteristik tersebut menyebabkan proses identifikasi pelanggan yang berpotensi melakukan reload menjadi semakin menantang apabila hanya menggunakan pendekatan analisis konvensional.

Berbagai penelitian sebelumnya telah menerapkan algoritma Machine Learning, seperti Logistic Regression dan Support Vector Machine (SVM), untuk memprediksi perilaku pelanggan pada industri telekomunikasi. Namun, metode-metode tersebut memiliki keterbatasan dalam menangkap hubungan nonlinier dan interaksi kompleks antarfitur ketika dihadapkan pada data transaksi berskala besar, sehingga performa prediksi dapat menurun pada dataset yang kompleks [1], [2]. Di sisi lain, sebagian besar penelitian lebih banyak berfokus pada prediksi *customer churn*, sedangkan penelitian mengenai prediksi keputusan reload pelanggan masih relatif terbatas. Padahal, keputusan reload merupakan indikator perilaku positif pelanggan yang dapat dimanfaatkan perusahaan untuk meningkatkan retensi pelanggan sebelum terjadi churn. Dengan demikian, masih terdapat peluang penelitian untuk mengembangkan model prediksi reload yang lebih akurat menggunakan algoritma Machine Learning yang sesuai dengan karakteristik data transaksi telekomunikasi.

Salah satu pendekatan yang sesuai untuk mengatasi karakteristik data tersebut adalah Ensemble Learning, yaitu teknik Machine Learning yang mengombinasikan beberapa model (base learners) sehingga mampu meningkatkan kemampuan generalisasi model. Pendekatan ini efektif dalam menangani data berdimensi tinggi, hubungan nonlinier, interaksi antarfitur yang kompleks, serta ketidakseimbangan distribusi data yang sering dijumpai pada data pelanggan telekomunikasi. Algoritma seperti Random Forest, Gradient Boosting, dan XGBoost juga memiliki kemampuan mengurangi varians maupun bias model sehingga menghasilkan prediksi yang lebih stabil dibandingkan penggunaan model tunggal [3]. Oleh karena itu, pendekatan Ensemble Learning dinilai sesuai untuk membangun model prediksi keputusan reload pelanggan.

Berdasarkan uraian tersebut, penelitian ini bertujuan membangun dan mengevaluasi model Machine Learning berbasis Ensemble Learning untuk memprediksi keputusan reload pelanggan pada perusahaan telekomunikasi X, di mana keputusan reload didefinisikan sebagai kondisi ketika pelanggan melakukan pengisian ulang pulsa atau paket layanan pada periode prediksi berikutnya. Penelitian ini juga membandingkan kinerja beberapa algoritma Ensemble Learning untuk menentukan model dengan performa terbaik sehingga dapat menjadi dasar bagi perusahaan dalam menyusun strategi pemasaran, pemberian promosi yang lebih tepat sasaran, serta meningkatkan retensi dan pendapatan pelanggan.

## 2. Landasan Teori

### 2.1 Keputusan Reload Pelanggan

Keputusan *reload* merupakan keputusan pelanggan untuk melakukan pengisian ulang pulsa maupun paket data setelah masa penggunaan sebelumnya berakhir. Dalam industri telekomunikasi, keputusan ini menjadi indikator keberlanjutan penggunaan layanan karena menunjukkan bahwa pelanggan masih aktif menggunakan produk yang ditawarkan operator. Berbeda dengan *customer churn* yang menggambarkan pelanggan berhenti menggunakan layanan, keputusan *reload* merepresentasikan perilaku positif pelanggan

yang berkontribusi langsung terhadap pendapatan perusahaan. Oleh karena itu, kemampuan memprediksi keputusan *reload* memungkinkan perusahaan mengidentifikasi pelanggan yang berpotensi melakukan pengisian ulang maupun pelanggan yang memerlukan strategi promosi agar tetap menggunakan layanan.

Perilaku *reload* dipengaruhi oleh berbagai faktor, seperti frekuensi penggunaan layanan, jumlah transaksi sebelumnya, lama penggunaan kartu, jenis paket yang digunakan, pengeluaran pelanggan, serta karakteristik demografis. Hubungan antarvariabel tersebut bersifat kompleks dan sering kali tidak linear sehingga sulit dimodelkan menggunakan pendekatan statistik konvensional. Oleh karena itu, diperlukan metode yang mampu mempelajari pola perilaku pelanggan berdasarkan data historis sehingga menghasilkan prediksi yang lebih akurat.

## 2.2 Machine Learning untuk Prediksi Perilaku Pelanggan

Machine Learning merupakan cabang Artificial Intelligence yang memungkinkan komputer mempelajari pola dari data sehingga mampu menghasilkan prediksi atau pengambilan keputusan tanpa diprogram secara eksplisit. Berdasarkan jenis pembelajarannya, Machine Learning dibedakan menjadi supervised learning, unsupervised learning, dan reinforcement learning. Pada penelitian ini digunakan pendekatan supervised learning karena data latih telah memiliki label keputusan reload, yaitu pelanggan melakukan reload atau tidak melakukan reload pada periode prediksi.

Dalam industri telekomunikasi, Machine Learning telah banyak diterapkan untuk berbagai permasalahan seperti prediksi customer churn, segmentasi pelanggan, deteksi penipuan (fraud detection), rekomendasi layanan, hingga prediksi perilaku pelanggan. Kemampuannya dalam mengolah data berskala besar serta menangkap hubungan yang kompleks antarvariabel menjadikan Machine Learning sebagai pendekatan yang sesuai untuk menganalisis data transaksi pelanggan yang dihasilkan setiap hari. Namun demikian, performa model sangat dipengaruhi oleh karakteristik data dan algoritma yang digunakan sehingga diperlukan metode yang memiliki kemampuan generalisasi yang baik.

## 2.3 Ensemble Learning

*Ensemble Learning* merupakan pendekatan *Machine Learning* yang menggabungkan beberapa model pembelajaran (base learners) menjadi satu model prediksi yang lebih kuat (strong learner). Prinsip dasar metode ini adalah mengombinasikan hasil prediksi dari beberapa model sehingga mampu meningkatkan akurasi, mengurangi varians maupun bias, serta menghasilkan kemampuan generalisasi yang lebih baik dibandingkan model tunggal. Pendekatan ini banyak digunakan pada permasalahan klasifikasi karena mampu meningkatkan stabilitas model ketika dihadapkan pada data yang kompleks.

Data pelanggan telekomunikasi memiliki karakteristik berupa volume data yang besar, kombinasi atribut numerik dan kategorikal, hubungan antarfitur yang bersifat nonlinier, serta kemungkinan distribusi kelas yang tidak seimbang. Kondisi tersebut menyebabkan model tunggal sering mengalami penurunan performa atau overfitting. Ensemble Learning dipilih pada penelitian ini karena mampu mengatasi karakteristik tersebut melalui kombinasi beberapa model sehingga menghasilkan prediksi yang lebih robust dan stabil.

Secara umum, prediksi Ensemble Learning dinyatakan sebagai:

$$\hat{y} = \Phi(f_1(x), f_2(x), \dots, f_M(x)). \quad (1)$$

dengan  $\hat{y}$  merupakan hasil prediksi akhir,  $f_m(x)$  adalah prediksi model dasar ke- $m$ ,  $M$  menyatakan jumlah *base learner*, sedangkan  $\Phi(\cdot)$  merupakan fungsi agregasi yang menggabungkan seluruh hasil prediksi.

### 1. Random Forest

Random Forest merupakan algoritma berbasis teknik *bagging* yang membangun sejumlah *decision tree* menggunakan data hasil *bootstrap sampling*. Setiap pohon menghasilkan prediksi secara independen, kemudian seluruh prediksi digabungkan menggunakan mekanisme *majority voting* pada kasus klasifikasi. Pendekatan ini mampu mengurangi varians model sehingga menghasilkan prediksi yang lebih stabil.

Bagging bertujuan mengurangi varians dari model. Fungsi agregasinya menggunakan rata-rata atau *majority voting*:

$$\hat{y}_{bag} = \frac{1}{M} \sum_{m=1}^M f_m(x). \quad (2)$$

dengan  $M$  merupakan jumlah pohon keputusan dan  $f_m(x)$  adalah prediksi dari pohon ke- $m$ .

Setiap node pada *decision tree* dalam Random Forest dipilih berdasarkan penurunan impurity, umumnya menggunakan Gini Index:

$$Gini(D) = 1 - \sum_{i=1}^c p_i^2. \quad (3)$$

dengan  $D$  merupakan himpunan data pada suatu node,  $p_i$  adalah proporsi data pada kelas ke- $i$ , sedangkan  $c$  menyatakan jumlah kelas.

### 2. XGBoost

XGBoost (*Extreme Gradient Boosting*) merupakan pengembangan dari Gradient Boosting yang menambahkan mekanisme regularisasi untuk mengendalikan kompleksitas model. Selain itu, XGBoost menerapkan optimasi komputasi, *tree pruning*, dan *parallel processing* sehingga mampu menghasilkan waktu komputasi yang lebih cepat serta mengurangi risiko *overfitting*. Fungsi objektif XGBoost dirumuskan sebagai

$$Obj(\theta) = \sum_{i=1}^n l(y_i, \hat{y}_i) + \sum_{m=1}^M \Omega(f_k). \quad (4)$$

Dengan  $l(y_i, \hat{y}_i)$  sebagai fungsi loss antara nilai aktual dan hasil prediksi dan  $\Omega(f_k)$  sebagai fungsi regularisasi yang mengontrol kompleksitas model.

#### 2.4 Evaluasi Model

Kinerja model dievaluasi menggunakan *confusion matrix*, yaitu matriks yang membandingkan hasil prediksi dengan label aktual. Dari matriks tersebut diperoleh empat komponen utama, yaitu *True Positive* (TP), *True Negative* (TN), *False Positive* (FP), dan

*False Negative* (FN). Berdasarkan keempat komponen tersebut dihitung nilai Accuracy, Precision, Recall, dan F1-Score.

$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (5)$$

Accuracy menunjukkan proporsi prediksi yang benar terhadap seluruh data, Precision mengukur ketepatan model dalam mengklasifikasikan kelas positif, sedangkan Recall menunjukkan kemampuan model menemukan seluruh data positif. Untuk memperoleh evaluasi yang seimbang, penelitian ini menggunakan F1-Score yang merupakan rata-rata harmonis antara Precision dan Recall.

dengan

$$Precision = \frac{TP}{TP + FP}$$

dan

$$Recall = \frac{TP}{TP + FN}$$

di mana  $TP$  adalah jumlah data positif yang diprediksi benar,  $FP$  adalah data negatif yang diprediksi positif, sedangkan  $FN$  adalah data positif yang diprediksi negatif.

## 2.5 Penelitian Terdahulu

Beberapa penelitian telah menunjukkan bahwa Machine Learning mampu menghasilkan performa yang baik pada berbagai permasalahan di industri telekomunikasi. Penelitian oleh [4] menggunakan kombinasi Ensemble SMOTE-ENN dan XGBoost untuk memprediksi *customer churn* pada dataset yang tidak seimbang. Hasil penelitian menunjukkan bahwa pendekatan tersebut mampu meningkatkan akurasi hingga 96%, sehingga efektif dalam menangani permasalahan ketidakseimbangan kelas. Namun, penelitian tersebut hanya berfokus pada prediksi pelanggan yang berhenti menggunakan layanan (*churn*) dan belum mengkaji perilaku pelanggan yang tetap aktif melalui keputusan *reload*.

Selanjutnya, penelitian oleh [5] membandingkan Random Forest, Gradient Boosting Machine (GBM), dan XGBoost untuk memprediksi perilaku pelanggan telekomunikasi. Hasil penelitian menunjukkan bahwa metode Ensemble Learning menghasilkan performa yang lebih baik dibandingkan model tunggal karena mampu menangkap hubungan nonlinier antarfitur. Meskipun demikian, penelitian tersebut belum membahas karakteristik keputusan *reload* sebagai salah satu indikator loyalitas pelanggan.

Penelitian lain yang dilakukan oleh [6] menerapkan Random Forest pada data pelanggan telekomunikasi dan menunjukkan bahwa algoritma berbasis pohon memiliki kemampuan yang baik dalam mengolah data berskala besar dan berdimensi tinggi. Akan tetapi, penelitian tersebut hanya mengevaluasi satu algoritma sehingga belum memberikan gambaran mengenai perbandingan performa beberapa metode Ensemble Learning.

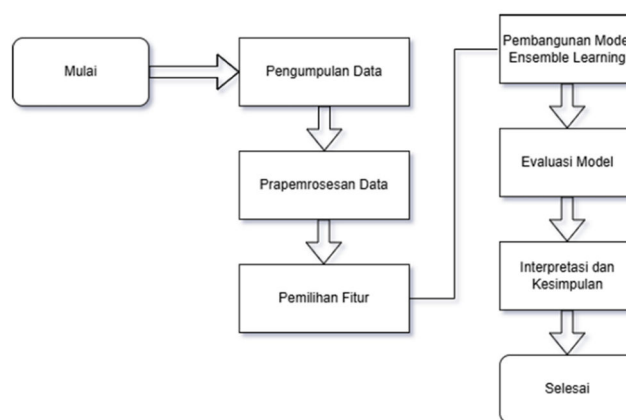
Berdasarkan penelitian-penelitian tersebut dapat disimpulkan bahwa sebagian besar penelitian masih berfokus pada prediksi *customer churn*, sedangkan penelitian mengenai prediksi keputusan *reload* pelanggan masih relatif terbatas. Selain itu, belum banyak penelitian yang membandingkan performa Random Forest, Gradient Boosting, dan

XGBoost secara khusus untuk memprediksi keputusan *reload* pada data pelanggan telekomunikasi. Oleh karena itu, penelitian ini mengisi kesenjangan tersebut dengan membandingkan ketiga algoritma Ensemble Learning untuk memperoleh model yang paling efektif dalam memprediksi keputusan *reload* pelanggan perusahaan telekomunikasi X.

### 3. Metodologi Penelitian

#### 3.1 Gambaran Umum Penelitian

Pada penelitian ini dilakukan dengan beberapa tahapan. Tahapan tersebut dimulai dari mengidentifikasi masalah, pengambilan data, pre-processing data, pemilihan dan pelatihan model, evaluasi model, hingga proses analisis dan hasil.



Gambar 1. Diagram Alir Penelitian

a. Pengumpulan data

Tahap ini melibatkan data pelanggan perusahaan telekomunikasi yang berisi informasi mengenai durasi berlangganan, jumlah dan denominasi isi ulang, saldo sisa, dan keterlibatan pelanggan selama interval tertentu. Data ini diperoleh dari sistem internal perusahaan telekomunikasi.

b. Pra-Pemrosesan data

Pada tahap ini dilakukan pembersihan data (data cleaning), meliputi mengatasi missing value dan menghapus duplikasi.

c. Pemilihan Fitur

Pemilihan fitur ini bertujuan untuk mengurangi kompleksitas model dan meningkatkan akurasi prediksi. Teknik yang digunakan dapat berupa feature importance (pada model ensemble seperti Random Forest atau XGBoost).

d. Pembangunan Model

Machine Learning Pada tahap ini dilakukan pembangunan model prediksi menggunakan beberapa algoritma Ensemble Learning, seperti:

- Random Forest Classifier
- Gradient Boosting Classifier
- XGBoost Classifier

Model dilatih menggunakan data latih (training data) dan divalidasi dengan data uji (testing data).

## e. Evaluasi Model

Model yang telah dibangun dievaluasi menggunakan metrik performa seperti Accuracy, Precision, Recall, F1-Score. Hasil evaluasi akan menunjukkan model mana yang paling efektif dalam memprediksi keputusan reload pelanggan.

## f. Interpretasi dan Kesimpulan

Tahap terakhir melibatkan interpretasi hasil model dan penarikan kesimpulan. Ini faktor utama yang mempengaruhi keputusan reload serta rekomendasi bagi perusahaan untuk meningkatkan kepuasan pelanggan.

### 3.2 Dataset Penelitian

Dataset yang digunakan dalam penelitian ini terdiri dari total 36.870 catatan pelanggan, mencakup 27 variabel yang menjelaskan keterlibatan pelanggan, status akun, dan riwayat transaksi reload pelanggan. Meskipun dataset terdiri dari 25 variabel, penelitian ini tidak menggunakan seluruh variabel sebagai input model. Untuk meningkatkan performa prediksi dan mengurangi kompleksitas model, dilakukan pemilihan fitur menggunakan metode feature importance. Proses ini bertujuan memilih variabel yang paling berpengaruh terhadap keputusan reload pelanggan.

| No | Variabel           | Tipe Data               |
|----|--------------------|-------------------------|
| 1  | current_tier       | Kategorikal             |
| 2  | available_points   | Numerik (Float)         |
| 3  | vlr_attached_p3d   | Numerik (Integer)       |
| 4  | flag_arpu_90d      | Kategorikal             |
| 5  | flag_arpu_last_30d | Kategorikal             |
| 6  | tenure_rgu         | Numerik (Integer)       |
| 7  | rgu_flag           | Kategorikal             |
| 8  | rld_30d            | Numerik (Integer)       |
| 9  | rld_60d            | Numerik (Integer)       |
| 10 | rld_90d            | Numerik (Integer)       |
| 11 | rld_tot_30d        | Numerik (Float)         |
| 12 | rld_tot_60d        | Numerik (Float)         |
| 13 | rld_tot_90d        | Numerik (Float)         |
| 14 | reload_p90d        | Numerik (Float)         |
| 15 | tot_month_rld      | Numerik (Float)         |
| 16 | denom_30d          | Numerik (Float)         |
| 17 | denom_60d          | Numerik (Float)         |
| 18 | denom_90d          | Numerik (Float)         |
| 19 | curr_balance       | Numerik (Float)         |
| 20 | active_pack        | Numerik (Integer/Float) |
| 21 | status             | Kategorikal             |
| 22 | n_days             | Numerik (Float)         |
| 23 | arpu_rld           | Numerik (Float)         |
| 24 | cust_flag          | Kategorikal             |
| 25 | rld_nm (Target)    | Numerik (Integer)       |

**Gambar 2.** Variabel Penelitian

Sebelum dilakukan pemodelan, dilakukan analisis terhadap sebaran kelas pada variabel target. Hasil distribusi menunjukkan adanya ketidakseimbangan data (class imbalance) sebagai berikut:



| Status Reload Bulan Berikutnya | Jumlah |
|--------------------------------|--------|
| Tidak Reload (0)               | 40.000 |
| Reload (1)                     | 20.000 |

**Gambar 2.** Variabel Penelitian

Berdasarkan gambar 2 tersebut, terlihat bahwa hanya sekitar 33% dari total pelanggan yang melakukan reload kembali di bulan berikutnya.

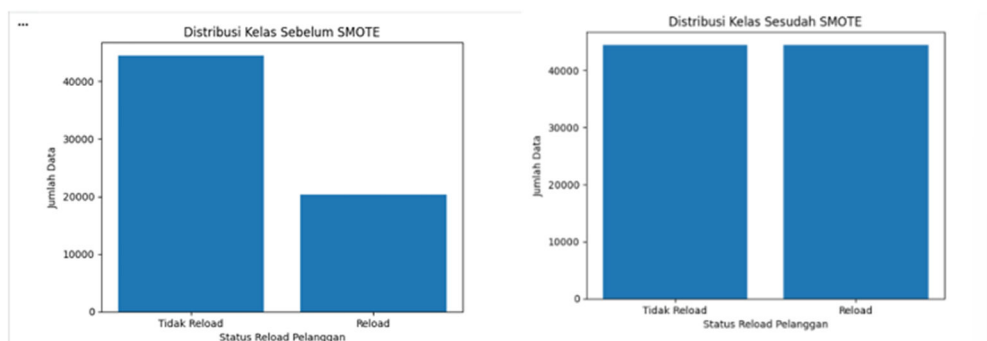
### 3.3 Data Preprocessing

Tahap pra-pemrosesan data diawali dengan proses pembersihan data (data cleaning) untuk memastikan kualitas dataset sebelum digunakan dalam pembangunan model machine learning. Proses ini meliputi pemeriksaan nilai hilang (missing value) dan penghapusan data duplikat.

Hasil pemeriksaan menunjukkan bahwa setelah proses pembersihan dilakukan, seluruh variabel pada dataset telah memiliki nilai missing sebesar nol. Hal ini menunjukkan bahwa data telah bersih dan siap digunakan untuk proses analisis lebih lanjut. Nilai hilang pada variabel numerik diisi menggunakan nilai median, sedangkan variabel kategorikal diisi menggunakan nilai modus sehingga tidak mengganggu distribusi data.

Selain itu, data duplikat yang ditemukan juga dihapus untuk mencegah bias pada model akibat pengulangan data pelanggan yang sama. Setelah seluruh proses cleaning dilakukan tidak ada lagi missing value pada setiap variabel, dataset akhir yang digunakan dalam penelitian berjumlah 998.934 data pelanggan dengan 25 variabel aktif.

Pada dataset awal, ditemukan adanya ketidakseimbangan kelas (class imbalance) yang signifikan antara pelanggan yang melakukan reload dan yang tidak. Berdasarkan visualisasi pada Gambar 4.1, kelas mayoritas ('Tidak Reload') memiliki sampel di atas 40.000, sementara kelas minoritas ('Reload') hanya berjumlah sekitar 20.000 sampel, sehingga membentuk rasio sekitar 2:1. Kondisi ini berisiko menyebabkan model mengalami bias; sistem mungkin menghasilkan akurasi tinggi secara keseluruhan, namun gagal secara fungsional dalam mengidentifikasi pelanggan yang benar-benar akan melakukan reload (low recall pada kelas minoritas). Hal ini sangat merugikan bagi tujuan strategis perusahaan telekomunikasi yang membutuhkan akurasi prediksi pada perilaku reload pelanggan.

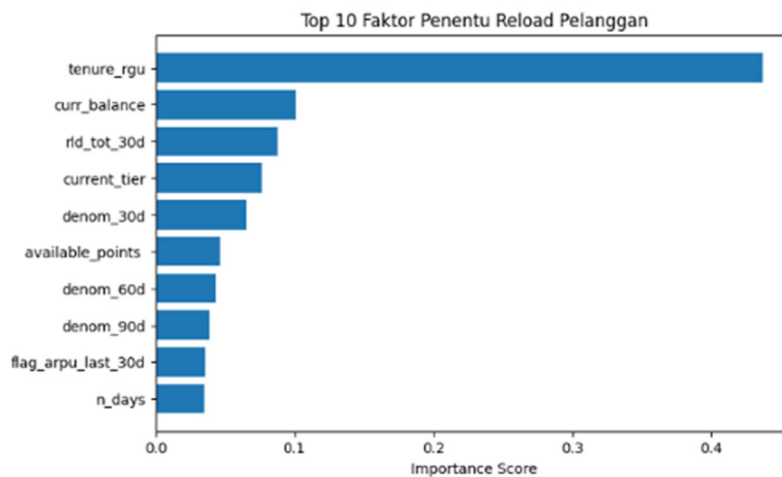


**Gambar 3.** Distribusi kelas sebelum dan sesudah SMOTE

Untuk mengatasi permasalahan tersebut, diterapkan metode Synthetic Minority Over-sampling Technique (SMOTE) pada data latih. Teknik ini bekerja dengan menghasilkan data sintetis pada kelas minoritas, sehingga jumlah sampel menjadi seimbang tanpa sekadar menduplikasi data yang sudah ada. Dengan distribusi data yang telah seimbang (balanced), model diharapkan dapat mempelajari karakteristik kedua kelas secara lebih adil, sehingga performa metrik evaluasi seperti presisi, recall, dan F1-Score dapat meningkat secara optimal.

### 3.4 Pemilihan Fitur

Berdasarkan gambar tersebut, terlihat bahwa variabel `tenure_rgu` merupakan faktor yang paling dominan dalam memengaruhi keputusan reload pelanggan dibandingkan variabel lainnya. Fitur ini memiliki skor kepentingan (importance score) tertinggi, yakni mencapai lebih dari 0,4. Hal ini mengindikasikan bahwa masa aktif pelanggan atau durasi pelanggan menggunakan layanan merupakan indikator terkuat; semakin lama loyalitas pelanggan terbentuk, semakin tinggi pula probabilitas perilaku reload yang dapat diprediksi oleh model.



**Gambar 4.** Faktor penentu Reload pelanggan

Faktor berikutnya yang memberikan kontribusi signifikan adalah `curr_balance` dan `rld_tot_30d`. Keberadaan saldo saat ini serta total aktivitas reload dalam 30 hari terakhir menunjukkan bahwa pola konsumsi pulsa jangka pendek sangat menentukan apakah pelanggan akan melakukan pengisian kembali dalam waktu dekat.

Selain itu, variabel seperti `current_tier` dan `denom_30d` juga masuk dalam jajaran fitur utama, yang menandakan bahwa status segmentasi pelanggan dan riwayat nominal pengisian sebelumnya memiliki pengaruh yang stabil terhadap model prediksi. Sementara itu, fitur-fitur seperti `available_points`, riwayat nominal jangka panjang (`denom_60d` dan `denom_90d`), hingga `n_days` berada di posisi bawah, yang berarti kontribusinya tetap diperhitungkan namun tidak sekuat faktor loyalitas dan saldo terkini.

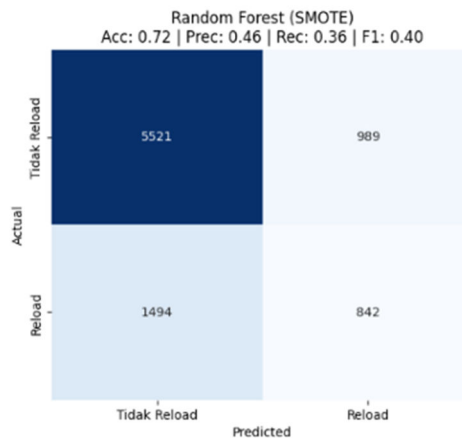
## 4. Hasil dan Pembahasan

Berdasarkan pengujian yang telah dilakukan pada sistem *customer reload prediction* yang telah dibuat akan dijelaskan sebagai berikut.

#### 4.1 Hasil Model Ensemble Learning

##### 1. Random Forest

Pada model algoritma Random Forest dengan SMOTE di tahap evaluasi, data testing menjadi input untuk diproses dan menghasilkan nilai prediksi. Data tersebut diproses sesuai konsep algoritma Random Forest dengan model yang telah dibangun menggunakan teknik SMOTE untuk menyeimbangkan kelas data pada tahap training. Kemudian, data hasil prediksi dicocokkan dengan data asli dari data testing seperti pada Gambar 4.7.

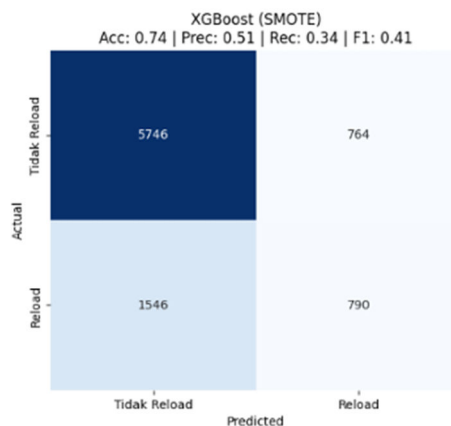


**Gambar 5.** Konfusi Matriks Random Forest

Dari Gambar 5 dapat dilihat bahwa penggunaan teknik SMOTE memberikan hasil yang identik dengan versi tanpa SMOTE pada model ini. Algoritma Random Forest dengan SMOTE menghasilkan accuracy (akurasi), recall, precision, dan f1-score sebesar 72%, 36%, 46%, dan 40%.

##### 2. XGBoost

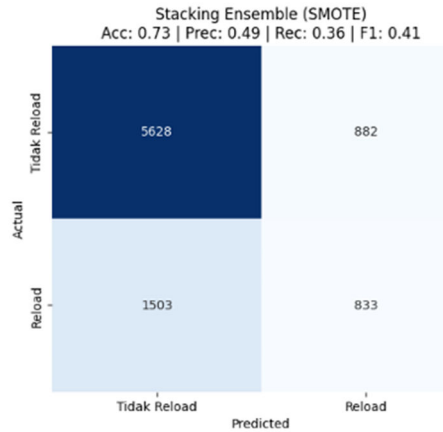
Pada model algoritma XGBoost dengan SMOTE di tahap evaluasi, data testing menjadi input untuk diproses dan menghasilkan nilai prediksi. Data tersebut diproses sesuai konsep algoritma XGBoost yang merupakan pengembangan dari Gradient Boosting yang dioptimalkan melalui regularisasi, namun pada tahap ini model dibangun dengan data training yang telah diproses menggunakan teknik SMOTE untuk menangani ketidakseimbangan kelas. Kemudian, data hasil prediksi dicocokkan dengan data asli dari data testing seperti pada Gambar 4.9



**Gambar 6.** Konfusi Matriks XGBoost

### 3. Stacking Ensemble

Pada model algoritma Stacking Ensemble dengan smote di tahap evaluasi, data testing menjadi input untuk diproses dengan mengombinasikan beberapa model melalui model meta. Data diproses dengan model yang telah dibangun sebelumnya menggunakan teknik SMOTE. Kemudian, data hasil prediksi dicocokkan dengan data asli seperti pada Gambar 4. 10.



**Gambar 7.** Konfusi Stacking Ensemble

Dari Gambar 7 dapat dilihat bahwa model Stacking memberikan hasil yang sangat stabil dan seimbang. Algoritma Stacking Ensemble dengan smote menghasilkan accuracy (akurasi), recall, precision, dan f1-score sebesar 73%, 36%, 49%, dan 41%.

#### 4.2 Kinerja Model Ensemble Learning

Berikut adalah tabel Evaluasi Model Ensemble Learning yang menggunakan algoritma Random Forest, dan XGBoost:

**Tabel 1.** Evaluasi Model *Ensemble Learning*

| No | Model             | Accuracy | Precision | Recall | F1-Score |
|----|-------------------|----------|-----------|--------|----------|
| 1  | Random Forest     | 72%      | 46%       | 36%    | 40%      |
| 2  | XGBoost           | 74%      | 51%       | 34%    | 41%      |
| 3  | Stacking Ensemble | 73%      | 49%       | 36%    | 41%      |

Berdasarkan Tabel 1, *XGBoost* merupakan algoritma yang memberikan performa klasifikasi terbaik dari sisi Akurasi (74%) dan Presisi (51%). Kemampuannya yang sangat baik dalam mengidentifikasi pelanggan yang tidak aktif serta keberhasilannya dalam menekan tingkat kesalahan prediksi keliru membuktikan efisiensi algoritma ini dalam menangani data berdimensi besar dan mencegah overfitting melalui mekanisme regularisasinya yang kuat [7]. Akurasi sebesar 74% dianggap sebagai pencapaian yang sangat realistis dan solid, mengingat keputusan pelanggan merupakan variabel dinamis yang juga sangat dipengaruhi oleh faktor eksternal (seperti kondisi ekonomi atau aksi kompetitor) yang tidak tercakup dalam dataset ini.

Di sisi lain, model *Gradient Boosting* mencatatkan nilai *Recall* tertinggi sebesar 50%, dibarengi dengan skor keseimbangan *F1-Score* sebesar 48%. *Metrik Recall* yang tinggi

mengindikasikan bahwa model ini merupakan pendekatan yang paling sensitif untuk menangkap sinyal atau mendeteksi sebanyak mungkin pelanggan aktual yang benar-benar akan melakukan pengisian ulang. Bagi perusahaan telekomunikasi, meminimalkan *False Negative* (pelanggan yang diprediksi tidak *reload* padahal sebenarnya *reload*) dengan mengandalkan model sensitif seperti *Gradient Boosting* dapat secara signifikan mengurangi risiko hilangnya peluang pendapatan (*revenue loss*).

Model *Stacking Ensemble* juga memberikan prediksi yang seimbang dan stabil dengan tingkat akurasi 73%, menunjukkan bahwa penggabungan beberapa *base learners* sukses dalam menyaring pelanggan yang tidak aktif sembari tetap mendeteksi pelanggan potensial.

#### 4.3 Analisis Feature Importance

Analisis *feature importance* dilakukan untuk mengetahui kontribusi masing-masing variabel terhadap proses prediksi keputusan *reload*. Hasil analisis menunjukkan bahwa variabel *rld\_30* memiliki tingkat kepentingan tertinggi dibandingkan variabel lainnya. Hal ini menunjukkan bahwa aktivitas *reload* pelanggan dalam 30 hari terakhir merupakan indikator paling kuat dalam menentukan kemungkinan pelanggan melakukan *reload* pada periode berikutnya.

Variabel *rld\_60* juga memberikan kontribusi yang cukup besar terhadap model, namun pengaruhnya lebih rendah dibandingkan *rld\_30*. Sementara itu, variabel *rld\_90* memiliki tingkat kepentingan yang lebih kecil, yang mengindikasikan bahwa riwayat transaksi yang lebih lama memiliki pengaruh yang lebih rendah terhadap keputusan *reload*. Temuan ini menunjukkan bahwa perilaku pelanggan yang lebih baru lebih relevan dibandingkan riwayat transaksi yang telah lama terjadi.

Dari perspektif bisnis, hasil ini memberikan informasi bahwa perusahaan telekomunikasi perlu memberikan perhatian lebih kepada pelanggan yang dalam 30 hari terakhir menunjukkan penurunan aktivitas *reload*. Kelompok pelanggan tersebut dapat menjadi sasaran program promosi, pemberian insentif, atau penawaran paket yang sesuai untuk meningkatkan kemungkinan pelanggan melakukan pengisian ulang pada periode berikutnya. Sebaliknya, pelanggan yang diprediksi memiliki kecenderungan tinggi untuk melakukan *reload* dapat menjadi target strategi *cross-selling* atau *up-selling* guna meningkatkan nilai transaksi.

### 5. Kesimpulan

Berdasarkan hasil analisis, durasi berlangganan (*tenure\_rgu*), saldo akhir pelanggan (*curr\_balance*), dan riwayat akumulasi *reload* 30 hari (*rld\_tot\_30d*) merupakan prediktor paling dominan dalam memengaruhi keputusan pelanggan untuk melakukan isi ulang (*reload*) pada bulan berikutnya. Implementasi algoritma *Ensemble Learning* dengan bantuan teknik *balancing* SMOTE terbukti sangat andal dalam memprediksi perilaku tersebut. Model XGBoost memberikan akurasi klasifikasi tertinggi (74%) yang cocok digunakan untuk penargetan kampanye efisiensi tinggi, sementara *Gradient Boosting* menghasilkan sensitivitas (*Recall*) terbaik (50%) yang menjadikannya ideal untuk mengidentifikasi seluas mungkin populasi pelanggan aktif agar perusahaan tidak

kehilangan potensi pendapatan. Temuan ini dapat menjadi landasan strategis bagi pengambilan keputusan pemasaran berbasis data.

## Referensi

- [1] S. Rahmawati, A. Wibowo, and A. F. Nur Masruriyah, "Improving Diabetes Prediction Accuracy in Indonesia : A Comparative Analysis of SVM , Logistic Regression , and Naive Bayes with SMOTE and," *JURNAL RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 5, no. 158, pp. 7–10, 2024, doi: <https://doi.org/10.29207/resti.v8i5.5980>.
- [2] E. J. Otse, G. N. Obunadike, and A. Abubakar, "Linear Regression Approach to Solving Multicollinearity and Overfitting in Predictive Analysis," *Journal of Science Research and Reviews*, vol. 2, no. 1, pp. 108–117, 2025, doi: <https://doi.org/10.70882/josrar.2025.v2i1.35>.
- [3] Z. Liu, "Investigating an Ensemble Classifier Based on Multi-Objective Genetic Algorithm for Machine Learning Applications," *(IJACSA) International Journal of Advanced Computer Science and Applications*, vol. 15, no. 5, pp. 883–889, 2024.
- [4] S. Ouf, K. T. Mahmoud, and M. A. A. Fattah, "A proposed hybrid framework to improve the accuracy of customer churn prediction in telecom industry," *Journal of Big Data*, 2024, doi: 10.1186/s40537-024-00922-9.
- [5] L. Saha, H. K. Tripathy, T. Gaber, and H. El-gohary, "Deep Churn Prediction Method for Telecommunication Industry," *sustainability*, pp. 1–21, 2023.
- [6] U. N. Surabaya, "Prediction and Analysis Customer Churn at Telkomsel Using a Machine Learning Approach Achmad," *Journal of Emerging Information Systems and Business Intelligence*, vol. 6, no. 2, pp. 230–240, 2025.
- [7] T. Chen and C. Guestrin, "XGBoost : A Scalable Tree Boosting System," in *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 2016.