

Prediksi Customer Churn Berbasis Nilai Pelanggan Menggunakan Explainable Machine Learning Untuk Mendukung Strategi Retensi

Amat Damuri^{*1}, Novi Wulandari², Iwan Mulyana³

^{*1,2}Manajemen Informatika STMIK Al Muslim, Bekasi

³Ilmu Komputer STMIK Al Muslim, Bekasi

e-mail: ^{*1} amatdamuri@gmail.com, ²wulandari.novi1@gmail.com, ³ iwanmulyanamkom@gmail.com

Abstrak

Customer churn dapat menurunkan pendapatan dan nilai pelanggan pada industri telekomunikasi. Penelitian ini bertujuan mengembangkan model prediksi churn yang akurat dan dapat dijelaskan serta menentukan prioritas retensi berdasarkan risiko dan nilai pelanggan. Iranian Churn Dataset yang digunakan memuat 3.150 pelanggan dan 13 prediktor. Ketidakseimbangan kelas ditangani menggunakan SMOTENC, sedangkan Logistic Regression, Random Forest, dan XGBoost dioptimasi melalui GridSearchCV dengan validasi silang terstratifikasi. XGBoost memberikan performa terbaik dengan ROC-AUC pengujian 0,9912. Pada ambang 0,5444, model menghasilkan akurasi 96,83%, presisi 88,35%, recall 91,92%, dan F1-score 90,10% untuk kelas churn. Analisis SHAP mengidentifikasi frekuensi penggunaan, status, durasi penggunaan, kegagalan panggilan, lama berlangganan, dan keluhan sebagai prediktor utama. Segmentasi risiko dan Customer Value menetapkan 74 pelanggan atau 2,35% sebagai prioritas retensi tertinggi. Hasil ini menunjukkan bahwa interpretabilitas model dapat dihubungkan langsung dengan keputusan retensi pelanggan. Pendekatan ini mendukung alokasi sumber daya retensi secara transparan, terarah, dan efisien.

Kata Kunci: customer churn; explainable machine learning; nilai pelanggan; SHAP; strategi retensi

Abstract

Customer churn can reduce revenue and customer value in the telecommunications industry. This study develops an accurate and explainable churn prediction model and determines retention priorities based on customer risk and value. The Iranian Churn Dataset contains 3,150 customers and 13 predictors. Class imbalance was addressed using SMOTENC, while Logistic Regression, Random Forest, and XGBoost were optimized through GridSearchCV with stratified cross-validation. XGBoost achieved the best performance, with a test ROC-AUC of 0.9912. At a threshold of 0.5444, the model obtained 96.83% accuracy, 88.35% precision, 91.92% recall, and a 90.10% F1-score for the churn class. SHAP analysis identified frequency of use, customer status, seconds of use, call failures, subscription length, and complaints as the main predictors. Risk and Customer Value segmentation classified 74 customers, or 2.35%, as the highest retention priority. This approach supports transparent, targeted, and efficient retention resource allocation.

Keywords: customer churn; explainable machine learning; customer value; SHAP; retention strategy

I. PENDAHULUAN

Persaingan industri telekomunikasi meningkat seiring bertambahnya pilihan layanan, kemudahan pelanggan berpindah penyedia, serta tingginya tuntutan terhadap kualitas dan harga. Kondisi tersebut menyebabkan customer churn, yaitu keadaan ketika pelanggan menghentikan penggunaan layanan atau berpindah kepada penyedia lain, menjadi persoalan penting dalam pengelolaan hubungan pelanggan. Churn dapat menurunkan pendapatan, profitabilitas, dan keberlanjutan bisnis perusahaan

telekomunikasi (Suryana et al., 2021). Oleh karena itu, perusahaan perlu mengidentifikasi pelanggan yang berpotensi churn lebih awal agar intervensi retensi dapat dilakukan tepat waktu.

Pemanfaatan machine learning memungkinkan pola churn dipelajari dari data demografis, riwayat penggunaan, keluhan, durasi berlangganan, dan karakteristik layanan. Sejumlah penelitian di Indonesia telah menggunakan algoritma klasifikasi untuk memprediksi churn. Arifin (2015) menggabungkan information gain dan K-Nearest

Neighbor untuk prediksi churn telekomunikasi, sedangkan Alfarez et al. (2024) menerapkan Naïve Bayes pada pelanggan PT Hutchison 3 Indonesia. Iskandar dan Latifa (2023) juga mengimplementasikan prediksi churn dalam aplikasi berbasis web untuk mendukung retensi pelanggan.

Prediksi churn umumnya menghadapi ketidakseimbangan kelas karena jumlah pelanggan churn lebih sedikit daripada pelanggan yang tetap menggunakan layanan. Kondisi tersebut dapat menyebabkan model lebih banyak memprediksi kelas mayoritas. Suryana et al. (2021) menunjukkan bahwa kombinasi Synthetic Minority Over-sampling Technique (SMOTE) dan boosting dapat meningkatkan performa prediksi churn. Perbandingan algoritma juga diperlukan karena setiap model mempunyai karakteristik berbeda. Mufida et al. (2025) membandingkan Logistic Regression dan Random Forest, sedangkan Husein dan Harahap (2021) menggunakan pendekatan data science untuk menemukan pola churn pelanggan.

Akurasi prediksi belum sepenuhnya memenuhi kebutuhan keputusan bisnis karena model yang berperforma tinggi belum tentu menjelaskan alasan pelanggan diprediksi churn. Permasalahan ini semakin penting pada Random Forest, XGBoost, dan model ensemble yang cenderung bersifat black box. Explainable machine learning dapat menjelaskan kontribusi setiap variabel terhadap keluaran model, baik pada tingkat keseluruhan maupun pelanggan individual.

Salah satu metode interpretasi yang banyak digunakan adalah Shapley Additive Explanations (SHAP). Asif et al. (2025) menunjukkan bahwa penerapan SHAP dan LIME pada prediksi churn telekomunikasi dapat meningkatkan transparansi dan menghasilkan informasi untuk strategi retensi. Dalam konteks penelitian Indonesia, Astarani dan Putri (2025) mengintegrasikan XGBoost, GridSearchCV, SMOTE, dan SHAP, sedangkan Ardhani dan Tania (2025) membandingkan sejumlah algoritma dan menggunakan SHAP untuk menemukan faktor perilaku yang memengaruhi churn.

Selain probabilitas churn, nilai pelanggan perlu dipertimbangkan ketika menentukan prioritas retensi. Pelanggan dengan risiko churn tinggi tidak selalu memberikan nilai ekonomi yang sama. Jika retensi hanya didasarkan pada probabilitas churn,

perusahaan dapat mengalokasikan sumber daya kepada pelanggan bernilai rendah dan mengabaikan pelanggan bernilai tinggi. Verbeke et al. (2012) menegaskan bahwa pengelolaan churn sebaiknya berorientasi pada keuntungan, biaya intervensi, dan nilai pelanggan.

Penelitian ini menggunakan Iranian Churn Dataset yang memuat 3.150 pelanggan dan karakteristik seperti kegagalan panggilan, keluhan, durasi berlangganan, biaya, penggunaan layanan, kelompok usia, paket tarif, status, dan Customer Value (UCI Machine Learning Repository, 2020). Hossain (2025) telah menggunakan dataset yang sama untuk membandingkan sejumlah model, tetapi belum secara khusus mengintegrasikan nilai pelanggan, SHAP, dan penyusunan prioritas retensi dalam satu kerangka analisis.

Berdasarkan kesenjangan tersebut, penelitian ini bertujuan membangun dan mengevaluasi model machine learning untuk memprediksi customer churn, menjelaskan faktor yang memengaruhi prediksi menggunakan SHAP, serta mengelompokkan prioritas retensi berdasarkan kombinasi probabilitas churn dan Customer Value. Kebaruan penelitian terletak pada integrasi prediksi, interpretabilitas, dan matriks risiko-nilai untuk menghasilkan rekomendasi retensi yang lebih terarah.

II. METODE PENELITIAN

Desain Penelitian

Penelitian menggunakan pendekatan kuantitatif eksperimental. Tahapannya meliputi pengumpulan dan pemahaman data, prapemrosesan, pembagian data, penanganan ketidakseimbangan kelas, pelatihan dan optimasi model, evaluasi, interpretasi menggunakan SHAP, serta penyusunan prioritas retensi berdasarkan risiko churn dan nilai pelanggan. Model yang dibandingkan adalah Logistic Regression, Random Forest, dan XGBoost.

Sumber dan Variabel Data

Data diperoleh dari Iranian Churn Dataset pada UCI Machine Learning Repository dan tersedia melalui Kaggle. Dataset dikumpulkan selama 12 bulan. Variabel perilaku berasal dari sembilan bulan pertama, sedangkan status churn ditetapkan pada akhir bulan ke-12 sehingga tersedia planning gap

selama tiga bulan (UCI Machine Learning Repository, 2020).

Tabel 1. Variabel penelitian

Kelompok	Variabel	Keterangan
Kualitas layanan	Call Failure	Jumlah kegagalan panggilan
Pengaduan	Complains	Status keluhan pelanggan
Hubungan	Subscription Length	Lama pelanggan berlangganan
Finansial	Charge Amount	Tingkat biaya yang dibayarkan
Penggunaan	Seconds of Use	Total durasi penggunaan layanan
Penggunaan	Frequency of Use	Frekuensi penggunaan layanan
Penggunaan	Frequency of SMS	Frekuensi pengiriman SMS
Penggunaan	Distinct Called Numbers	Jumlah nomor berbeda yang dihubungi
Demografis	Age dan Age Group	Usia dan kelompok usia pelanggan
Layanan	Tariff Plan	Jenis paket tarif pelanggan
Aktivitas	Status	Status aktif atau tidak aktif
Nilai pelanggan	Customer Value	Nilai pelanggan dalam dataset
Target	Churn	1 = churn; 0 = tidak churn

Prapemrosesan dan Pembagian Data

Prapemrosesan mencakup pemeriksaan tipe data, nilai hilang, kombinasi fitur identik, dan distribusi kelas. Atribut identitas dikeluarkan apabila tersedia. Complains, Tariff Plan, Status, dan Age Group diperlakukan sebagai fitur kategoris, sedangkan fitur lainnya diperlakukan sebagai numerik. Standardisasi hanya diterapkan pada Logistic Regression. Data dibagi menjadi 80% data latih dan 20% data uji menggunakan stratified sampling dengan random_state 42.

Ketidakseimbangan kelas ditangani menggunakan SMOTENC. Oversampling hanya dilakukan di dalam fold data latih untuk mencegah data leakage. Optimasi hiperparameter dilakukan menggunakan GridSearchCV dengan Stratified 5-Fold Cross-Validation. Model dipilih berdasarkan ROC-AUC validasi, sedangkan recall dan F1-score

kelas churn digunakan sebagai pertimbangan tambahan

Tabel 2. Ruang pencarian hiperparameter

Model	Parameter	Nilai kandidat
Logistic Regression	C; penalty	0,1; 1; 10 — L1; L2
Random Forest	n_estimators; max_depth; min_samples_split	200; 500 — None; 10; 20 — 2; 5
XGBoost	n_estimators; learning_rate; max_depth; subsample	100; 200 — 0,05; 0,1 — 3; 5 — 0,8; 1,0

Evaluasi Model

Performa model diukur menggunakan accuracy, precision, recall, F1-score, dan ROC-AUC. Recall menjadi penting karena menunjukkan kemampuan model menemukan pelanggan yang benar-benar churn. Metrik dihitung berdasarkan true positive (TP), true negative (TN), false positive (FP), dan false negative (FN).

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (1)$$

$$Recall = \frac{TP}{TP+FN} \quad (2)$$

$$F1 - Score = 2X \frac{Precision \times Recall}{Precision+Recall} \quad (3)$$

Interpretasi dan Prioritas Retensi

Model terbaik diinterpretasikan menggunakan SHAP pada tingkat global dan individual. Nilai SHAP positif menunjukkan kontribusi ke arah churn, sedangkan nilai negatif menurunkan keluaran churn. Probabilitas churn untuk segmentasi dihasilkan menggunakan prediksi out-of-fold. Pelanggan dikategorikan berisiko tinggi jika probabilitasnya mencapai ambang optimal dan bernilai tinggi jika Customer Value mencapai atau melebihi median data latih.

Tabel 3. Matriks prioritas retensi

Risiko	Nilai	Prioritas	Strategi
Tinggi	Tinggi	I	Intervensi segera, penyelesaian keluhan, dan penawaran personal
Rendah	Tinggi	II	Program loyalitas dan pemeliharaan hubungan
Tinggi	Rendah	III	Retensi berbiaya rendah atau otomatis
Rendah	Rendah	IV	Pemantauan rutin

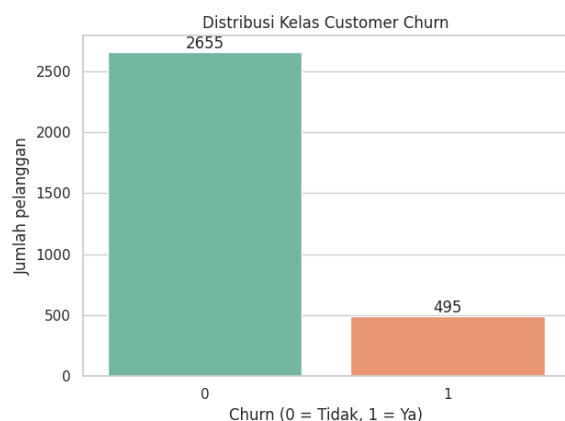
III. HASIL DAN PEMBAHASAN

Karakteristik Dataset

Dataset terdiri atas 3.150 data pelanggan, 13 prediktor, dan satu target churn. Seluruh variabel mempunyai data lengkap. Pemeriksaan menemukan 314 baris dengan kombinasi fitur identik yang tidak dihapus karena dataset tidak menyediakan identitas pelanggan dan kesamaan tersebut belum dapat dipastikan sebagai duplikasi pencatatan.

Tabel 4. Distribusi kelas customer churn

Status pelanggan	Jumlah	Persentase
Tidak churn	2.655	84,29%
Churn	495	15,71%
Total	3.15	100,00%



Gambar 1. Distribusi kelas customer churn

Kelas tidak churn berjumlah 84,29%, sedangkan kelas churn hanya 15,71%, dengan rasio sekitar 5,36:1. Ketimpangan tersebut mendukung penggunaan SMOTENC di dalam pipeline. Secara deskriptif, rata-rata lama berlangganan adalah 32,54 bulan dan median Customer Value adalah 228,48. Data latih berjumlah 2.520 dan data uji 630, dengan proporsi churn tetap 15,71% pada keduanya.

Perbandingan Model

Tabel 5. Hiperparameter terbaik setiap model

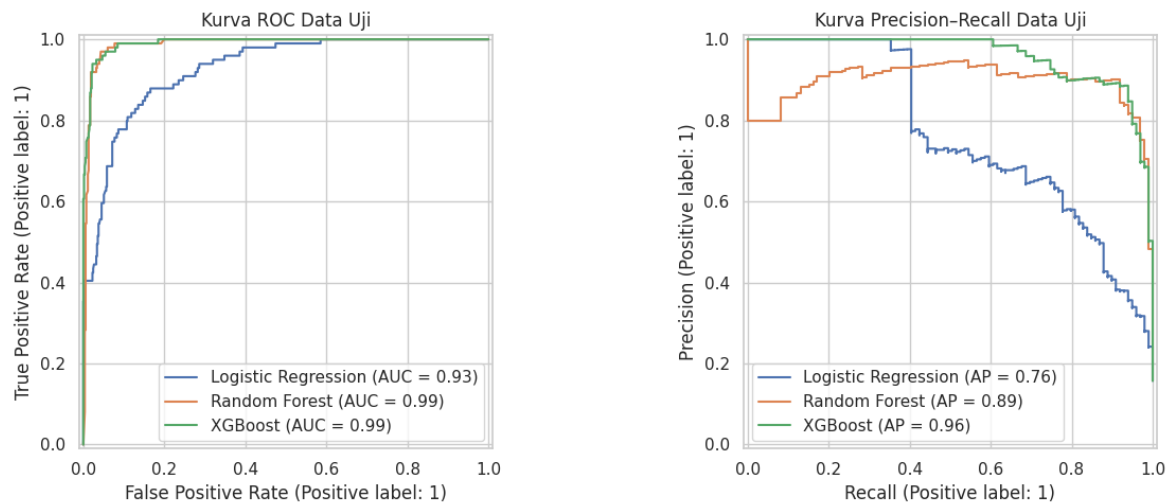
Model	Hiperparameter terbaik	ROC-AUC validasi
Logistic Regression	C = 10; penalty L1	0,9326
Random Forest	500 pohon; max_depth=None; min_samples_split=2	0,9837
XGBoost	200 estimator; learning_rate=0,1; max_depth=5; subsample=1,0	0,9868

XGBoost memperoleh ROC-AUC validasi tertinggi sebesar 0,9868, diikuti Random Forest sebesar 0,9837 dan Logistic Regression sebesar 0,9326. Hasil pengujian pada 630 data uji disajikan pada Tabel 6.

Tabel 6. Perbandingan performa model pada data uji

Model	Accuracy	Precision	Recall	F1-Score	ROC-AUC
Logistic Regression	0,8492	0,5120	0,8586	0,6415	0,9268
Random Forest	0,9667	0,8750	0,9192	0,8966	0,9873
XGBoost	0,9698	0,8846	0,9293	0,9064	0,9912

XGBoost menghasilkan performa tertinggi pada seluruh metrik. Recall 0,9293 menunjukkan bahwa model dapat menemukan 92,93% pelanggan yang benar-benar churn. Logistic Regression mempunyai recall relatif tinggi tetapi precision hanya 0,5120 sehingga menghasilkan lebih banyak false positive. Kinerja model berbasis pohon menunjukkan bahwa pola churn melibatkan hubungan nonlinier dan interaksi antarvariabel. Temuan ini konsisten dengan Ardhani dan Tania (2025) serta Astarani dan Putri (2025), yang melaporkan performa kuat model ensemble dalam prediksi churn.



Gambar 2. Kurva ROC dan precision-recall pada data uji

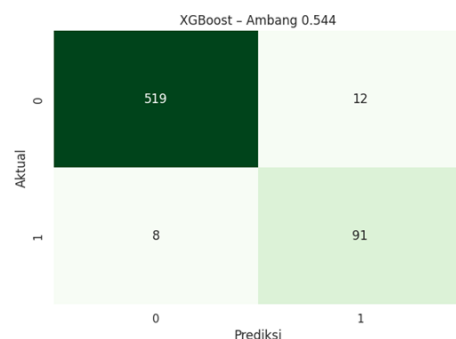
Evaluasi XGBoost pada Ambang Risiko

Ambang klasifikasi ditentukan dari probabilitas out-of-fold pada data latih. Nilai yang menghasilkan F1-score tertinggi adalah 0,5444 dengan F1-score out-of-fold 0,8781. Penggunaan ambang tersebut pada data uji menghasilkan accuracy 0,9683 dan ROC-AUC 0,9912.

Model mengklasifikasikan 519 pelanggan non-churn dan 91 pelanggan churn secara benar. Sebanyak 12 pelanggan non-churn diprediksi churn, sedangkan delapan pelanggan churn tidak terdeteksi. Selisih kecil antara ROC-AUC validasi 0,9868 dan ROC-AUC pengujian 0,9912 menunjukkan kemampuan generalisasi yang baik pada pembagian data penelitian.

Tabel 6. Perbandingan performa model pada data uji

Kelas	Precision	Recall	F1-score	Support
Tidak churn	0,9848	0,9774	0,9811	531
Churn	0,8835	0,9192	0,9010	99
Macro average	0,9342	0,9483	0,9410	630
Weighted average	0,9689	0,9683	0,9685	630



Gambar 3. Confusion matrix XGBoost pada ambang 0,5444

Tabel 7. Kepentingan fitur berdasarkan mean absolute SHAP

Peringkat	Fitur	Mean absolute SHAP
1	Frequency of Use	1,8041
2	Status	1,7170
3	Seconds of Use	1,1626
4	Call Failure	0,9870
5	Subscription Length	0,6562
6	Complains	0,6558
7	Frequency of SMS	0,5844
8	Age	0,4485
9	Customer Value	0,4133
10	Distinct Called Numbers	0,3983
11	Charge Amount	0,2983
12	Age Group	0,0765
13	Tariff Plan	0,0000



Gambar 4. Ringkasan pengaruh fitur terhadap keluaran XGBoost

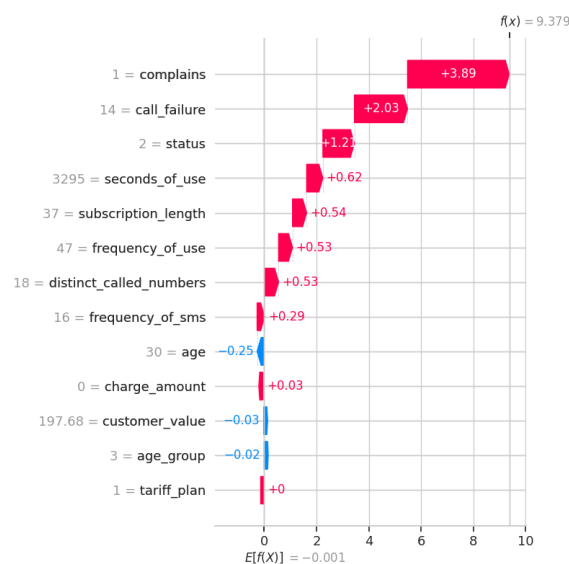
Frequency of Use menjadi fitur dengan kontribusi terbesar. Frekuensi penggunaan yang rendah cenderung mendorong prediksi ke arah churn, sedangkan penggunaan tinggi lebih banyak menurunkannya. Pola serupa terlihat pada Seconds of Use dan Frequency of SMS. Status tidak aktif, tingginya Call Failure, dan adanya Complains memberikan kontribusi positif kuat terhadap risiko churn. Customer Value berada pada peringkat kesembilan; nilai yang tinggi secara umum cenderung menurunkan risiko churn, tetapi tetap penting secara manajerial untuk menentukan prioritas retensi. Tariff Plan mempunyai nilai SHAP nol pada model ini.

Hasil tersebut menunjukkan bahwa penurunan keterlibatan dan masalah kualitas layanan dapat menjadi sinyal awal churn. Interpretasi SHAP menjembatani performa prediktif dan keputusan bisnis sebagaimana ditekankan oleh Asif et al. (2025) serta Ardhani dan Tania (2025).

Interpretasi Pelanggan Individual

Pelanggan yang dijelaskan pada Gambar 5 mempunyai probabilitas churn 0,9999. Faktor terbesar yang meningkatkan prediksi adalah adanya keluhan dengan nilai SHAP +3,89, diikuti 14 kegagalan panggilan sebesar +2,03 dan status tidak aktif sekitar +1,27. Usia 30 tahun dan Customer Value 197,68 sedikit menurunkan keluaran model, tetapi tidak mampu mengimbangi pengaruh keluhan dan gangguan layanan. Penjelasan individual

memungkinkan perusahaan memilih tindakan spesifik berupa penyelesaian keluhan, pemeriksaan kualitas jaringan, dan penawaran pemulihan layanan.



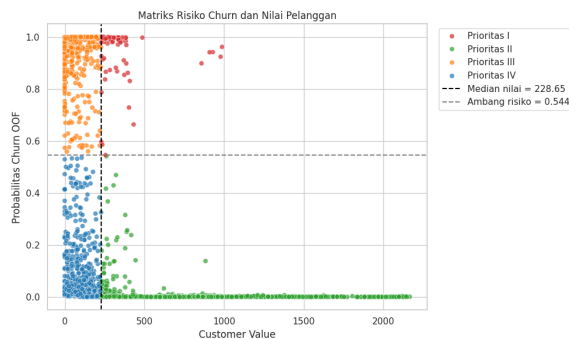
Gambar 5. SHAP pelanggan dengan probabilitas churn tertinggi

Segmentasi Prioritas Retensi

Segmentasi menggunakan ambang probabilitas churn 0,5444 dan median Customer Value data latih sebesar 228,6525. Probabilitas setiap pelanggan diperoleh secara out-of-fold agar pelanggan tidak dinilai oleh model yang dilatih menggunakan datanya sendiri.

Tabel 9. Hasil segmentasi prioritas retensi

Prioritas	Risiko	Nilai	Jumlah	Persentase	Strategi
I	Tinggi	Tinggi	74	2,35%	Intervensi segera dan penawaran personal
II	Rendah	Tinggi	1.499	47,59%	Program loyalitas dan pemeliharaan hubungan
III	Tinggi	Rendah	447	14,19%	Retensi berbiaya rendah atau otomatis
IV	Rendah	Rendah	1.130	35,87%	Pemantauan rutin



Gambar 6. Matriks risiko churn dan nilai pelanggan

Sebanyak 74 pelanggan atau 2,35% termasuk Prioritas I karena mempunyai risiko dan nilai tinggi. Kelompok ini perlu menerima intervensi segera melalui penyelesaian keluhan, peningkatan kualitas layanan, penawaran personal, dan insentif khusus. Prioritas II merupakan kelompok terbesar, yaitu 1.499 pelanggan atau 47,59%, yang perlu dipertahankan melalui program loyalitas. Sebanyak 447 pelanggan pada Prioritas III dapat ditangani dengan kanal retensi berbiaya rendah, sedangkan 1.130 pelanggan pada Prioritas IV cukup dipantau secara rutin.

Hasil segmentasi menunjukkan bahwa tidak seluruh pelanggan berisiko churn perlu ditangani dengan strategi yang sama. Integrasi probabilitas churn dan nilai pelanggan membantu membedakan pelanggan yang membutuhkan intervensi personal dari pelanggan yang cukup menerima retensi otomatis. Pendekatan ini mendukung pandangan Verbeke et al. (2012) bahwa pengelolaan churn perlu mempertimbangkan nilai ekonomi dan biaya intervensi, bukan hanya akurasi klasifikasi.

IV. KESIMPULAN

Penelitian berhasil mengembangkan model prediksi customer churn berbasis nilai pelanggan menggunakan explainable machine learning. XGBoost menghasilkan performa terbaik dengan ROC-AUC validasi 0,9868 dan ROC-AUC pengujian 0,9912. Pada ambang risiko 0,5444, model memperoleh accuracy 96,83%, precision 88,35%,

recall 91,92%, dan F1-score 90,10% untuk kelas churn.

SHAP menunjukkan bahwa Frequency of Use, Status, Seconds of Use, Call Failure, Subscription Length, dan Complains merupakan faktor paling berpengaruh. Frekuensi penggunaan yang rendah, status tidak aktif, tingginya kegagalan panggilan, dan adanya keluhan cenderung meningkatkan risiko churn. Penggabungan probabilitas churn dan Customer Value menghasilkan empat kelompok retensi, dengan 74 pelanggan atau 2,35% berada pada prioritas tertinggi.

Dengan demikian, pendekatan yang dikembangkan tidak hanya memprediksi pelanggan yang berpotensi churn secara akurat, tetapi juga menjelaskan faktor penyebab dan menerjemahkannya menjadi prioritas retensi. Penelitian selanjutnya disarankan menggunakan data beberapa perusahaan, validasi berbasis waktu, pengukuran Customer Lifetime Value dan biaya intervensi, serta eksperimen bisnis untuk menilai efektivitas strategi retensi.

V. REFERENSI

- Alfarez, R., Rianto, R., & Purwayoga, V. (2024). Penerapan Naïve Bayes untuk prediksi customer churn (Studi kasus: PT Hutchison 3 Indonesia). *Jurnal Riset dan Aplikasi Mahasiswa Informatika*, 5(2), 301–307. <https://doi.org/10.30998/jrami.v5i2.8556>
- Ardhani, D. A., & Tania, K. D. (2025). Knowledge discovery on e-commerce customer churn using interpretable machine learning: A comparative study of SHAP-based classifiers. *Journal of Applied Informatics and Computing*, 9(5), 2695–2702. <https://doi.org/10.30871/jaic.v9i5.10811>
- Arifin, M. (2015). IG-KNN untuk prediksi customer churn telekomunikasi. *Simetris: Jurnal Teknik Mesin, Elektro dan Ilmu Komputer*, 6(1), 1–10. <https://doi.org/10.24176/simet.v6i1.230>
- Asif, D., Arif, M. S., & Mukheimer, A. (2025). A data-driven approach with explainable artificial intelligence for customer churn prediction in the telecommunications industry. *Results in*

- Engineering*, 26, 104629.
<https://doi.org/10.1016/j.rineng.2025.104629>
- Astarani, I. G. A. R., & Putri, L. A. A. R. (2025). Klasifikasi customer churn menggunakan XGBoost dengan optimasi GridSearchCV berbasis Shapley Additive Explanations. *Jurnal Nasional Teknologi Informasi dan Aplikasinya*, 4(1), 1–10.
<https://doi.org/10.24843/JNATIA.2025.v04.i01.p01>
- Hossain, M. R. (2025). Predicting customer churn in telecommunications with machine learning models. *Asian Journal of Research in Computer Science*, 18(1), 53–66.
<https://doi.org/10.9734/ajrcos/2025/v18i1548>
- Husein, A. M., & Harahap, M. (2021). Pendekatan data science untuk menemukan churn pelanggan pada sektor perbankan dengan machine learning. *Data Sciences Indonesia*, 1(1), 8–13.
<https://doi.org/10.47709/dsi.v1i1.1169>
- Iskandar, M. A., & Latifa, U. (2023). Website prediksi customer churn untuk mempertahankan pelanggan pada perusahaan telekomunikasi. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 7(2), 1308–1316.
<https://doi.org/10.36040/jati.v7i2.6639>
- Mufida, E., Andriansyah, D., & Hertyana, H. (2025). Customer churn prediction pada sektor perbankan dengan model logistic regression dan random forest. *Computer Science (CO-SCIENCE)*, 5(1), 58–66.
<https://doi.org/10.31294/coscience.v5i1.7576>
- Suryana, N., Pratiwi, P., & Prasetyo, R. T. (2021). Penanganan ketidakseimbangan data pada prediksi customer churn menggunakan kombinasi SMOTE dan boosting. *IJCIT (Indonesian Journal on Computer and Information Technology)*, 6(1), 31–37.
<https://doi.org/10.31294/ijcit.v6i1.9545>
- UCI Machine Learning Repository. (2020). *Iranian churn dataset* [Data set]. University of California, Irvine.
<https://doi.org/10.24432/C5JW3Z>
- Verbeke, W., Dejaeger, K., Martens, D., Hur, J., & Baesens, B. (2012). New insights into churn prediction in the telecommunication sector: A profit-driven data mining approach. *European Journal of Operational Research*, 218(1), 211–229.
<https://doi.org/10.1016/j.ejor.2011.09.031>