

KOMPARASI ALGORITMA MACHINE LEARNING DALAM KLASIFIKASI LOYALITAS NASABAH BANK BERBASIS PARTICLE SWARM OPTIMIZATION

Khaerul Anam, Ade Rizki Rinaldi, Fathurrohman

Teknik Informatika, STMIK IKMI Cirebon
Jl. Perjuangan No 10 B Majasem Kota Cirebon
jodiust9@gmail.com

ABSTRAK

Konsumen atau nasabah dalam bisnis adalah aset berharga untuk keberlanjutan bisnis di masa depan. Memperoleh kepercayaan dan meningkatkan loyalitas nasabah memerlukan perhatian dan kepedulian yang tulus. Kemajuan teknologi telah membuat persaingan bisnis lebih ketat, meningkatkan potensi *churn* atau penghentian penggunaan produk oleh nasabah. Salah satu upaya untuk mengatasi penurunan loyalitas adalah dengan melakukan klasifikasi loyalitas nasabah bank menggunakan teknik *machine learning*. Penelitian ini bertujuan untuk membandingkan algoritma *Random Forest*, *Decision Tree*, dan *Artificial Neural Network* yang dioptimasi dengan metode *boosting Particle Swarm Optimization* (PSO) pada dataset nasabah bank XYZ yang diperoleh dari situs www.kaggle.com. Hasil penelitian menunjukkan bahwa algoritma *Random Forest* memiliki nilai akurasi tertinggi sebesar 86,34%, nilai AUC sebesar 0,854, rata-rata precision sebesar 83,55%, dan rata-rata recall sebesar 70,42%. Atribut yang paling signifikan dalam *Random Forest* adalah *Geography*, *Gender*, *Tenure*, dan *Estimated Salary*. Pada *Decision Tree*, atribut signifikan adalah *Tenure* dan *Estimated Salary*, sementara pada *Artificial Neural Network*, atribut signifikan adalah *Geography*, *Has Card*, dan *Estimated Salary*.

Kata kunci : Loyalitas Nasabah, Klasifikasi, Perbandingan Algoritma, Machine Learning, Particle Swarm Optimization

1. PENDAHULUAN

Konsumen atau nasabah dalam sebuah bisnis merupakan aset yang sangat berharga bagi keberlangsungan bisnis di masa mendatang. Nasabah perlu mendapatkan perhatian dan kepedulian yang sungguh-sungguh demi mendapatkan kepercayaan serta meningkatkan loyalitas nasabah [1]. Meningkatkan tingkat loyalitas nasabah sebuah perusahaan merupakan salah satu tantangan tersendiri dalam bidang CRM (*Customer Relationship Management*) di sebuah perusahaan [2].

Perkembangan teknologi berdampak pada persaingan usaha dalam dunia perbankan menjadi lebih ketat [3]. Persaingan ini berpotensi terhadap terjadinya *churn* atau berhentinya nasabah dalam menggunakan suatu produk dan menurunnya loyalitas nasabah. Beberapa faktor yang mengakibatkan ketidakloyalan nasabah diantaranya adalah kinerja CRM yang belum optimal, tidak terpenuhinya kebutuhan nasabah terhadap kemudahan kredit, suku bunga yang kompetitif dan kualitas layanan yang belum optimal [4].

Salah satu upaya dalam mengatasi penurunan loyalitas nasabah adalah dengan melakukan klasifikasi loyalitas nasabah bank menggunakan teknik *machine learning*. Klasifikasi loyalitas ini memberikan manfaat yang besar dalam sektor perbankan dalam membantu divisi kinerja *Customer Relationship Management* (CRM) dalam membuat program dalam rangka menghasilkan kualitas layanan yang lebih baik [5].

Penelitian mengenai klasifikasi loyalitas sebelumnya pernah dilakukan oleh [6] yang terbit pada tahun 2023. Penelitian ini menggunakan data *Credit Card Customers*

<https://www.kaggle.com/datasets/whenamancodes/credit-card-customers-prediction>.

Hasil penelitian menunjukkan perbedaan metode *feature selection* dengan *variance threshold* dan *correlation coefficient* tidak signifikan mempengaruhi performa ketiga algoritma klasifikasi dalam mengidentifikasi potensi hilangnya nasabah bank. Algoritma Logistic Regression (LR) sedikit lebih baik dengan data hasil preprocessing, sementara *Decision Tree* (DT) dan *Naïve Bayes* (NB) lebih unggul pada data tanpa preprocessing. *Naïve Bayes* juga menunjukkan stabilitas performa pada berbagai jenis data. Dari ketiga algoritma, *Decision Tree* dengan *feature selection correlation coefficient* menghasilkan F1 Score dan akurasi tertinggi, yaitu 0,96 dan 0,93.

Penelitian lain dilakukan oleh [7] terbit pada tahun 2020. Data yang digunakan dalam penelitian didapatkan langsung dari PT. Erdika Elit Sekuritas berupa data nasabah selama 3 tahun. Berdasarkan hasil perhitungan *entropy* dan *gain node akar* terhadap 5 atribut didapatkan hasil bahwa atribut latar belakang pendidikan memiliki pengaruh yang besar terhadap klasifikasi loyalitas nasabah bank dengan nilai *gain* sebesar 1.545292721.

Penelitian lain dilakukan oleh [8] pada tahun 2021. Penelitian ini membahas prediksi *churn* terhadap dataset pelanggan telekomunikasi yang berasal dari Kaggle.com yang memiliki 7.403 record data dan 21 atribut. Dengan distribusi dataset rasio 80% data training dan 20% data testing untuk model klasifikasi dengan algoritma *Random Forest*, *Adaboost* dan *Xgboost*. Dari hasil pengujian didapat nilai akurasi tertinggi untuk algoritma *adaboost* sebesar 80%.

Berdasarkan latar belakang yang dijelaskan sebelumnya, penelitian yang akan dilakukan bertujuan untuk melakukan perbandingan klasifikasi loyalitas nasabah bank dengan studi kasus dataset nasabah perbankan XYZ menggunakan teknik *machine learning* yang dioptimasi dengan metode *boosting Particle Swarm Optimization* (PSO). Dataset yang digunakan adalah dataset nasabah perbankan yang diperoleh dari website <https://www.kaggle.com/> untuk dilakukan perbandingan klasifikasi menggunakan algoritma *Random Forest*, *Decision Tree*, dan *Artificial Neural Network* yang dioptimasi dengan metode *boosting Particle Swarm Optimization* (PSO). PSO digunakan dalam rangka meningkatkan kinerja model prediksi dengan cara melakukan fitur seleksi sehingga hasil akurasi yang dihasilkan bisa lebih baik.

2. TINJAUAN PUSTAKA

2.1. Teknik Klasifikasi

Klasifikasi merupakan model *supervised learning* dalam disiplin ilmu *machine learning*. Dalam *supervised learning*, model dipresentasikan menggunakan dataset $D=\{x,y(_n=1\wedge N).\}$ dari input x dan label pasangan y . Nilai y disini dapat menjadi berbagai bentuk sesuai dengan *learning task* atau model pembelajaran yang digunakan. Dalam kasus klasifikasi, nilai y merupakan skala yang mewakili label daripada kelas. *Supervised learning* biasanya digunakan untuk menemukan model dengan parameter Θ yang paling baik dalam memprediksi data berdasarkan fungsi loss $L(y, \hat{y})$. Nilai $\hat{y}$ disini menunjukkan output dari model yang diperoleh dengan mengoperasikan data pada poin x pada fungsi $f(x, \Theta)$ sebagai representasi daripada model [9].

Dalam penelitian ini, kami menggunakan tiga algoritma klasifikasi: *Decision Tree*, *Random Forest*, dan *Neural Network*. *Decision Tree* adalah metode yang mempartisi data secara rekursif berdasarkan fitur tertentu untuk membentuk struktur pohon keputusan yang sederhana dan interpretatif. Algoritma ini bekerja dengan membagi dataset ke dalam subset yang lebih kecil berdasarkan fitur yang memberikan gain informasi tertinggi, hingga mencapai kondisi berhenti yang telah ditentukan. *Random Forest*, sebagai metode *ensemble learning*, terdiri dari kombinasi beberapa pohon keputusan untuk meningkatkan akurasi prediksi. Algoritma ini bekerja dengan membangun banyak pohon keputusan selama pelatihan dan output dari model adalah mode (untuk klasifikasi) atau rata-rata (untuk regresi) dari prediksi pohon-pohon tersebut. *Random Forest* memiliki keunggulan dalam mengurangi *overfitting* dan meningkatkan generalisasi model, terutama pada dataset dengan fitur yang tinggi [10].

Neural Network adalah model yang terdiri dari lapisan-lapisan neuron yang saling terhubung dan mampu menangkap hubungan kompleks dalam data. Model ini belajar dengan menyesuaikan bobot koneksi antar neuron berdasarkan error antara prediksi dan

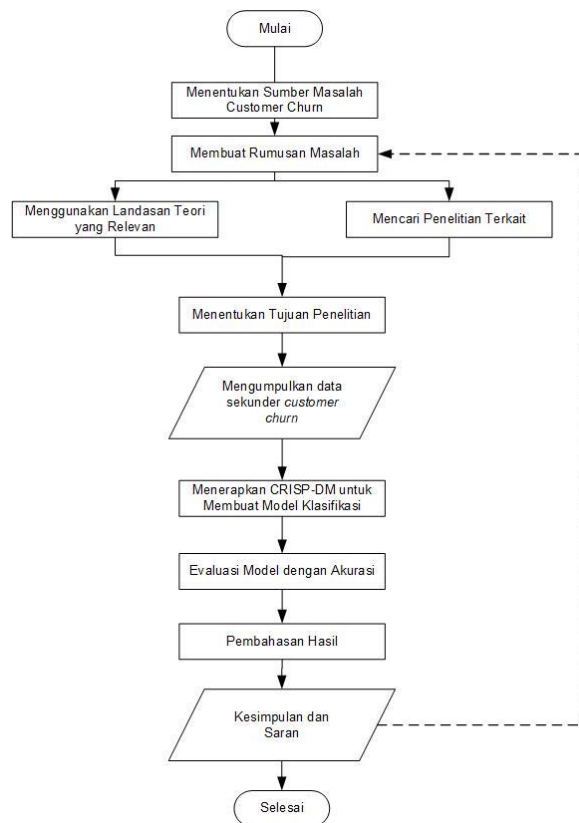
nilai sebenarnya melalui proses *backpropagation*. *Neural Network* sangat efektif dalam menangani data dengan struktur non-linear yang kompleks. Untuk mengevaluasi performa ketiga model klasifikasi tersebut, kami menggunakan beberapa metrik evaluasi seperti akurasi, *precision*, *recall*, dan *AUC* (*Area Under Curve*). Akurasi mengukur seberapa banyak prediksi yang benar dari total prediksi yang dilakukan, sementara *precision* dan *recall* memberikan informasi tentang ketepatan dan sensitivitas model dalam mendeteksi kelas target. *AUC* digunakan untuk mengevaluasi kinerja model dalam membedakan antara kelas positif dan negatif pada berbagai threshold [11]. Selanjutnya, kami juga melakukan validasi silang (*cross-validation*) untuk memastikan bahwa model yang dikembangkan tidak *overfitting* dan memiliki kemampuan generalisasi yang baik pada data yang belum pernah dilihat sebelumnya. Teknik ini membagi dataset menjadi beberapa subset, melatih model pada subset tertentu, dan mengujinya pada subset lainnya. Hasil dari validasi silang memberikan estimasi kinerja model yang lebih akurat [12].

2.2. Tahapan Penelitian

Penelitian ini disusun berdasarkan tahapan penelitian menggunakan *Cross-Industry Standard Process For Data Mining* (CRISP-DM) yang merupakan sebuah tahapan proses *data mining* dan *machine learning*. Tahapan penelitian yang digunakan dalam studi ini digambarkan melalui sebuah diagram alur yang sistematis, dimulai dari identifikasi masalah hingga penyusunan kesimpulan dan saran. Langkah pertama dalam tahapan ini adalah menentukan sumber masalah terkait customer churn, diikuti dengan penyusunan rumusan masalah yang komprehensif. Langkah ini didukung oleh penggunaan landasan teori yang relevan dan pencarian penelitian terkait untuk memperkuat kerangka konseptual penelitian. Selanjutnya, tujuan penelitian ditetapkan dengan jelas untuk memberikan arah yang terfokus dalam pengumpulan data sekunder terkait customer churn. Data yang terkumpul kemudian dianalisis menggunakan metodologi CRISP-DM (*Cross-Industry Standard Process for Data Mining*) untuk membuat model klasifikasi yang bertujuan memprediksi churn customer. Model yang dihasilkan dievaluasi berdasarkan akurasi untuk memastikan keandalannya.

Tahap akhir dari penelitian ini melibatkan pembahasan hasil yang diperoleh dari model klasifikasi, di mana temuan utama diinterpretasikan dalam konteks teori dan penelitian sebelumnya. Kesimpulan dan saran disusun sebagai refleksi dari hasil penelitian, menawarkan wawasan praktis dan akademik yang dapat berkontribusi pada pemahaman yang lebih baik mengenai customer churn serta menyarankan arah penelitian masa depan. Secara keseluruhan, tahapan penelitian ini dirancang untuk memberikan pendekatan yang holistik dan terstruktur dalam memahami dan memprediksi customer churn,

dengan menggunakan teknik data mining yang robust dan didukung oleh kerangka teori yang kuat. Langkah-langkah penelitian ini secara bertahap kerangka kerjanya dapat dijelaskan pada gambar 1.



Gambar 1. Tahapan Penelitian

2.3. Confusion Matrix

Confusion matrix merupakan alat ukur evaluasi yang terdapat pada weka classifier dengan tujuan untuk mempermudah dalam menganalisis performa algoritma karena *confusion matrix* memberikan informasi dalam bentuk angka sehingga dapat dihitung rasio keberhasilan klasifikasi [13]. Berikut ini pada tabel 1 disajikan bentuk *confusion matrix* :

Tabel 1. *Confusion Matrix*

		Kelas Hasil Prediksi Tidak (+)	Kelas Hasil Prediksi Ya (-)
	Tidak (+)	TP (True Positive)	FP (False Positive)
Kelas Asli	Ya (-)	FN (False Negative)	TN (True Negative)

Confusion matrix merupakan alat yang sangat penting dalam evaluasi kinerja model klasifikasi dalam machine learning. Matriks ini menyediakan pandangan yang mendetail mengenai prediksi yang dilakukan oleh model dibandingkan dengan nilai aktualnya. Confusion matrix terdiri dari empat komponen utama: True Positive (TP), False Positive (FP), True Negative (TN), dan False Negative (FN). True Positive adalah jumlah prediksi positif

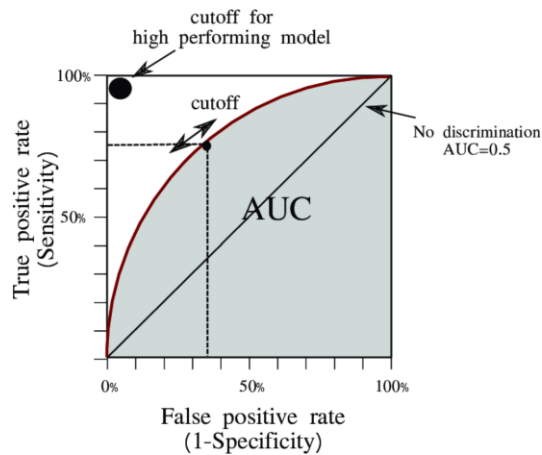
yang benar, sedangkan False Positive adalah jumlah prediksi positif yang salah. Sebaliknya, True Negative adalah jumlah prediksi negatif yang benar, dan False Negative adalah jumlah prediksi negatif yang salah.

Keempat komponen ini digunakan untuk menghitung berbagai metrik evaluasi yang penting dalam menilai kinerja model klasifikasi. Akurasi, yang didefinisikan sebagai proporsi prediksi yang benar dari total prediksi, dapat dihitung menggunakan rumus $(TP+TN)/(TP+FP+TN+FN)$. Precision, yang mengukur ketepatan prediksi positif, dihitung sebagai $TP/(TP+FP)$. Recall, yang mengukur sensitivitas model dalam mendeteksi kelas positif, dihitung sebagai $TP/(TP+FN)$. Selain itu, metrik F1-Score, yang merupakan harmonisasi antara precision dan recall, dihitung sebagai $2 \times (\text{precision} \times \text{recall}) / (\text{precision} + \text{recall})$.

Confusion matrix juga memungkinkan perhitungan metrik yang lebih lanjut seperti Specificity dan AUC (Area Under the Curve). Specificity mengukur kemampuan model dalam mendeteksi kelas negatif, dihitung sebagai $TN/(TN+FP)$. AUC, yang sering digunakan dalam evaluasi model klasifikasi, terutama pada skenario dengan data yang tidak seimbang, menggambarkan kemampuan model dalam membedakan antara kelas positif dan negatif pada berbagai threshold [14][15]. Analisis mendalam menggunakan confusion matrix dapat memberikan wawasan yang lebih baik tentang kinerja model pada tingkat granularitas yang lebih tinggi. Misalnya, model dengan akurasi tinggi mungkin masih memiliki precision atau recall yang rendah jika data yang digunakan tidak seimbang. Oleh karena itu, penting untuk menggunakan berbagai metrik evaluasi untuk mendapatkan gambaran yang komprehensif tentang kinerja model [16].

2.4. Area Under Curve

Cara lain yang dapat digunakan dalam mengukur kinerja dari sebuah model klasifikasi adalah menggunakan kurva ROC (Receiving Operating Characteristic) yang menggambarkan nilai Area Under Curve (AUC) [17]. Nilai AUC menggunakan sensitivitas atau spesifisitas sebagai dasar pengukuran. Nilai AUC berada diantara 0 dan 1. Apabila nilai AUC semakin mendekati 1, maka model klasifikasi yang terbentuk semakin akurat [18]. Ilustrasi kurva AUC seperti pada gambar 2.



ROC curves and area under curve (AUC).

Gambar 2. Kurva ROC dan Area Under Curve (AUC)

3. METODE PENELITIAN

Teknik analisis yang digunakan adalah *framework* dari *Cross-Industry Standard Process For Data Mining* (CRISP-DM). CRISP-DM telah lama diakui secara luas menjadi sebuah proses terstruktur atau kerangka kerja yang digunakan dalam memandu proyek *data mining* dan *machine learning* [19]. Kerangka kerja ini digunakan untuk menerjemahkan berbagai permasalahan bisnis kedalam sebuah model *data mining* dan menjalankan prosesnya secara independen baik pada sisi aplikatif maupun pada penggunaan teknologi [20]. Kerangka kerja ini terdiri dari enam fase utama: *business understanding*, *data understanding*, *data preparation*, *modeling*, *evaluation*, dan *deployment* [21] seperti pada gambar 3.



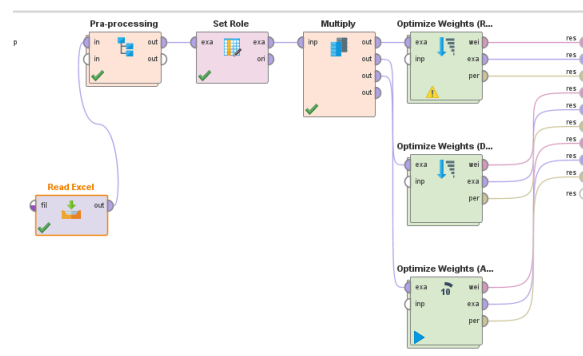
Gambar 3. Proses dalam CRISP-DM

Fase pertama, *business understanding*, melibatkan identifikasi tujuan bisnis dan pemahaman masalah yang ingin dipecahkan. Pada tahap ini, kami menentukan tujuan penelitian terkait customer churn serta mengidentifikasi metrik keberhasilan yang relevan. Selanjutnya, dalam fase *data understanding*, kami mengumpulkan dan mengeksplorasi data

sekunder yang relevan untuk mendapatkan wawasan awal mengenai data yang tersedia dan untuk mengidentifikasi masalah kualitas data yang perlu ditangani. Fase *data preparation* melibatkan pembersihan dan transformasi data, yang merupakan langkah krusial dalam memastikan kualitas data sebelum memasuki tahap pemodelan. Pada fase *modeling*, kami menerapkan berbagai algoritma klasifikasi dan memilih model terbaik berdasarkan kriteria akurasi dan nilai AUC. Tahap evaluasi kemudian dilakukan untuk memastikan model memenuhi tujuan bisnis dan memberikan hasil yang dapat diandalkan. Pada fase ini, kami juga membandingkan model yang dihasilkan dengan hasil penelitian sebelumnya untuk menilai kekuatan dan kelemahan model yang dikembangkan. Terakhir, fase *deployment* melibatkan penerapan model dalam lingkungan operasional dan pemantauan kinerja model secara berkelanjutan [22].

4. HASIL DAN PEMBAHASAN

Implementasi model klasifikasi *machine learning* ini dimulai dengan tahap membaca data menggunakan operator *Read excel*, melakukan preprocessing untuk cleansing data dan mengubah data *nominal* ke data *numerical* menggunakan operator *nominal to numerical*, menetapkan atribut target dengan operator *Set Role*, melakukan distribusi dataset menggunakan operator *multiply*, optimasi atribut menggunakan PSO, melakukan *K-Fold Cross Validation* untuk membagi dataset menjadi *data training* dan *data testing*. Adapun model proses dari penelitian ini menggunakan tools *rapidminer* seperti pada gambar 4.



Gambar 4. Model Klasifikasi Loyalitas Nasabah dengan Particle Swarm Optimization

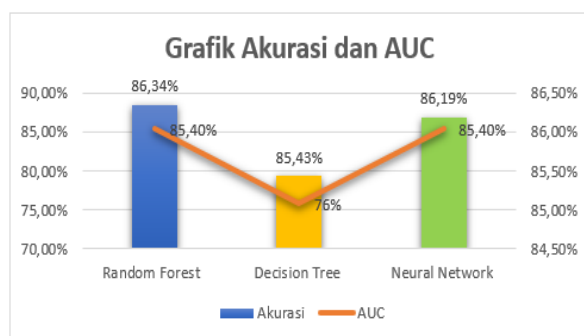
Hasil eksperimen terhadap tiga algoritma ini dengan menerapkan PSO secara lengkap ditampilkan dalam tabel 2. Dimana dari tabel tersebut terlihat bahwa algoritma yang memiliki nilai akurasi yang paling baik adalah algoritma *Random Forest* dengan nilai akurasi 86,30% dan nilai AUC 0,854.

Tabel 2. Hasil Akurasi dan AUC

No	Algoritma	Akurasi	AUC	Avg. Precision	Avg. Recall
1	Random Forest	86,30%	0,854	83,55%	70,42%
2	Decision Tree	85,43%	0,758	80,36%	70,13%
3	Neural Network	86,19%	0,854	81,86%	71,60%

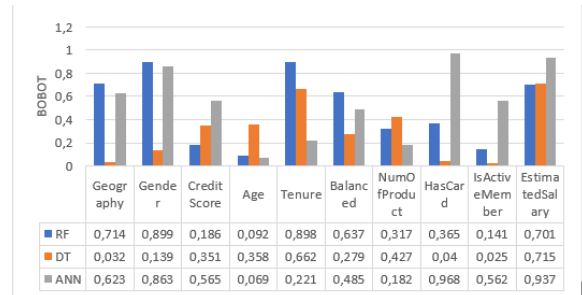
Nilai kesalahan atau error klasifikasi ditunjukkan pada banyaknya nilai *False Positive* (FP) dan *False Negative* (FN) atau kesalahan dalam klasifikasi data. Kesalahan klasifikasi dapat disebabkan karena beberapa alasan, salah satunya adalah karena jumlah data untuk setiap kelas yang tidak seimbang. Diketahui bahwa dataset dalam penelitian ini terdiri 20,36% data dengan label 1 dan 79,64% data dengan label 0, hal ini membuat kecenderungan data akan diklasifikasikan ke dalam kelas yang dominan atau kelas label 0. Selain itu, adanya outlier atau percikan dalam data juga dapat mempengaruhi performa klasifikasi.

Dari tabel pengujian dapat disimpulkan bahwa secara keseluruhan algoritma baik *Random Forest*, *Decision Tree* dan *Artificial Neural Network* memiliki akurasi yang cukup baik, namun yang memiliki akurasi paling baik adalah algoritma *Random Forest* dengan nilai akurasi 86,30%. Visualisasi nilai akurasi dan AUC dari percobaan ini dapat dilihat pada gambar 5.



Gambar 5. Grafik Perbandingan Nilai Akurasi dan AUC

Atribut yang memiliki bobot yang tinggi merupakan atribut signifikan yang mempunyai pengaruh besar terhadap hasil prediksi loyalitas nasabah bank. Eksperimen menggunakan algoritma yang berbeda menghasilkan bobot atribut yang berbeda dimana pada *Random Forest* atribut yang paling signifikan adalah *Geography*, *Gender*, *Tenure* dan *Estimated Salary*. Adapun pada analisis menggunakan algoritma *Decision Tree* atribut yang signifikan adalah *Tenure* dan *Estimated Salary*. Sedangkan pada algoritma *Neural Network* atribut yang signifikan adalah *Geography*, *Has Card* dan *Estimated Salary*. Secara lengkap bobot setiap atribut ini dapat dilihat pada gambar 6.



Gambar 6. Grafik bobot atribut

Eksperimen dalam penelitian ini dibandingkan dengan penelitian terkait yang dilakukan sebelumnya membahas penerapan teknik *machine learning* dalam menangani masalah *Customer Churn* dalam *Customer Relation Management*, dalam bidang bisnis telekomunikasi, supermarket dan perbankan. Sedangkan dalam penelitian ini, dilakukan penerapan algoritma *Random Forest*, *Decision Tree* dan *Artificial Neural Network* dalam melakukan klasifikasi loyalitas nasabah bank. Perbedaan dengan penelitian sebelumnya dengan penelitian yang dilakukan adalah pada cara pengujian analisis data serta penerapan *ensemble optimize parameters* dengan teknik *Feature selection Particle Swarm Optimization* (PSO).

Pada penelitian terkait sebelumnya teknik *boosting* yang banyak digunakan adalah teknik *bagging* karena fokus penelitian lebih kepada performa algoritma, sedangkan pada penelitian ini menggunakan PSO karena selain bertujuan untuk meningkatkan performa juga untuk mengetahui atribut-atribut signifikan yang perlu dianalisis. Analisis terhadap atribut dilakukan untuk menghasilkan *knowledge* yang dapat dimanfaatkan untuk menyelesaikan masalah loyalitas nasabah perbankan. Nasabah yang memiliki tingkat loyalitas yang rendah adalah nasabah dengan kriteria berasal dari *Germany* dengan kemungkinan *churn* sebesar 32,44% dan nasabah *Female* dengan kemungkinan *churn* sebesar 25,07%. Selain itu juga nasabah masa *Tenure* 1, 2 dan 10 tahun dengan kemungkinan *churn* diatas 22% dan nasabah dengan *Estimated Salary* 61-80k, 101-120k, 141-160k, 161-180k dan 181-200k yang memiliki peluang *churn* diatas 20, 37%. Pihak bank perlu menyusun strategi dan program yang difokuskan berdasarkan analisis loyalitas diatas sebagai langkah strategis berdasarkan analisis ilmiah yang didukung oleh referensi dan teori yang relevan.

Model *machine learning* klasifikasi loyalitas nasabah bank tidak terbatas hanya diaplikasikan untuk subjek penelitian ini saja, namun juga dapat digunakan untuk regional dan data yang lain. Dimana akan diperoleh hasil evaluasi yang berbeda sesuai dengan

kerakteristik data inputan. Atribut-atribut yang digunakan dalam penelitian ini adalah atribut yang secara global terdapat dalam dunia perbankan termasuk Indonesia. Sehingga hasil penelitian dapat pula diaplikasikan di Indonesia dengan menambahkan *data training* yang berasal dari Indonesia sehingga hasil yang didapatkan lebih relevan dengan keadaan di Indonesia.

Berdasarkan hasil penelitian, beberapa rekomendasi disusun untuk meningkatkan loyalitas nasabah bank. Untuk nasabah dari Jerman, yang memiliki tingkat churn sebesar 32.44%, bank sebaiknya meningkatkan branding melalui iklan yang menargetkan masyarakat Jerman dan meningkatkan layanan nasabah di cabang Jerman. Nasabah wanita, dengan tingkat *churn* sebesar 25.07%, dapat dipertahankan dengan meningkatkan branding melalui iklan yang menargetkan nasabah wanita dan bekerja sama dengan merchant belanja untuk memberikan promo khusus menggunakan kartu debit/kredit pada produk-produk wanita. Nasabah dengan masa kerja 1-2 tahun yang memiliki tingkat churn di atas 22%, disarankan untuk mendapatkan layanan maksimal sebagai nasabah baru dan program keuntungan bagi mereka yang menggunakan layanan lebih dari 2 tahun. Nasabah dengan masa kerja 10 tahun, juga dengan tingkat churn di atas 22%, perlu mendapatkan program tabungan masa tua dan bonus bagi mereka yang telah menjadi nasabah selama lebih dari 10 tahun. Selain itu, nasabah dengan perkiraan gaji 61-80k, 101-120k, 141-160k, 161-180k, dan 181-200k, yang memiliki tingkat *churn* di atas 20.37%, disarankan untuk mendapatkan suku bunga yang lebih menguntungkan, layanan khusus bagi nasabah prioritas, dan peningkatan layanan bagi pemegang kartu kredit prioritas/platinum.

5. KESIMPULAN DAN SARAN

Kesimpulan penelitian ini menegaskan efektivitas penggunaan algoritma Random Forest dalam memprediksi customer churn, khususnya pada data nasabah bank. Hasil eksperimen menunjukkan bahwa algoritma Random Forest memiliki performa terbaik dengan akurasi sebesar 86,34% dan nilai AUC (Area Under Curve) sebesar 0,854. Nilai AUC ini menunjukkan bahwa model mampu membedakan antara nasabah yang berpotensi churn dan yang tidak dengan tingkat keandalan yang tinggi. Selain itu, nilai rata-rata precision sebesar 83,55% menandakan bahwa model ini sangat tepat dalam prediksinya, sementara rata-rata recall sebesar 70,42% menunjukkan bahwa model ini cukup baik dalam menemukan kembali informasi yang relevan, meskipun ada ruang untuk peningkatan dalam mendeteksi semua nasabah yang berisiko churn. Lebih lanjut, penelitian ini juga mengungkapkan bahwa eksperimen dengan algoritma yang berbeda menghasilkan bobot atribut yang bervariasi, yang menandakan pentingnya pemilihan algoritma yang sesuai dengan karakteristik data yang digunakan. Model klasifikasi yang dihasilkan tidak

hanya relevan untuk data nasabah bank pada studi ini, tetapi juga memiliki potensi untuk diaplikasikan pada data dan regional lain dengan karakteristik berbeda. Penerapan ini, tentu saja, memerlukan evaluasi ulang terhadap performa model untuk memastikan bahwa model tetap akurat dan andal dalam konteks data yang berbeda. Secara keseluruhan, penelitian ini memberikan kontribusi signifikan dalam bidang analisis customer churn dengan menunjukkan bahwa algoritma Random Forest adalah alat yang kuat untuk prediksi churn. Implikasi praktis dari penelitian ini mencakup peningkatan strategi retensi nasabah, yang dapat membantu bank dan institusi keuangan lainnya dalam mengurangi tingkat churn dan meningkatkan loyalitas nasabah. Selain itu, hasil penelitian ini membuka peluang untuk penelitian lebih lanjut, termasuk eksplorasi algoritma lain dan penyesuaian model untuk berbagai jenis data dan konteks yang lebih luas. Dengan demikian, penelitian ini tidak hanya memperkaya literatur akademik tetapi juga memberikan panduan praktis bagi praktisi dalam industri keuangan.

DAFTAR PUSTAKA

- [1] M. Yosefina Muku Usu, R. P.C.Fanggiadae, M. Kurniawati, and M. Bunga, "PENGARUH CUSTOMER RELATIONSHIP MANAGEMENT (CRM) TERHADAP LOYALITAS PELANGGAN DENGAN KEPUASAN PELANGGAN SEBAGAI VARIABEL INTERVENING (STUDI PADA PT TELKOMSEL INDONESIA TBK CABANG MBAY) The Effect of Customer Relationship Management (CRM) on Customer," *Glory: JurnalEkonomi&IlmuSosia*, vol. 4, no. 2, pp. 319–333, 2023.
- [2] A. Wicaksono and T. N. Padilah, "Uji Performa Teknik Klasifikasi untuk Memprediksi Customer Churn Uji Performa Teknik Klasifikasi untuk Memprediksi Customer Churn," no. March, pp. 36–45, 2021, doi: 10.31294/bi.v9i1.9992.
- [3] Muslim, E. Rahmat Taufik, and Lutfi, "Mempertahankan Loyalitas Nasabah Melalui Kepuasan Dan Kepercayaan Nasabahe," *SAINS J. Manaj. dan Bisnis*, vol. 12, no. 2, 2020.
- [4] S. Syahnita, N. Nofriadi, and ..., "Rancang Bangun Customer Relationship Management (CRM) Dalam Penjualan Berbasis E-Commerce," *Build. Informatics ...*, vol. 4, no. 1, pp. 18–27, 2022, doi: 10.47065/bits.v4i1.1485.
- [5] A. F. Azmi and A. Voutama, "KOMPUTA : Jurnal Ilmiah Komputer dan Informatika PREDIKSI CHURN NASABAH BANK MENGGUNAKAN KLASIFIKASI RANDOM FOREST DAN DECISION TREE DENGAN EVALUASI CONFUSION MATRIX KOMPUTA : Jurnal Ilmiah Komputer dan Informatika," vol. 13, no. 1, 2024.
- [6] M. Farid Naufal, Subrata, A. Fernando Susanto, and C. Nathaneil Kansil, "Analisis Perbandingan

- Algoritma Machine Learning untuk Application of Machine Learning to Predict Potential Loss of Bank Customer,” vol. 22, no. 1, pp. 1–11, 2023.
- [7] K. Umam, D. Puspitasari, and A. Nurhadi, “Penerapan Algoritma C4.5 Untuk Prediksi Loyalitas Nasabah PT Erdika Elit Jakarta,” *J. Media Inform. Budidarma*, vol. 4, no. 1, p. 65, 2020, doi: 10.30865/mib.v4i1.1652.
- [8] I. M. Latief, A. Subekti, and W. Gata, “Prediksi Tingkat Pelanggan Churn Pada Perusahaan Telekomunikasi Dengan Algoritma Adaboost,” *J. Inform.*, vol. 21, no. 1, pp. 34–43, 2021, doi: 10.30873/ji.v21i1.2867.
- [9] F. Baharuddin and A. Tjahyanto, “Peningkatan Performa Klasifikasi Machine Learning Melalui Perbandingan Metode Machine Learning dan Peningkatan Dataset,” vol. 11, pp. 25–31, 2022.
- [10] C. Chen, A. Liaw, and L. Breiman, “Random forest: A comprehensive review,” *Journal of Machine Learning Research*, vol. 22, no. 1, pp. 15–21, 2021, doi: 10.1145/3122009.
- [11] T. T. Wong and P. Y. Yeh, “Reliable accuracy estimates from k-fold cross-validation,” *IEEE Transactions on Knowledge and Data Engineering*, vol. 32, no. 8, pp. 1586–1594, 2020, doi: 10.1109/TKDE.2019.2912814.
- [12] Z. Wu, S. Pan, F. Chen, G. Long, C. Zhang, and P. Yu, “A comprehensive survey on graph neural networks,” *IEEE Transactions on Neural Networks and Learning Systems*, vol. 32, no. 1, pp. 4–24, 2021, doi: 10.1109/tnnls.2020.2978386.
- [13] D. Normawati and S. Allit Prayogi, “Implementasi Naïve Bayes Classifier Dan Confusion Matrix Pada Analisis Sentimen Berbasis Teks Pada Twitter,” vol. 5, no. September, pp. 697–711, 2021.
- [14] T. Saito and M. Rehmsmeier, “The Precision-Recall Plot Is More Informative than the ROC Plot When Evaluating Binary Classifiers on Imbalanced Datasets,” *PLOS ONE*, vol. 10, no. 3, p. e0118432, 2015, doi: 10.1371/journal.pone.0118432.
- [15] D. Chicco and G. Jurman, “The advantages of the Matthews correlation coefficient (MCC) over F1 score and accuracy in binary classification evaluation,” *BMC Genomics*, vol. 21, no. 1, p. 6, 2020, doi: 10.1186/s12864-019-6413-7.
- [16] C. Ferri, J. Hernández-Orallo, and R. Modroiu, “An experimental comparison of performance measures for classification,” *Pattern Recognition Letters*, vol. 30, no. 1, pp. 27–38, 2021, doi: 10.1016/j.patrec.2008.08.010.
- [17] F. Yusa Rahman, I. Indah Purnomo, and N. Hijriana, “PENERAPAN ALGORITMA DATA MINING UNTUK KLASIFIKASI KUALITAS AIR,” vol. 13, no. 3, pp. 228–232, 2022.
- [18] M. M. Rodríguez-Hernández, R. E. Pruneda, and J. M. Rodríguez-Díaz, “Statistical analysis of the evolutive effects of language development in the resolution of mathematical problems in primary school education,” *Mathematics*, vol. 9, no. 10, 2021, doi: 10.3390/math9101081.
- [19] C. Schröer, F. Kruse, and J. M. Gómez, “A systematic literature review on applying CRISP-DM process model,” *Procedia Comput. Sci.*, vol. 181, no. 2019, pp. 526–534, 2021, doi: 10.1016/j.procs.2021.01.199.
- [20] T. Wuriyanto, H. B. Setiawan, and A. B. Tjandrarini, “Penerapan Model CRISP-DM pada Prediksi Nasabah Kredit yang Berisiko Menggunakan Algoritma Support Vector Machine,” vol. 10, 2022.
- [21] G. Gunawan, “DATA MINING USING CRISP-DM PROCESS FRAMEWORK ON OFFICIAL STATISTICS: A CASE STUDY OF EAST JAVA PROVINCE,” *Jurnal Ekonomi Dan Pembangunan*, vol. 29, no. 2, pp. 183–198, 2021, doi: 10.14203/.
- [22] “Data Clustering Menggunakan Metodologi CRISP-DM Untuk Pengenalan Pola Proporsi Pelaksanaan Tridharma,” *Jurnal Sistem Informasi Bisnis*, vol. 1, no. 3, pp. 129–134, 2014, doi: 10.21456/vol1iss3pp129-134.