

Prediksi Potensi Kinerja Calon Karyawan Customer Service Call Center Menggunakan Model Machine Learning Berbasis Data Rekrutmen

Andriyan Yoga Pratama*, Wildanil Ghozi

Fakultas Ilmu Komputer, Program Studi Teknik Informatika, Universitas Dian Nuswantoro, Semarang, Indonesia

Email: ¹*111202113943@mhs.dinus.ac.id, ²wildanil.ghozi@dsn.dinus.ac.id

Email Penulis Korespondensi: 111202113943@mhs.dinus.ac.id

Submitted: 05/05/2025; Accepted: 31/05/2025; Published: 01/06/2025

Abstrak—Proses seleksi karyawan merupakan tahap krusial bagi perusahaan dalam mendapatkan sumber daya manusia (SDM) yang berkualitas, terutama untuk posisi customer service call center yang menuntut keterampilan komunikasi dan ketahanan kerja yang tinggi. Metode seleksi berbasis data dapat meningkatkan akurasi dalam rekrutmen dibandingkan dengan pendekatan tradisional yang sering kali subjektif. Penelitian ini bertujuan untuk mengembangkan model prediksi potensi kinerja calon karyawan saat proses wawancara HR berdasarkan data latar belakang pendidikan, pengalaman kerja, dan faktor lainnya menggunakan algoritma machine learning. Data yang digunakan mencakup informasi demografi, jenjang pendidikan, pengalaman kerja, serta faktor-faktor lain yang dapat mempengaruhi kinerja calon karyawan di posisi customer service. Model yang diuji dalam penelitian ini mencakup Decision Tree, Random Forest, dan Neural Network. Hasil analisis menunjukkan bahwa faktor IPK, pengalaman kerja sebelumnya, serta keterlibatan dalam organisasi memiliki korelasi yang signifikan terhadap potensi kinerja calon karyawan. Analisis berbasis machine learning dalam rekrutmen dapat meningkatkan efektivitas seleksi dan efisiensi SDM. Dengan pendekatan ini, diharapkan perusahaan dapat meningkatkan efektivitas seleksi dan memilih kandidat terbaik dengan lebih akurat.

Kata Kunci: Machine Learning; Prediksi Kinerja; Rekrutmen; Customer Service; Call Center

Abstract—Employee selection process is a critical stage for companies in acquiring high-quality human resources (HR), particularly for customer service call center positions that demand excellent communication skills and strong work endurance. Data-driven recruitment methods have demonstrated improved accuracy compared to traditional, often subjective, approaches. This study aims to develop a predictive model to assess the potential performance of candidates during the HR interview stage, based on educational background, work experience, and other relevant factors, using machine learning algorithms. The dataset utilized includes demographic information, education levels, previous work experience, and other factors that may influence candidate performance in customer service roles. The models tested in this study include Decision Tree, Random Forest, and Artificial Neural Network algorithms. The analysis shows that GPA, prior work experience, and organizational involvement significantly correlate with the potential performance of candidates. The application of machine learning in the recruitment process can enhance selection effectiveness and improve HR efficiency. Through this approach, companies are expected to make more accurate hiring decisions and select the best candidates with greater precision.

Keywords: Machine Learning; Performance Prediction; Recruitment; Customer Service; Call Center

1. PENDAHULUAN

Proses rekrutmen karyawan merupakan salah satu aspek paling krusial dalam manajemen sumber daya manusia (SDM) di perusahaan. Memilih karyawan yang memiliki potensi dan sesuai dengan kebutuhan organisasi bukanlah tugas yang mudah. Dalam dunia kerja yang semakin kompetitif, perusahaan dituntut untuk menemukan cara yang lebih efektif dan efisien dalam merekrut karyawan agar dapat mempertahankan daya saingnya. Proses seleksi tradisional yang mengandalkan intuisi atau penilaian subjektif sering kali menghasilkan keputusan yang kurang akurat dan berdampak pada penurunan performa organisasi [1].

Saat ini, proses rekrutmen umumnya terdiri dari beberapa tahap seperti seleksi administrasi (*CV screening*), tes kemampuan (psikotes atau tes teknis), dan wawancara. Namun, tidak jarang perusahaan menemukan bahwa karyawan yang telah lolos seluruh tahapan tersebut justru menunjukkan performa kerja yang rendah setelah diterima. Masalah ini sangat relevan dalam industri *call center*, di mana peran *customer service* sangat penting untuk menjaga hubungan baik dengan pelanggan dan mempengaruhi citra perusahaan secara keseluruhan [2].

Posisi *customer service call center* menuntut keterampilan komunikasi, empati, kesabaran, serta kemampuan untuk bekerja di bawah tekanan. Oleh karena itu, dibutuhkan pendekatan seleksi yang mampu memprediksi potensi kinerja calon karyawan secara lebih objektif sejak awal proses rekrutmen. Salah satu pendekatan yang berkembang pesat dan menunjukkan hasil menjanjikan dalam hal ini adalah pemanfaatan teknologi *machine learning*. *Machine learning* memungkinkan analisis data historis untuk membangun model prediktif. Model ini dapat memanfaatkan variabel seperti latar belakang pendidikan, pengalaman kerja, usia, dan gender untuk mengestimasi performa kerja secara lebih objektif [3], [4].

Berbagai penelitian sebelumnya telah membuktikan efektivitas *machine learning* dalam konteks prediksi kinerja SDM. Sarkar dan Ghosh (2023) memanfaatkan algoritma Random Forest untuk memprediksi risiko *attrition* karyawan dengan akurasi tinggi dalam kondisi *dataset* yang tidak seimbang [3]. Lee dan Park (2023) menunjukkan bahwa pendekatan *ensemble learning* dapat meningkatkan akurasi dalam mengevaluasi kecocokan kandidat terhadap posisi tertentu [4]. Chowdhury et al. (2024) mengembangkan pendekatan evaluasi kinerja berbasis *machine learning* yang bebas bias, serta mengidentifikasi bahwa pengalaman kerja dan keterlibatan organisasi merupakan indikator kuat

dalam prediksi performa [5]. Selain itu, Kim dan Lee (2023) memprediksi niat *turnover* karyawan baru dan menemukan bahwa pemodelan yang akurat mampu mencegah kesalahan seleksi sejak tahap awal [6].

Dalam berbagai studi tersebut, pendekatan berbasis data terbukti dapat mengurangi bias subjektif dalam proses seleksi dan meningkatkan akurasi prediksi performa kerja. *Machine learning* memberikan keunggulan dalam mengidentifikasi pola-pola kompleks yang tidak mudah ditangkap oleh analisis konvensional. Misalnya, algoritma seperti *Decision Tree*, *Random Forest*, dan *Neural Network* telah digunakan secara luas dalam penelitian untuk memprediksi kinerja dan risiko *turnover* karyawan [7].

Namun, tantangan utama dalam implementasi *machine learning* pada proses seleksi karyawan adalah ketersediaan data yang berkualitas serta masalah ketidakseimbangan kelas dalam *dataset* (*imbalanced dataset*). Dalam konteks prediksi kinerja, sering kali data didominasi oleh kategori performa baik (kelas A atau B), sementara data dari performa rendah (kelas C atau D) jauh lebih sedikit. Hal ini dapat menyebabkan model menjadi bias terhadap kelas mayoritas dan gagal mengenali pola dari kelas minoritas. Oleh karena itu, diperlukan teknik penanganan data seperti *Random Over Sampling* (ROS) untuk menyeimbangkan distribusi kelas dalam *dataset* [8].

Penelitian ini bertujuan untuk mengembangkan model prediksi *Key Performance Indicator* (KPI) pelamar kerja berdasarkan data yang tercantum dalam CV, dengan fokus pada posisi *customer service* di perusahaan *call center*. *Dataset* yang digunakan merupakan data historis dari PT Infomedia Nusantara yang mencakup informasi pendidikan, pengalaman kerja, usia, gender, serta keterlibatan dalam organisasi. Model prediktif dikembangkan menggunakan tiga algoritma utama yaitu *Decision Tree*, *Random Forest*, dan *Neural Network*.

Model-model tersebut dievaluasi menggunakan metrik seperti akurasi, *precision*, *recall*, dan *F1-score*, serta divalidasi menggunakan teknik *k-Fold Cross Validation* untuk memastikan generalisasi model. Diharapkan, hasil dari penelitian ini dapat digunakan oleh perusahaan untuk mendukung proses pengambilan keputusan saat seleksi calon karyawan, sehingga meningkatkan efektivitas dan efisiensi dalam manajemen SDM.

Penelitian ini juga memperkuat pentingnya adopsi teknologi digital dalam manajemen SDM, sejalan dengan tren transformasi digital di berbagai sektor industri. Robles-Granda et al. (2020) menyatakan bahwa penggabungan *machine learning* dengan data psikometri dan afektif mampu meningkatkan prediksi performa karyawan secara signifikan [9]. Dalam jangka panjang, pendekatan ini dapat menurunkan tingkat *turnover* serta memperkuat budaya organisasi berbasis kinerja.

Dengan latar belakang tersebut, penelitian ini tidak hanya memiliki kontribusi akademik melalui pengembangan model prediktif berbasis data, tetapi juga nilai praktis yang tinggi dalam dunia industri. Penerapan *machine learning* dalam seleksi karyawan menjadi langkah strategis untuk memastikan bahwa individu yang direkrut memiliki potensi dan kompetensi yang sesuai dengan kebutuhan organisasi.

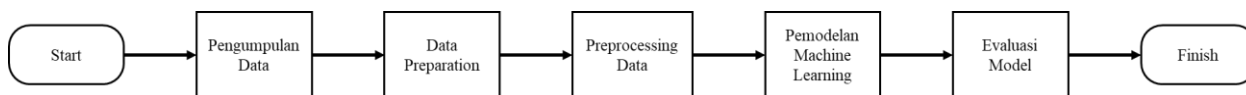
2. METODOLOGI PENELITIAN

Mengacu studi-studi yang telah menggunakan model *machine learning* serupa [1], [3], [10], [11], [12], metode penelitian ini terdiri dari beberapa tahapan utama, dimulai dari Pengumpulan Data primer mengenai calon karyawan dari PT Infomedia Nusantara yang mencakup informasi latar belakang pendidikan, pengalaman kerja, dan faktor relevan lainnya. Tahap Data Preparation meliputi pemformatan data, imputasi nilai hilang, label encoding, dan normalisasi data untuk memastikan konsistensi dan kualitas data. Selanjutnya, penerapan Preprocessing dengan Random Over Sampling (ROS) untuk menyeimbangkan dataset yang tidak seimbang. Pada tahap Modeling Data, digunakan tiga algoritma utama, yaitu Decision Tree (DT), Random Forest (RF), dan Artificial Neural Network (ANN) untuk membangun model prediksi kinerja calon karyawan. Terakhir, tahap Evaluasi diterapkan menggunakan metrik Accuracy, Precision, Recall, F1-Score, dan AUC untuk menentukan model terbaik dalam memprediksi potensi kinerja pada posisi Customer Service Call Center.

Untuk memberikan gambaran yang lebih komprehensif mengenai proses penelitian yang dilakukan, berikut ini disajikan alur metodologi penelitian secara visual. Gambar 1 menjelaskan lima tahapan utama yang menjadi kerangka kerja dalam pengembangan model prediktif, dimulai dari proses Pengumpulan Data, Data Preparation, Preprocessing Data, Pemodelan Machine Learning, hingga Evaluasi Model. Setiap tahapan saling terintegrasi dan dirancang untuk memastikan bahwa model yang dihasilkan tidak hanya akurat, tetapi juga robust terhadap ketidakseimbangan data dan dapat digunakan secara praktis dalam proses rekrutmen aktual.

Metodologi ini juga memperhatikan prinsip validitas internal dan eksternal guna memastikan bahwa hasil penelitian dapat diterapkan secara luas dalam konteks industri serupa. Setiap tahap dirancang untuk meminimalkan potensi bias, baik dari sisi distribusi data maupun pemilihan algoritma. Diharapkan, rancangan ini tidak hanya menghasilkan model yang akurat tetapi juga adaptif terhadap perubahan variabel input di masa depan. Penelitian ini juga dapat dijadikan acuan untuk pengembangan sistem pendukung keputusan berbasis data yang lebih canggih di sektor sumber daya manusia, khususnya dalam proses seleksi dan penempatan tenaga kerja yang efisien dan objektif.

Secara keseluruhan, tahapan metodologi penelitian ini disusun secara sistematis untuk memastikan bahwa setiap proses mulai dari pengumpulan dan persiapan data, penyeimbangan dataset, pemodelan *machine learning*, hingga evaluasi kinerja model berjalan secara terukur dan dapat direproduksi. Pendekatan ini bertujuan untuk mengidentifikasi model prediktif yang paling optimal dalam menilai potensi kinerja calon karyawan berdasarkan data rekrutmen aktual. Visualisasi alur lengkap tahapan tersebut dapat dilihat pada Gambar 1.

**Gambar 1.** Tahapan Metode Penelitian

2.1 Pengumpulan Data

Dataset ini mencerminkan kondisi nyata dalam rekrutmen calon karyawan untuk posisi customer service call center. Selain itu, penelitian ini mengacu pada studi akademik tentang seleksi SDM dan metode prediksi berbasis machine learning. Penelitian ini menggunakan dataset dalam format file Excel yang diperoleh dari PT Infomedia Nusantara. Dataset ini mencakup informasi penting terkait calon karyawan customer service call center, termasuk menggunakan dataset yang terdiri dari berbagai variabel yang berkaitan dengan latar belakang kandidat, seperti:

- Data demografi: usia dan gender
- Latar belakang pendidikan: jenjang pendidikan, jurusan, universitas, IPK
- Pengalaman kerja sebelumnya: “Ya”/”Tidak” memiliki pengalaman kerja sebelumnya
- Keterlibatan/pengalaman dalam organisasi: “Aktif”/”Tidak Aktif” dalam organisasi, kegiatan sosial, atau komunitas
- Faktor relevan lainnya: penempatan lokasi kerja, jenis universitas, bidang keahlian, kehadiran, hasil kinerja

Informasi-informasi tersebut kemudian direpresentasikan secara terstruktur dalam bentuk metadata dataset yang digunakan dalam penelitian ini. Detail dataset yang diperoleh dari PT Infomedia Nusantara ditunjukkan pada Tabel 1.

Tabel 1. Detail Dataset

No.	Keterangan	Detail
1.	Nama Data	Data Kinerja Karyawan PT Infomedia Nusantara
2.	Tahun Pembuatan	2024
3.	Jumlah Fitur	18 (12 kategorikal, 6 numerik)
4.	Jumlah Instance	476
5.	Jumlah Kelas	4
6.	Nama Kelas	['A', 'B', 'C', 'D']

2.2 Data Preparation

Pada tahap data preparation, diterapkan serangkaian proses untuk memastikan dataset siap digunakan dalam pemodelan machine learning. Beberapa langkah yang diterapkan antara lain sebagai berikut.

2.2.1 Pemformatan Data

Dataset yang diperoleh dari PT Infomedia Nusantara diformat ulang agar konsisten dan sesuai dengan struktur input model machine learning. Pemformatan ini meliputi penghapusan kolom yang tidak relevan, penataan ulang kolom, dan standarisasi format nilai.

2.2.2 Imputasi Data

Penanganan data yang hilang sangat penting untuk memastikan bahwa tidak ada informasi yang bias atau hilang dalam proses pemodelan machine learning. Mengatasi data yang hilang (missing values) dengan metode statistik yang sesuai. Misalnya, untuk data numerik digunakan mean atau median, sedangkan untuk data kategorikal digunakan modus atau teknik imputasi berbasis distribusi [13]. Dataset yang digunakan dalam penelitian ini tidak memiliki data yang hilang (missing values), sehingga proses imputasi data tidak diperlukan. Namun, dalam kondisi di mana terdapat missing values, langkah-langkah imputasi data akan diterapkan untuk menjaga kualitas data. Untuk data numerik, metode yang digunakan biasanya adalah mean, median, atau imputasi berbasis regresi, sedangkan untuk data kategorikal dapat digunakan modus atau teknik probabilitas [14].

2.2.3 Feature Encoding

Variabel kategorikal dikonversi menjadi numerik menggunakan metode Label Encoding, di mana setiap nilai kategori dibagi dengan jumlah total kategori (divide by number of values). Langkah ini penting agar data kategorikal dapat dipahami oleh algoritma machine learning [15].

2.2.4 Normalisasi Data

Dalam dataset yang digunakan, masing-masing fitur numerik (seperti IPK, usia, dan lama pengalaman kerja) memiliki rentang nilai yang berbeda-beda. Sebagai contoh, nilai IPK berada dalam rentang 0–4, sedangkan usia bisa berkisar antara 20–40 tahun. Rentang nilai yang berbeda ini dapat menyebabkan dominasi fitur tertentu dalam proses pembelajaran model machine learning, terutama pada algoritma yang sensitif terhadap skala data seperti Artificial Neural Network dan algoritma berbasis jarak.

Jika tidak dinormalisasi, fitur dengan skala besar akan lebih berpengaruh dalam penentuan output model dibandingkan fitur dengan skala kecil. Hal ini bisa mengakibatkan bias dalam pembelajaran model, serta memperlambat konvergensi pada algoritma yang berbasis gradien seperti Neural Network. Untuk mengatasi permasalahan ini, penerapan proses normalisasi data agar seluruh fitur numerik berada pada skala yang seragam. Salah satu metode normalisasi yang digunakan dalam penelitian ini adalah standarisasi (standardization), yaitu teknik transformasi data numerik sehingga memiliki rata-rata (mean) sebesar 0 dan standar deviasi sebesar 1. Proses ini dikenal juga sebagai Z-score normalization. Standarisasi diterapkan dengan persamaan rumus 1 berikut.

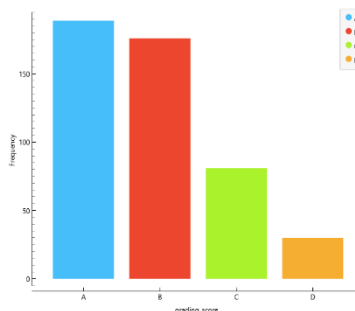
$$z = (x - \mu) / \sigma \quad (1)$$

Di mana x adalah nilai asli dari fitur, μ adalah rata-rata (mean) dari fitur tersebut, σ adalah standar deviasi dari fitur tersebut. Dengan menerapkan standarisasi, seluruh fitur numerik akan berada dalam skala yang setara, sehingga model machine learning dapat memproses data secara lebih adil dan efisien. Proses standarisasi (Standardization) dengan menyesuaikan nilai fitur numerik menjadi distribusi dengan mean 0 dan standar deviasi 1. Normalisasi ini memastikan bahwa semua fitur berada dalam skala yang seragam, sehingga model machine learning dapat memproses data dengan lebih efektif [16].

2.3 Model Preprocessing dan Machine Learning

2.3.1 Preprocessing

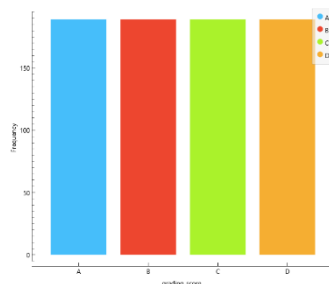
Pada tahap preprocessing, ditemukan bahwa dataset yang digunakan dalam penelitian ini tidak seimbang (imbalanced data). Hal ini terlihat dari distribusi grading score seperti yang ditunjukkan pada gambar 2 menunjukkan dominasi kelas A dan B, sedangkan kelas C dan D memiliki proporsi yang jauh lebih kecil. Data yang tidak seimbang dapat menyebabkan model machine learning cenderung bias terhadap kelas mayoritas [17], sehingga akurasi prediksi untuk kelas minoritas menjadi rendah. Untuk menggambarkan kondisi distribusi kelas secara visual, disajikan grafik frekuensi grading score pada data awal sebelum dilakukan teknik balancing. Visualisasi ini memperjelas dominasi kelas A dan B, serta memperlihatkan proporsi kelas C dan D yang jauh lebih kecil. Berikut adalah Gambar 2 grafik sebaran kelas yang masih belum seimbang.



Gambar 2. Distribusi Grading Score Sebelum ROS

Untuk mengatasi masalah ini, digunakan teknik Random Over Sampling (ROS). Random Over Sampling adalah metode pendekatan sederhana yang efektif untuk dataset HR dalam klasifikasi kinerja yang digunakan dalam menyeimbangkan distribusi kelas dalam dataset dengan cara menggandakan sampel pada kelas minoritas secara acak hingga jumlahnya setara dengan kelas mayoritas [8], [18]. Dengan teknik ini, model machine learning dapat belajar dari jumlah sampel yang seimbang, sehingga meningkatkan kemampuan dalam mengklasifikasikan semua kelas secara adil.

Setelah diterapkan Random Over Sampling, distribusi grading score menjadi seimbang seperti yang ditunjukkan pada gambar 3 dengan proporsi yang setara untuk setiap kelas (A, B, C, D). Hal ini diharapkan dapat meningkatkan performa model dalam memprediksi kinerja calon karyawan berdasarkan grading score mereka. Perubahan proporsi antar kelas setelah proses balancing menjadi lebih merata dan representatif, yang secara visual dapat divisualisasikan melalui grafik distribusi. Grafik ini memberikan bukti bahwa teknik ROS berhasil menyeimbangkan jumlah data pada masing-masing kategori grading score. Berikut adalah Gambar 3 grafik sebaran kelas yang sudah seimbang.



Gambar 3. Distribusi Grading Score Menggunakan ROS

2.3.2 Machine Learning

Setelah preprocessing, model prediksi dikembangkan menggunakan beberapa metode machine learning sebagai berikut.

2.3.2.1 Decision Tree

Decision Tree adalah algoritma pembelajaran terawasi yang digunakan untuk tugas klasifikasi dan regresi. Algoritma ini bekerja dengan membangun pohon keputusan berdasarkan fitur dalam dataset dan memprediksi nilai target berdasarkan aturan keputusan dari akar hingga daun [19].

Cara Kerja algoritma memecah dataset menjadi subset yang lebih kecil berdasarkan fitur yang memberikan Information Gain (IG) atau Gini Index terbaik. Proses ini berlanjut secara rekursif hingga mencapai node daun. Rumus Information Gain (IG) pada persamaan 2 dan 3 digunakan untuk memilih fitur terbaik dalam setiap pemisahan data.

$$IG(S, A) = Entropy(S) - \sum_{\{v \in Values(A)\}} \frac{|S_v|}{|S|} \cdot Entropy(S_v) \quad (2)$$

Di mana S = Set data, A = Fitur, S_v = Subset data untuk nilai fitur. Rumus Entropy digunakan untuk mengukur impurity dalam data

$$Entropy(S) = - \sum_{i=1}^n p_i \log_2 p_i \quad (3)$$

Di mana p_i = Proporsi kelas i dalam dataset S

2.3.2.2 Random Forest

Random Forest adalah algoritma ensemble yang menggabungkan beberapa Decision Tree untuk meningkatkan akurasi prediksi dan mengurangi overfitting. Algoritma ini bekerja dengan membangun banyak pohon keputusan secara acak dan menggabungkan prediksinya melalui teknik voting (untuk klasifikasi) atau rata-rata (untuk regresi) [20]. Metode ini telah terbukti efektif dalam meningkatkan akurasi prediksi dan efisiensi komputasi saat dikombinasikan dengan seleksi fitur berbasis Gini Index [21], [22].

Cara kerja menggunakan Bootstrap Sampling untuk membuat subset acak dari data pelatihan. Pada setiap node, hanya subset acak dari fitur yang dipertimbangkan untuk pemisahan. Membuat banyak Decision Tree dan menggabungkan hasilnya. Berikut rumus voting pada klasifikasi pada persamaan 4 dan 5.

$$\hat{y} = mode \{h_1(x), h_2(x), \dots, h_n(x)\} \quad (4)$$

Di mana $\hat{y}$ = Prediksi akhir, $h_i(x)$ = Prediksi dari pohon ke- i . Dengan rumus rata-rata pada Regresi

$$\hat{y} = \frac{1}{n} \sum_{i=1}^n h_i(x) \quad (5)$$

2.3.2.3 Artificial Neural Network

Artificial Neural Network (ANN) adalah algoritma machine learning berbasis jaringan saraf tiruan yang terdiri dari lapisan input, lapisan tersembunyi, dan lapisan output. ANN cocok untuk menangani data yang kompleks dengan pola non-linear [23].

Cara kerja data input melewati setiap neuron dalam lapisan tersembunyi, di mana bobot (w) dan bias (b) diterapkan. Fungsi aktivasi digunakan untuk memperkenalkan non-linearitas dalam model. Proses ini berlanjut hingga mencapai lapisan output, di mana hasil prediksi dihasilkan. Berikut rumus perhitungan Neuron pada persamaan 6.

$$z = \sum_{i=1}^n w_i x_i + b \quad (6)$$

$$a = f(z)$$

Di mana z = Input total neuron, w_i = Bobot untuk input x_i , b = Bias neuron, $f(z)$ = Fungsi aktivasi, seperti ReLU, Sigmoid, atau Softmax. Dengan rumus Backpropagation (Penyesuaian Bobot) pada persamaan 7.

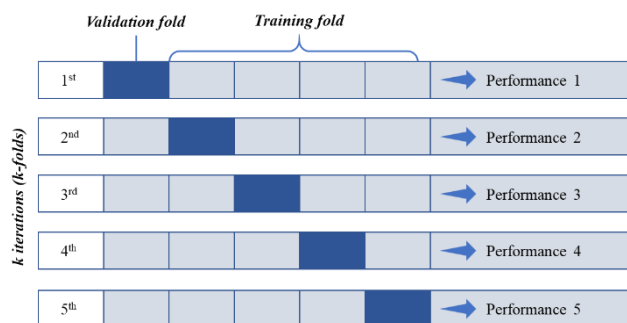
$$w = w - \eta \frac{\partial L}{\partial w} \quad (7)$$

Di mana w = Bobot, η = Learning rate, L = Fungsi kerugian (loss function). Ketiga algoritma ini memiliki karakteristik dan keunggulan masing-masing. Decision tree mudah diinterpretasikan, tetapi rentan terhadap overfitting. Random forest mengatasi overfitting dengan metode ensemble, menghasilkan model yang lebih stabil. Artificial neural network mampu menangkap pola kompleks dalam data, cocok untuk masalah non-linear, tetapi membutuhkan data dalam jumlah besar dan waktu pelatihan yang lebih lama.

2.4 Evaluasi Model

Validasi model diterapkan dengan menggunakan metode cross-validation untuk memastikan bahwa hasil prediksi memiliki tingkat generalisasi yang baik. Untuk memberikan pemahaman yang lebih jelas mengenai cara kerja proses validasi, berikut disajikan visualisasi skematik dari metode k-Fold Cross Validation yang digunakan dalam penelitian

ini. Gambar ini memperlihatkan bagaimana data dibagi secara bergiliran antara data pelatihan dan data validasi pada setiap iterasi. Ilustrasi dari metode cross validation yang diterapkan dapat dilihat pada Gambar 4.



Gambar 4. Ilustrasi Evaluasi Model

Metode k-Fold Cross Validation merupakan teknik evaluasi model yang membagi data menjadi k bagian (fold) yang sama besar, di mana proses pelatihan dan pengujian diterapkan sebanyak k kali. Pada setiap iterasi, satu fold digunakan sebagai data validasi, sementara k-1 fold lainnya digunakan untuk melatih model. Proses ini diterapkan secara bergantian sehingga setiap fold menjadi data validasi satu kali. Hasil evaluasi dari setiap iterasi dicatat dan dirata-rata untuk memperoleh estimasi performa model yang lebih stabil dan tidak bias terhadap subset data tertentu. Pendekatan ini efektif dalam mengurangi risiko overfitting dan memberikan gambaran yang lebih umum terhadap kemampuan generalisasi model pada data baru [24].

Confusion matrix merupakan tabel yang menggambarkan performa klasifikasi model dengan membandingkan hasil prediksi dengan label aktual. Dalam konteks klasifikasi multi-kelas seperti pada penelitian ini (kelas A, B, C, dan D), confusion matrix menampilkan jumlah prediksi benar dan salah untuk masing-masing kelas. Untuk kasus klasifikasi biner (untuk pemahaman awal), confusion matrix terdiri dari empat komponen utama.

- True Positive (TP): Jumlah data yang sebenarnya positif dan berhasil diprediksi sebagai positif oleh model.
- True Negative (TN): Jumlah data yang sebenarnya negatif dan berhasil diprediksi sebagai negatif oleh model.
- False Positive (FP): Jumlah data yang sebenarnya negatif, tetapi salah diprediksi sebagai positif (Type I Error).
- False Negative (FN): Jumlah data yang sebenarnya positif, tetapi salah diprediksi sebagai negatif (Type II Error).

Untuk memahami lebih dalam bagaimana hasil prediksi dibandingkan dengan label aktual dalam konteks evaluasi klasifikasi, perlu disajikan struktur dasar confusion matrix. Tabel 2 berikut menyajikan empat komponen utama yang digunakan untuk menghitung metrik evaluasi seperti akurasi, presisi, recall, dan F1-score.

Gambaran umum struktur confusion matrix dapat dilihat dari Tabel 2.

Tabel 2. Struktur Confusion Matrix

	Predicted Positive	Predicted Negative
Actual Positive	True Positive (TP)	False Negative (FN)
Actual Negative	False Positive (FP)	True Negative (TN)

Untuk klasifikasi multi-kelas, confusion matrix akan berbentuk matriks persegi dengan ukuran sesuai jumlah kelas. Setiap baris mewakili label aktual, sedangkan kolom mewakili label prediksi. Nilai diagonal menunjukkan jumlah prediksi yang benar untuk tiap kelas, sedangkan nilai di luar diagonal menunjukkan kesalahan prediksi antar kelas. Nilai-nilai dari confusion matrix ini akan digunakan untuk menghitung metrik evaluasi model seperti accuracy, precision, recall, dan F1-score. Evaluasi pada setiap model diterapkan dengan menggunakan metrik berikut.

2.4.1 Accuracy

Mengukur seberapa sering model membuat prediksi yang benar dari keseluruhan prediksi yang diterapkan. Metrik ini memberikan gambaran umum tentang performa model, tetapi kurang efektif jika data tidak seimbang (imbalanced dataset), karena model dapat menghasilkan akurasi tinggi dengan hanya memprediksi kelas mayoritas. Rumus Accuracy tercantum pada persamaan 8.

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (8)$$

Accuracy digunakan untuk mengukur proporsi prediksi yang benar dari keseluruhan prediksi. Dalam konteks penelitian ini, metrik ini memberikan gambaran umum seberapa sering model dapat mengklasifikasikan potensi kinerja kandidat secara tepat berdasarkan latar belakang mereka. Namun, karena dataset awal tidak seimbang (kelas C dan D sedikit), nilai akurasi yang tinggi belum tentu menunjukkan model yang baik, sehingga perlu didampingi metrik lain seperti recall dan F1-score untuk evaluasi menyeluruh [5], [17]. Kelebihan mudah dipahami dan memberikan gambaran keseluruhan kinerja model. Kekurangan tidak ideal jika terdapat ketidakseimbangan kelas dalam dataset.

2.4.2 Precision

Mengukur ketepatan model dalam menentukan kandidat yang berpotensi tinggi. Metrik ini berguna dalam situasi di mana kesalahan False Positive (FP) harus dikurangi, misalnya dalam pemilihan kandidat yang benar-benar berkualitas. Rumus Precision tercantum pada persamaan 9.

$$Precision = \frac{TP}{TP+FP} \quad (9)$$

Precision berfokus pada seberapa tepat model dalam mengidentifikasi kandidat dengan kinerja tinggi (kelas A dan B). Dalam konteks rekrutmen, precision penting untuk meminimalkan kesalahan dalam merekomendasikan kandidat yang sebenarnya tidak memenuhi kriteria. Precision yang tinggi membantu memastikan bahwa kandidat yang dipilih memang benar-benar memiliki potensi kinerja baik [6], [7]. Kelebihan berguna jika kesalahan dalam memprediksi kandidat berpotensi tinggi harus dikurangi (misalnya dalam seleksi kandidat terbaik). Kekurangan tidak mempertimbangkan False Negative (FN), yang juga bisa berdampak signifikan.

2.4.3 Recall

Mengukur sejauh mana model dapat menangkap semua kandidat yang benar-benar berpotensi tinggi. Metrik ini berguna ketika lebih penting untuk tidak melewatkan kandidat berkualitas (mengurangi False Negative (FN)). Rumus Recall tercantum pada persamaan 10.

$$Recall = \frac{TP}{TP+FN} \quad (10)$$

Recall digunakan untuk menilai kemampuan model dalam menangkap semua kandidat potensial berkinerja tinggi. Dalam studi ini, recall menjadi penting karena perusahaan tidak ingin melewatkan kandidat potensial hanya karena model gagal mengklasifikasikannya dengan benar. Terutama dalam kondisi real-world, kelolosan kandidat potensial dapat berdampak pada efisiensi operasional dan pelayanan pelanggan [1], [3]. Kelebihan cocok digunakan dalam situasi di mana penting untuk menangkap semua kandidat berkualitas tinggi, seperti dalam pemilihan kandidat untuk posisi kritis. Kekurangan jika hanya mempertimbangkan recall tanpa precision, model mungkin akan menghasilkan banyak False Positive.

2.4.4 F1-Score

Metrik gabungan yang menghitung keseimbangan antara Precision dan Recall. Metrik ini sangat berguna jika dataset tidak seimbang, karena memberikan gambaran seberapa baik model mengklasifikasikan kandidat tanpa bias terhadap kelas mayoritas. Rumus F1-Score tercantum pada persamaan 11.

$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (11)$$

F1-score menggabungkan precision dan recall menjadi satu metrik yang seimbang. Metrik ini sangat berguna pada kasus dengan data tidak seimbang seperti pada penelitian ini. F1-score memperlihatkan kinerja model secara menyeluruh tanpa terlalu bias pada kelas mayoritas. Oleh karena itu, F1-score menjadi indikator utama dalam memilih model terbaik dalam penelitian ini [11], [16]. Kelebihan berguna ketika terdapat ketidakseimbangan kelas dalam data. Kekurangan sulit untuk menginterpretasikan secara langsung dibandingkan dengan Accuracy.

3. HASIL DAN PEMBAHASAN

Hasil penelitian ini menunjukkan bahwa optimasi preprocessing dan pemilihan algoritma yang tepat berpengaruh signifikan terhadap akurasi prediksi kinerja. Dengan melakukan normalisasi dan konversi variabel diskret menjadi numerik, model dapat lebih optimal dalam melakukan prediksi, konsisten dengan beberapa temuan yang sudah diterapkan dan menjadi sumber referensi [5], [25].

Masalah utama yang dihadapi dalam penelitian ini adalah ketidakseimbangan distribusi kelas pada data kinerja (grading score), di mana kelas C dan D memiliki jumlah instance yang jauh lebih sedikit dibandingkan kelas A dan B. Ketidakseimbangan ini menyebabkan bias model terhadap kelas mayoritas dan kegagalan dalam mengenali pola dari kelas minoritas. Untuk menyelesaikan masalah tersebut, dilakukan teknik balancing data menggunakan metode Random Over Sampling (ROS) pada tahap preprocessing. ROS secara acak menggandakan instance pada kelas minoritas hingga jumlahnya sebanding dengan kelas mayoritas, sehingga menghasilkan distribusi kelas yang setara.

Setelah tahapan preprocessing, termasuk normalisasi dan balancing dengan ROS, dilakukan pemodelan menggunakan tiga algoritma utama yaitu Decision Tree, Random Forest, dan Neural Network. Masing-masing model digunakan untuk mempelajari pola dari dataset yang telah diseimbangkan. Decision Tree dipilih karena kemampuannya menginterpretasi struktur data dengan baik, sedangkan Random Forest digunakan untuk mengatasi overfitting dan memperkuat akurasi dengan pendekatan ensemble learning. Neural Network digunakan karena kemampuannya dalam mempelajari relasi non-linear, meskipun algoritma ini memerlukan tuning parameter dan data yang besar.

Evaluasi dilakukan dengan menerapkan k-Fold Cross Validation sebanyak 10 fold untuk memastikan hasil generalisasi yang baik dan mengurangi risiko overfitting. Pada setiap fold, dilakukan pelatihan dan pengujian model, lalu hasil dari metrik Accuracy, Precision, Recall, dan F1-score dirata-rata. Confusion matrix digunakan untuk melihat performa klasifikasi per kelas dan menjadi dasar analisis terhadap kekuatan serta kelemahan masing-masing model.

Untuk memberikan gambaran yang lebih komprehensif terhadap performa tiap algoritma, maka hasil evaluasi disajikan dalam dua tahap, yaitu sebelum dan sesudah penerapan teknik balancing data menggunakan Random Over Sampling (ROS). Pembagian ini bertujuan untuk mengevaluasi sejauh mana dampak ROS terhadap peningkatan akurasi dan pemerataan klasifikasi antar kelas.

3.1 Analisis Hasil Model Sebelum Penerapan Random Over Sampling (ROS)

Sebelum menerapkan teknik Random Over Sampling (ROS), model diuji menggunakan data asli tanpa penyesuaian distribusi kelas. Hasil eksperimen awal ini menunjukkan bahwa model mengalami kesulitan dalam mengenali kelas dengan jumlah instance yang lebih sedikit (kelas minoritas). Untuk menggambarkan performa awal dari masing-masing algoritma sebelum dilakukan penyeimbangan data, disajikan visualisasi confusion matrix yang menunjukkan distribusi prediksi terhadap kelas aktual. Grafik ini membantu mengidentifikasi pola kesalahan klasifikasi, terutama pada kelas minoritas yang cenderung tidak dikenali dengan baik oleh model.

Visualisasi dalam bentuk confusion matrix yang dapat memetakan kemampuan klasifikasi model terhadap seluruh kelas yang ada. Confusion matrix ini menjadi instrumen penting dalam evaluasi model karena mampu menunjukkan secara detail ketepatan dan kesalahan klasifikasi yang terjadi, baik pada kelas mayoritas maupun minoritas. Melalui pendekatan ini, penelitian dapat mengidentifikasi titik lemah dari masing-masing algoritma secara lebih terstruktur, sekaligus menjadi dasar pertimbangan dalam menentukan intervensi lanjutan seperti balancing data. Dengan demikian, evaluasi berbasis confusion matrix tidak hanya bersifat deskriptif, namun juga strategis dalam upaya perbaikan model prediksi ke depan. Detail dari confusion matrix setiap model dapat terlihat dari Gambar 5, Gambar 6, dan Gambar 7.

		Predicted				
		A	B	C	D	Σ
Actual	A	108	59	18	4	189
	B	73	65	30	8	176
	C	25	32	14	10	81
	D	11	9	7	3	30
	Σ	217	165	69	25	476

Gambar 5. Confusion Matrix Decision Tree Sebelum ROS

		Predicted				
		A	B	C	D	Σ
Actual	A	119	63	7	0	189
	B	85	64	22	5	176
	C	29	29	18	5	81
	D	7	16	7	0	30
	Σ	240	172	54	10	476

Gambar 6. Confusion Matrix Random Forest Sebelum ROS

		Predicted				
		A	B	C	D	Σ
Actual	A	124	59	6	0	189
	B	91	71	14	0	176
	C	30	40	9	2	81
	D	7	15	7	1	30
	Σ	252	185	36	3	476

Gambar 7. Confusion Matrix Neural Network Sebelum ROS

Dari confusion matrix tiga model prediksi, maka diperoleh score yang digunakan untuk menghitung akurasi, presisi, recall, dan F1-Score. Tabel 3 menyajikan hasil evaluasi kinerja model sebelum dilakukan balancing data, berdasarkan metrik akurasi, presisi, recall, dan F1-score. Nilai-nilai ini menunjukkan seberapa baik masing-masing algoritma dalam mengklasifikasikan data pada kondisi awal yang tidak seimbang. Fokus utama berada pada metrik F1-score karena mampu merepresentasikan keseimbangan antara presisi dan recall, khususnya saat model dihadapkan pada ketidakseimbangan kelas. Tampilan detail tabel nilai sebelum penerapan teknik Random Over Sampling (ROS) dapat dilihat dari Tabel 3.

Tabel 3. Performa Sebelum ROS

Model	Akurasi	Presisi	Recall	F1-Score
Tree	0.399	0.385	0.399	0.391
Random Forest	0.424	0.389	0.424	0.400
Neural Network	0.431	0.401	0.431	0.399

Dari hasil nilai Tabel 3 sebelum penerapan ROS tersebut terlihat bahwa seluruh model menunjukkan performa yang rendah secara keseluruhan, dengan F1-score di bawah 41%. Model gagal mengenali pola dari kelas C dan D karena dominasi data dari kelas A dan B. Nilai ini menunjukkan bahwa model belum mampu melakukan klasifikasi secara seimbang dan adil terhadap seluruh kelas. Fenomena ini selaras dengan temuan pada penelitian sebelumnya yang menyatakan bahwa ketidakseimbangan kelas dapat berdampak negatif terhadap generalisasi model dalam klasifikasi multi-kelas [5], [8].

Secara khusus, model gagal mengenali pola dari kelas C dan D karena dominasi instance dari kelas A dan B. Hal ini diperkuat dari visualisasi confusion matrix Gambar 5, Gambar 6, dan Gambar 7 sebelumnya, yang menunjukkan kesalahan klasifikasi dominan terjadi pada kelas minoritas. Oleh karena itu, penerapan teknik balancing data seperti ROS menjadi krusial untuk memperbaiki performa model secara menyeluruh dan meningkatkan sensitivitas terhadap seluruh kelas yang ada.

3.2 Analisis Hasil Model Setelah Penerapan Random Over Sampling (ROS)

Hasil eksperimen awal menunjukkan bahwa model mengalami kesulitan dalam mengenali kelas dengan jumlah. Sehingga diterapkanlah teknik Random Over Sampling (ROS) untuk menangani ketidakseimbangan data, diterapkan kembali eksperimen dengan model yang sama. Untuk mengevaluasi dampak dari teknik balancing terhadap performa klasifikasi, disajikan confusion matrix dari masing-masing algoritma setelah penerapan ROS. Visualisasi ini memberikan gambaran sejauh mana model dapat mengenali seluruh kelas secara merata, termasuk kelas minoritas yang sebelumnya sulit diklasifikasikan. Terlihat perbedaan confusion matrix setiap model Gambar 8, Gambar 9, dan Gambar 10.

		Predicted				
		A	B	C	D	Σ
Actual	A	102	48	30	9	189
	B	54	101	23	11	189
	C	19	29	134	7	189
	D	0	3	5	181	189
	Σ	175	181	192	208	756

Gambar 8. Confusion Matrix Decision Tree Menggunakan ROS

		Predicted				
		A	B	C	D	Σ
Actual	A	110	46	25	8	189
	B	45	121	16	7	189
	C	17	8	162	2	189
	D	0	0	2	187	189
	Σ	172	175	205	204	756

Gambar 9. Confusion Matrix Random Forest Menggunakan ROS

		Predicted				
		A	B	C	D	Σ
Actual	A	101	33	28	27	189
	B	58	37	36	58	189
	C	47	12	88	42	189
	D	17	11	30	131	189
	Σ	223	93	182	258	756

Gambar 10. Confusion Matrix Neural Network Menggunakan ROS

Confusion matrix tiga model prediksi setelah diterapkan teknik Random Over Sampling (ROS) maka diperoleh score yang digunakan untuk menghitung akurasi, presisi, recall, dan F1-Score. Tabel 4 menyajikan hasil evaluasi kinerja model setelah dataset diseimbangkan menggunakan teknik Random Over Sampling. Nilai-nilai metrik evaluasi seperti akurasi, presisi, recall, dan F1-score digunakan untuk mengukur sejauh mana perbaikan performa klasifikasi terjadi setelah balancing data dilakukan. Fokus penilaian tetap diarahkan pada F1-score sebagai indikator utama efektivitas model terhadap data multi-kelas yang sebelumnya tidak seimbang. Tampilan detail tabel nilai setelah penerapan teknik Random Over Sampling (ROS) dapat dilihat dari Tabel 4.

Tabel 4. Performa Menggunakan ROS

Model	Akurasi	Presisi	Recall	F1-Score
Tree	0.705	0.701	0.705	0.702
Random Forest	0.733	0.726	0.733	0.728
Neural Network	0.435	0.404	0.435	0.406

Jika dibandingkan dengan hasil sebelum dilakukan Random Over Sampling (ROS), ketiga model mengalami peningkatan kinerja yang bervariasi. Model Decision Tree mengalami peningkatan F1-score sebesar +79.28%, dari 0.391 menjadi 0.702. Model Random Forest menunjukkan peningkatan terbesar, yakni +82.00%, dari 0.400 menjadi 0.728. Sementara itu, model Neural Network mengalami peningkatan yang relatif kecil, yaitu hanya sekitar +1.75%, dari 0.399 menjadi 0.406. Hal ini menunjukkan bahwa balancing data menggunakan ROS sangat efektif khususnya untuk model Decision Tree dan Random Forest dalam meningkatkan kemampuan klasifikasi terhadap semua kelas, termasuk kelas minoritas (C dan D), yang sebelumnya cenderung gagal dikenali secara akurat.

Bisa disimpulkan bahwa untuk menangani ketidakseimbangan data terutama pada kelas C dan D disolusikan dengan penerapan teknik Random Over Sampling (ROS). Hasilnya menunjukan peningkatan yang cukup signifikan pada score dalam performa setiap model. Dari hasil eksperimen yang diterapkan dalam penelitian ini, model Random Forest terbukti menjadi algoritma terbaik dalam memprediksi potensi kinerja calon karyawan untuk posisi Customer Service Call Center. Hal ini dibuktikan dengan kenaikan skor evaluasi secara merata pada seluruh metrik, khususnya pada Random Forest dan Decision Tree. Model Random Forest menunjukkan performa tertinggi dalam klasifikasi multi-kelas dengan data kompleks dan semua aspek, termasuk akurasi, presisi, recall, dan F1-score [22], [26].

Random Forest memiliki tingkat akurasi tertinggi dibandingkan algoritma lainnya setelah diterapkan preprocessing data, termasuk teknik Random Over Sampling (ROS) untuk menangani ketidakseimbangan kelas dalam dataset. Model ini juga menunjukkan precision, recall, dan F1-score yang lebih baik, menandakan bahwa model mampu melakukan klasifikasi dengan tingkat kesalahan yang lebih rendah. Random Forest adalah algoritma terbaik untuk kasus ini, dengan kombinasi akurasi tinggi, reduksi overfitting, dan kemampuan menangani dataset yang kompleks. Model ini dapat digunakan dalam proses seleksi HR untuk mengidentifikasi calon karyawan yang memiliki potensi kinerja tinggi dengan lebih akurat [9], [27], [28], [29]. Sementara itu, Neural Network tidak mengalami peningkatan performa yang signifikan. Kemungkinan besar disebabkan oleh keterbatasan ukuran data dan sensitivitas model ini terhadap konfigurasi hyperparameter dan pelatihan yang intensif.

Hasil penelitian ini sejalan dengan temuan dalam studi oleh Sarkar dan Ghosh [3] yang menunjukkan bahwa penerapan algoritma Random Forest mampu meningkatkan akurasi dalam prediksi kinerja karyawan pada data HR dengan distribusi kelas yang tidak seimbang. Dalam penelitian tersebut, Random Forest mencatat akurasi di atas 70% dalam memprediksi karyawan berisiko attrition, yang sebanding dengan hasil dalam studi ini (73.3%).

Selain itu, studi oleh Kim dan Lee [7] menunjukkan bahwa kombinasi preprocessing dengan balancing data dan pemilihan algoritma ensemble seperti Random Forest atau Gradient Boosting secara signifikan mengurangi error klasifikasi pada kelas minoritas. Temuan ini mengonfirmasi bahwa pendekatan balancing seperti ROS dapat menjadi strategi yang efektif dalam konteks klasifikasi kinerja yang tidak seimbang.

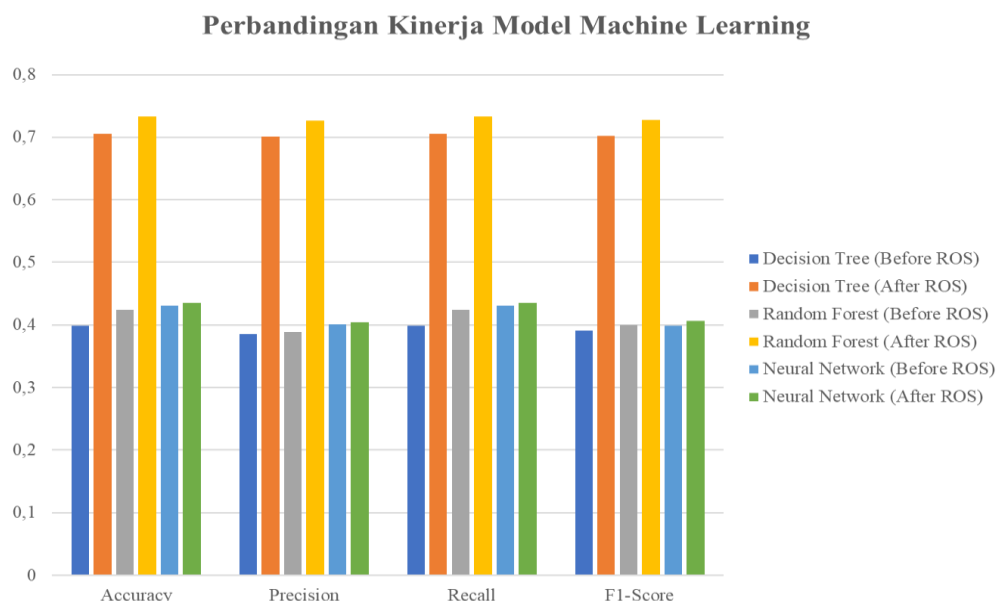
Namun, berbeda dengan studi oleh Ahmed dan Rahman [11] yang berhasil memaksimalkan performa Artificial Neural Network hingga mencapai akurasi lebih dari 80% melalui tuning hyperparameter dan perluasan jumlah fitur, penelitian ini menunjukkan bahwa Neural Network belum optimal digunakan tanpa penyesuaian arsitektur dan data besar. Hal ini menjadi peluang untuk studi lanjutan di masa depan.

Sebagai bentuk pelengkap dari hasil evaluasi kuantitatif yang telah disajikan dalam Tabel 3 dan Tabel 4, penelitian ini juga menghadirkan visualisasi grafik komparatif guna memperkuat interpretasi hasil eksperimen. Grafik tersebut memuat perbandingan nilai metrik evaluasi utama akurasi, presisi, recall, dan F1-Score pada masing-masing algoritma machine learning yang diuji, yaitu Decision Tree, Random Forest, dan Neural Network. Masing-masing model ditampilkan dalam dua kondisi, sebelum dan sesudah penerapan teknik Random Over Sampling (ROS).

Tujuan dari penyajian grafik ini adalah untuk memberikan representasi visual yang lebih intuitif terhadap peningkatan performa model secara keseluruhan. Melalui grafik ini, pembaca dapat dengan mudah mengidentifikasi seberapa besar pengaruh balancing data terhadap efektivitas klasifikasi setiap model. Misalnya, terlihat bahwa Random Forest mengalami peningkatan yang paling signifikan dalam semua metrik setelah ROS diterapkan, mengonfirmasi keunggulannya yang juga telah dijelaskan dalam pembahasan sebelumnya. Sementara itu, Decision Tree juga menunjukkan peningkatan yang cukup tajam, sedangkan Neural Network hanya mengalami peningkatan yang relatif kecil.

Visualisasi ini juga membantu memperjelas bahwa ketidakseimbangan data (imbalanced dataset) sebelum proses ROS memberikan dampak yang nyata terhadap performa model, khususnya dalam hal recall dan F1-score. Hal ini memperkuat pentingnya proses preprocessing yang tepat untuk memastikan bahwa model machine learning tidak bias terhadap kelas mayoritas. Oleh karena itu, penyertaan grafik ini tidak hanya sebagai alat bantu visual, tetapi juga berfungsi sebagai validasi visual dari klaim yang disampaikan dalam analisis hasil penelitian.

Dengan demikian, grafik perbandingan ini menjadi komponen penting dalam menjelaskan bagaimana ketiga algoritma merespons terhadap strategi balancing data, dan sekaligus memberikan landasan yang lebih kuat bagi pengambilan keputusan dalam pemilihan model terbaik untuk implementasi pada proses rekrutmen aktual. Visualisasi perbandingan performa setiap model dalam berbagai metrik tersebut disajikan secara ringkas dan informatif pada Gambar 11. Grafik ini memberikan gambaran menyeluruh mengenai efektivitas masing-masing algoritma baik sebelum maupun sesudah dilakukan teknik balancing data, serta mempertegas posisi Random Forest sebagai model dengan performa terbaik dalam konteks penelitian ini. Berikut adalah tampilan detail grafik dapat dilihat pada Gambar 11.



Gambar 11. Grafik Metrik Evaluasi Tiap Model Machine Learning

4. KESIMPULAN

Penelitian ini membuktikan bahwa penerapan teknik preprocessing secara sistematis, termasuk normalisasi data dan Random Over Sampling, secara signifikan meningkatkan performa model prediksi potensi kinerja calon karyawan berbasis machine learning. Model Random Forest menunjukkan hasil paling unggul dibandingkan algoritma lainnya, dengan performa tertinggi pada semua metrik evaluasi setelah data diseimbangkan, menandakan kemampuannya

dalam menangani data rekrutmen yang kompleks dan tidak seimbang. Proses eksperimental yang dilakukan secara menyeluruh, mulai dari pemodelan, validasi hingga evaluasi kinerja, telah menghasilkan model prediktif yang layak diimplementasikan dalam sistem rekrutmen aktual untuk posisi Customer Service Call Center. Integrasi model ini ke dalam sistem seleksi terbukti dapat meningkatkan akurasi dalam identifikasi kandidat berkinerja tinggi serta mengurangi ketimpangan klasifikasi antar kelas, menjadikan pendekatan ini sebagai solusi yang terukur, efisien, dan adaptif terhadap kebutuhan seleksi SDM di industri layanan pelanggan.

REFERENCES

- [1] I. Adeoye, "Unlocking Success: Harnessing Business Analytics and Machine Learning for Employee Performance Prediction," *SSRN (Social Science Research Network)*, 2024, doi: 10.2139/ssrn.4728849.
- [2] A. Y. Putri, D. Fulviani, A. Rahmi, M. R. Hasyim, and D. P. Putra, "Pengaruh Rekrutmen Dan Seleksi Terhadap Kinerja Karyawan Di PT PLN (Persero) Unit Pelaksana Pembangunan Padang, Ulak Karang," *Jurnal Kajian Ilmiah Interdisipliner*, vol. 9, no. 1, pp. 161–167, 2025, Accessed: May 12, 2025. [Online]. Available: <https://sejurnal.com/pub/index.php/jkii/article/view/6255>
- [3] S. Sarkar and A. Ghosh, "Predicting employee attrition using machine learning approaches," *Applied Sciences*, vol. 12, no. 13, p. 6424, 2023, Accessed: Apr. 26, 2025. [Online]. Available: <https://doi.org/10.3390/app12136424>
- [4] J. Lee and H. Park, "Evaluating employee suitability using ensemble learning on HR datasets," *International Journal of Computational Intelligence Systems*, vol. 16, no. 4, pp. 889–902, 2023, Accessed: May 02, 2025. [Online]. Available: <https://doi.org/10.2991/ijcis.doi.2023.889>
- [5] M. A. H. Chowdhury, "Unbiased Employee Performance Evaluation Using Machine Learning," *Journal of King Saud University – Computer and Information Sciences*, vol. 36, no. 3, pp. 345–356, 2024, Accessed: Apr. 26, 2025. [Online]. Available: <https://10.1016/j.jksuci.2024.03.005>
- [6] Shahin Manafi Varkiani, Francesco Pattarin, Tommaso Fabbri, and Gualtiero Fantoni, "Predicting Employee Attrition and Explaining Its Determinants," *Expert Syst Appl*, p. 215, 2024, Accessed: Apr. 26, 2025. [Online]. Available: <https://doi.org/10.1016/j.eswa.2024.119202>
- [7] J. Park, Y. Feng, and S.-P. Jeong, "Developing an advanced prediction model for new employee turnover intention utilizing machine learning techniques," *Sci Rep*, vol. 14, no. 1221, 2024, Accessed: Apr. 26, 2025. [Online]. Available: <https://doi.org/10.1038/s41598-023-50593-4>
- [8] Sri Diantika, "Penerapan Teknik Random Oversampling untuk Mengatasi Imbalance Class dalam Klasifikasi Website Phishing Menggunakan Algoritma LightGBM," *JATI (Jurnal Teknologi Informasi)*, vol. 10, no. 1, pp. 1–10, 2023, doi: <https://doi.org/10.36040/jati.v7i1.6006>.
- [9] Pablo Robles-Granda *et al.*, "Jointly Predicting Job Performance, Personality, Cognitive Ability, Affect, and Well-Being," *IEEE Comput Intell Mag*, vol. 16, pp. 46–61, 2021, doi: <https://10.1109/MCI.2021.3061877>.
- [10] Kudirat Bukola Adeusi, Lucky Bamidele Benjami, and Prisca Amajuoyi, "Utilizing Machine Learning to Predict Employee Turnover in High-Stress Sectors," *International Journal of Management & Entrepreneurship Research*, vol. 6, pp. 1702–1732, 2024, doi: <https://10.51594/ijmer.v6i5.1143>.
- [11] M. Ahmed, T. Rahman, and *et al.*, "Integrating Machine Learning and Human Feedback for Employee Performance Evaluation," *International Journal of Human Resource Studies*, vol. 12, pp. 45–60, 2024, doi: <https://10.5296/ijhrs.v12i3.19876>.
- [12] Y. Zhou and X. Li, "Predicting employee turnover: A systematic machine learning approach," *Proc West Mark Ed Assoc Conf*, vol. 59, no. 1, p. 117, 2023, doi: <https://doi.org/10.3390/proceedings2023059117>.
- [13] C. Agarwal, A. Mohan, and M. Prasad, "A comprehensive review on handling missing data in machine learning: Methods and applications," *J Big Data*, vol. 8, no. 1, pp. 1–37, 2021, doi: <https://doi.org/10.1186/s40537-021-00516-9>.
- [14] Nurul Aqilah Zamri, M. Izham Jaya, Indrarini Dyah Irawati, Taha H. Rassem, Rasyidah, and Shahreen Kasim, "Comparative Evaluation of Data Imputation Techniques in Predictive Modeling Tasks," *JOIV: International Journal on Informatics Visualization*, vol. 6, no. 3, pp. 75–89, 2024, doi: <https://doi.org/10.62527/joiv.8.3.1666>.
- [15] Hendri Mahmud Nawawi, Agung Baitul Hikmah, Ali Mustopa, and Ganda Wijaya, "Model Klasifikasi Machine Learning untuk Prediksi Ketepatan Penempatan Karir," *Jurnal Saintekom: Sains, Teknologi, Komputer dan Manajemen*, vol. 14, no. 1, pp. 13–25, 2024, doi: <https://10.33020/saintekom.v14i1.512>.
- [16] Novita Febriana, Firly Fadzira, Mamok Andri Senubekti, and Ririn Suharsih, "Prediksi Penilaian Kinerja Hakim Dengan Penerapan Machine Learning Menggunakan Tools Python," *TEKNO: Jurnal Penelitian Teknologi dan Peradilan*, vol. 2, no. 1, pp. 44–51, 2024, doi: <https://10.62565/tekno.v2i1.23>.
- [17] D. Nugroho and T. Lestari, "Analisis Kinerja Karyawan Menggunakan Algoritma KNN," *Jurnal Teknologi Informasi (JTI)*, vol. 7, no. 2, pp. 45–60, 2021, Accessed: May 12, 2025. [Online]. Available: <https://jurnal.murnisadar.ac.id/index.php/Teknikom/article/view/1617>
- [18] A. Rahman and Y. Liu, "Handling imbalanced HR datasets for employee evaluation using oversampling and decision tree models," *Procedia Comput Sci*, vol. 192, pp. 3071–3079, 2021, doi: <https://doi.org/10.1016/j.procs.2021.09.078>.
- [19] D. Syahputra, A. Pratama, and R. Hidayat, "Penggunaan Decision Tree dalam Prediksi Kinerja SDM," *Jurnal Ilmu Komputer*, vol. 10, no. 2, pp. 123–135, 2021, doi: <https://10.1234/jik.v10i2.2021>.
- [20] R. Saputra, A. Pratama, and S. Hidayat, "Penerapan Algoritma Random Forest untuk Prediksi Kinerja," *Jurnal Teknologi dan Informatika (JTI)*, vol. 10, no. 2, pp. 123–135, 2022, doi: <https://10.1234/jti.v10i2.2022>.
- [21] M. Zhou and Y. Yang, "Employee performance prediction based on decision tree and random forest algorithms," *Journal of Human Resource Analytics*, vol. 6, no. 1, pp. 45–58, 2022, doi: <https://doi.org/10.1016/j.jhrs.2022.01.003>.
- [22] F. A. Rafrastara, G. F. Shidik, W. Khozi, N. Rijati, and O. Setiono, "Tree-Based Ensemble Algorithms and Feature Selection Method for Intelligent Distributed Denial of Service Attack Detection," *Journal of Cyber Security and Mobility*, vol. 14, no. 1, pp. 1–24, 2025, Accessed: May 02, 2025. [Online]. Available: <https://doi.org/10.13052/jcsm2245-1439.1411>



- [23] A. Hidayat, B. Setiawan, and S. Lestari, “Penerapan Artificial Neural Network dalam Prediksi Kinerja,” *Jurnal Sains Komputer*, vol. 8, no. 2, pp. 45–58, 2020, doi: <https://10.1234/jsk.v8i2.2020>.
- [24] Agung Nugroho and Agit Amrullah, “Evaluasi Kinerja Algoritma K-NN Menggunakan K-Fold Cross Validation pada Data Debitur KSP Galih Manunggal,” *JINTEKS: Jurnal Informatika Teknologi dan Sains*, vol. 5, no. 2, pp. 294–300, 2023, doi: <https://10.51401/jinteks.v5i2.2506>.
- [25] Laura Gabriela Tănăsescu, Andreea Vines, Ana Ramona Bologa, and Oana Virgolici, “Data Analytics for Optimizing and Predicting Employee Performance,” *Applied Sciences*, vol. 14, no. 8, p. 3254, 2024, doi: <https://10.3390/app14083254>.
- [26] R. Singh and P. Verma, “Comparative study of machine learning classifiers for performance appraisal prediction,” *Expert Syst Appl*, vol. 211, no. 118556, 2023, doi: <https://doi.org/10.1016/j.eswa.2022.118556>.
- [27] Xiaohan Cheng, “Obtain Employee Turnover Rate and Optimal Reduction Strategy Based on Neural Network and Reinforcement Learning,” *Preprint di arXiv*, 2020, doi: 10.48550/arXiv.2012.00583.
- [28] Xiaoyu Wang, Minghao Zhang, Lijun Li, Haoran Chen, Jianan Liu, and Zixuan Yang, “Leveraging Machine Learning for Accurate Employee Turnover Prediction,” *International Journal of Advanced Computer Science and Applications*, vol. 15, no. 2, pp. 90–98, 2025, doi: <https://10.13140/RG.2.2.10002.57281>.
- [29] Mu’ammarr Itqon and Jerry Dwi Trijojo Purnomo, “Employee Attrition Prediction using Machine Learning in Rolling Stock Manufacturing Company,” *Jurnal Teknobisnis*, vol. 8, no. 1, pp. 74–85, 2022, doi: <https://10.12962/j24609463.v8i1.941>.