

KLASIFIKASI KUALITAS KINERJA REKANAN PENYEDIA BARANG MENGUNAKAN ALGORITMA *K-NEAREST NEIGHBORS* DI PT KRAKATAU BANDAR SAMUDERA

Ratu Keisha Jasmine Deanova¹, Luluk Muthoharoh^{2*}, Syah Rezzal³

^{1,2} Program Studi Sains Data, Institut Teknologi Sumatera, Lampung Selatan, Indonesia

³ PT Krakatau Bandar Samudera, Cilegon, Indonesia

*Penulis korespondensi: luluk.muthoharoh@sd.itera.ac.id

ABSTRAK

PT Krakatau Bandar Samudera (PT KBS) merupakan anak perusahaan PT Krakatau Steel yang bergerak di bidang jasa kepelabuhanan dan logistik dengan dukungan sistem *e-Procurement* dalam proses pengadaan barang dan jasa. Evaluasi kinerja rekanan menjadi aspek kritis untuk menjamin objektivitas dan kualitas proses pengadaan, namun proses penilaian yang dilakukan secara manual rentan terhadap subjektivitas dan ketidakkonsistenan. Oleh karena itu, penelitian ini mengimplementasikan algoritma *K-Nearest Neighbors* (K-NN) untuk mengklasifikasikan kualitas kinerja rekanan penyedia barang di PT KBS secara otomatis berbasis data. Dataset dibagi dengan proporsi 80:20 untuk data latih dan data uji, serta nilai parameter k ditentukan melalui pengujian sistematis terhadap beberapa nilai k . Hasil penelitian menunjukkan bahwa model K-NN dengan $k = 6$ memberikan kinerja klasifikasi terbaik dengan akurasi sebesar 91,67%. Evaluasi lebih lanjut menggunakan *confusion matrix* menghasilkan nilai presisi sebesar 86,11%, *recall* sebesar 91,67%, dan *F1-Score* sebesar 88,33%, yang mengindikasikan bahwa model memiliki performa klasifikasi yang baik dan seimbang. Dengan demikian, model K-NN terbukti efektif sebagai pendekatan klasifikasi otomatis yang dapat mendukung pengambilan keputusan berbasis data dalam penilaian dan seleksi rekanan di PT KBS.

Kata kunci: *confusion matrix*, *e-procurement*, evaluasi kinerja rekanan, klasifikasi, *k-nearest neighbors*.

1 PENDAHULUAN

PT Krakatau Bandar Samudera (PT KBS) merupakan perusahaan yang bergerak di bidang jasa kepelabuhanan dan logistik yang menerapkan sistem *e-Procurement* berbasis web melalui portal KIPOP (*Krakatau International Port Online Procurement*) untuk mendukung proses pengadaan barang dan jasa (Harianto, 2025.). Dalam proses tersebut, rekanan memiliki peran penting dalam menjamin ketersediaan serta kualitas barang dan jasa yang dibutuhkan perusahaan. Oleh karena itu, evaluasi kinerja rekanan perlu dilakukan secara berkelanjutan untuk memastikan bahwa rekanan yang bekerja sama mampu memenuhi standar kualitas yang telah ditetapkan.

Seiring bertambahnya jumlah rekanan dan data evaluasi yang dihasilkan, proses penilaian secara manual menjadi kurang efektif dan berpotensi menimbulkan subjektivitas dalam pengambilan keputusan. Berbagai penelitian menunjukkan bahwa metode *K-Nearest Neighbors* (K-NN) memiliki performa yang baik dalam permasalahan klasifikasi. Penelitian klasifikasi kinerja karyawan menggunakan K-NN dengan teknik SMOTE memperoleh akurasi 94% (Nuraeni dkk., 2024), sedangkan klasifikasi prestasi akademik siswa menggunakan K-NN dan *K-Fold Cross Validation* menghasilkan akurasi tertinggi 91,39% (Muhaimin dkk., 2024). Penelitian lain menunjukkan bahwa penerapan Gain Ratio pada K-NN mampu meningkatkan akurasi klasifikasi kinerja siswa dari 74,07% menjadi 75,11 (Rahmat dkk., 2020). Selain itu, K-NN berhasil diterapkan pada klasifikasi seleksi calon karyawan baru dengan akurasi 91% (Adhitya dkk., 2020) serta klasifikasi data penjualan produk untuk mendukung perencanaan

stok Perusahaan (Novitadewi dkk., 2023). Hasil-hasil tersebut menunjukkan bahwa K-NN mampu memberikan performa klasifikasi yang baik pada berbagai bidang dan jenis data.

Penerapan metode K-NN untuk klasifikasi kinerja rekanan pada sektor kepelabuhanan dan logistik masih relatif terbatas. Oleh karena itu, penelitian ini bertujuan untuk mengklasifikasikan kualitas kinerja rekanan penyedia barang di PT Krakatau Bandar Samudera menggunakan metode K-NN. Metode ini dipilih karena mampu melakukan klasifikasi berdasarkan kemiripan data tanpa memerlukan asumsi distribusi tertentu, dengan menentukan kelas suatu data berdasarkan mayoritas kelas dari k tetangga terdekatnya. Dalam penelitian ini, nilai k diuji secara bertahap untuk memperoleh model terbaik, sedangkan performa model dievaluasi menggunakan metrik akurasi, *precision*, *recall*, dan *F1-Score*. Hasil penelitian diharapkan dapat membantu perusahaan dalam meningkatkan objektivitas, efisiensi, dan akurasi evaluasi kinerja rekanan sehingga mendukung proses pengambilan keputusan terkait pemilihan mitra kerja yang berkinerja optimal.

2 METODE

2.1 Data

Data yang digunakan pada penelitian ini adalah data sekunder, yang dikumpulkan langsung dari formulir survei kepuasan kinerja rekanan pada periode Januari-Juni 2024 yang diisi oleh setiap *user* atau unit kerja yang ada di PT Krakatau Bandar Samudera. Penelitian ini menggunakan lima variabel prediktor, yaitu kualitas barang (x_1), waktu pengiriman (x_2), report pengiriman (x_3), respon terhadap komplain (x_4), dan nilai akhir evaluasi (x_5), serta satu variabel target (y) berupa nilai huruf yang merepresentasikan kategori kinerja rekanan.

Tabel 1. Deskripsi Data

Jenis Variabel	Nama Variabel	Tipe Data
x_1	Kualitas Barang	<i>float</i>
x_2	Waktu Pengiriman	<i>float</i>
x_3	Report Pengiriman	<i>float</i>
x_4	Respon Terhadap Komplain	<i>float</i>
x_5	Nilai Akhir	<i>float</i>
y	Nilai Huruf	<i>string</i>

2.2 Preprocessing Data

Preprocessing Data bertujuan untuk mengubah data mentah menjadi data yang berkualitas sehingga data layak untuk diolah pada tahapan selanjutnya. Tahapan ini dilakukan pada data mentah untuk menghilangkan data yang bermasalah atau inkonsisten (Wilujeng dkk., 2023). *Preprocessing Data* yang dilakukan meliputi pembersihan data, normalisasi data, dan transformasi data.

2.3 Klasifikasi

Klasifikasi adalah suatu metode analisis data yang digunakan untuk mengidentifikasi pola-pola yang menggambarkan kelas-kelas dalam data. Algoritma klasifikasi bertujuan untuk memprediksi kategori kelas dari suatu data dan mengklasifikasikan data tersebut ke dalam kelas yang sesuai berdasarkan label yang telah dikenali sebelumnya (Nugroho & Religia, 2021). Penelitian ini melakukan klasifikasi pada data berdasarkan dokumen Laporan Evaluasi Kinerja Rekanan PT Krakatau Bandar Samudera yang disusun oleh *Procurement Department* PT Krakatau Bandar Samudera dengan pembagian klasifikasi kelas dari nilai akhir penilaian setiap rekanan pada Tabel 2.

Tabel 2. Klasifikasi Nilai Akhir Kinerja Rekanan PT Krakatau Bandar Samudera

Huruf	Nilai	Keterangan
-------	-------	------------

A	90-100	Sangat Melebihi Ekspetasi
B	80-89,9	Melebihi Ekspetasi
C	70-79,9	Sesuai Ekspetasi
D	60-69,9	Kurang Dari Ekspetasi
E	0-59,9	Tidak Sesuai Ekspetasi

2.4 Model Klasifikasi

K-Nearest Neighbor (K-NN) adalah sebuah metode untuk melakukan klasifikasi terhadap objek berdasarkan data latih (*neighbor*) yang jaraknya paling dekat dengan objek tersebut (Nugroho & Religia, 2021). Dekat atau jauhnya *neighbor* biasanya dihitung berdasarkan jarak *Euclidean*. diperlukan suatu sistem klasifikasi sebagai sebuah sistem yang mampu mencari informasi. Berikut ini adalah langkah-langkah algoritma K-NN (Roihan dkk., 2020).

- 1) Menentukan nilai k
- 2) Nilai k dapat dihitung menggunakan persamaan
- 3) Melakukan perhitungan nilai jarak (*euclidean distance*) terhadap masing-masing objek data yang diberikan. Rumus untuk menghitung *euclidean distance* dapat dilihat pada Persamaan (1).

$$d_i = \sqrt{\sum_{j=1}^n (x_{ij} - x_{ji})^2} \quad (1)$$

Keterangan:

d_i : jarak *Euclidean*

x_{ij} : data latih ke- i

x_{ji} : data uji ke- j

- 4) Melakukan pengelompokkan data sesuai dengan perhitungan jarak (*euclidean distance*)
- 5) Melakukan pengelompokkan data sesuai dengan nilai tetangga terdekat (*nearest neighbor*) atau berdasarkan data yang mempunyai jarak *Euclidean* terkecil.
- 6) Memilih nilai mayoritas dari tetangga terdekat sebagai hasil klasifikasi.

Nilai k yang terbaik untuk K-NN tergantung pada data. Secara umum, nilai k yang tinggi akan mengurangi efek *noise* pada klasifikasi, tetapi membuat batasan antara setiap klasifikasi menjadi lebih kabur. Nilai k yang bagus dapat dipilih dengan optimasi parameter, misalnya dengan menggunakan *cross-validation*. Kasus khusus di mana klasifikasi diprediksikan berdasarkan data pembelajaran yang paling dekat (dengan kata lain, $k = 1$) disebut algoritma *nearest neighbor*.

2.5 Evaluasi Model

2.5.1 AUC (Area Under Curve)

AUC atau *Area Under Curve* merupakan sebuah area dibawah kurva yang menggambarkan evaluasi mengenai kemampuan metode klasifikasi membedakan antar kelas (Hansen & Hariyanto, 2023). Semakin tinggi AUC maka kinerja model tersebut semakin baik dalam membedakan diantara kelas positif dan negatif. Nilai AUC dapat dibagi menjadi beberapa kelompok yang disajikan pada Tabel 3.

Tabel 3. Interpretasi Nilai AUC

Nilai AUC	Keterangan
0.90 – 1.00	<i>Excellent Classification</i>
0.80 – 0.90	<i>Good Classification</i>
0.70 – 0.80	<i>Fair Classification</i>
0.60 – 0.70	<i>Poor Classification</i>

0.50 – 0.60

Failure

Sumber: (Wilujeng dkk., 2023)

2.5.2 Confusion Matrix

Tahap evaluasi adalah proses untuk menilai kinerja model dan perhitungan yang telah dilakukan dengan mengukur performa algoritma klasifikasi yang digunakan. Dalam penelitian ini, performa diukur menggunakan metode *confusion matrix*. *Confusion matrix* merupakan metode evaluasi yang banyak digunakan untuk mengukur kinerja suatu algoritma klasifikasi. *Confusion matrix* mengandung informasi perbandingan hasil klasifikasi yang diprediksi sistem dan hasil yang seharusnya (Baharuddin dkk., 2019). *Confusion matrix* merupakan sebuah teknik yang mudah dan efektif dalam mengukur kinerja sistem klasifikasi. Tujuan utama dari *confusion matrix* adalah untuk menilai performa atau akurasi dari suatu sistem klasifikasi dalam mengklasifikasikan data uji (Dwi Fasnuari dkk., 2022). Evaluasi dengan metode *confusion matrix* bertujuan untuk mengetahui nilai akurasi, presisi, *recall*, dan F1-Score juga mengidentifikasi kekuatan dan kelemahan model. Tabel 4 menunjukkan *confusion matrix* yang diterapkan pada penelitian ini untuk mengevaluasi hasil klasifikasi berdasarkan perbandingan antara kelas aktual dan kelas prediksi.

Tabel 4. Confusion Matrix

Aktual	Prediksi			
	A	B	C	E
A	C _{AA}	C _{AB}	C _{AC}	C _{AD}
B	C _{BA}	C _{BB}	C _{BC}	C _{BD}
C	C _{CA}	C _{CB}	C _{CC}	C _{CD}
E	C _{DA}	C _{DB}	C _{DC}	C _{DD}

Misalkan Matriks disebut sebagai C, di mana:

i : Baris (Prediksi)

j : Kolom (Aktual)

Maka rumus evaluasinya berubah menjadi lebih universal.

a) Akurasi

Akurasi menggambarkan tingkat ketepatan sistem dalam mengklasifikasikan data secara benar. Nilai akurasi diperoleh dari perbandingan antara jumlah data yang terklasifikasi benar dengan total keseluruhan data dan dirumuskan pada Persamaan (2).

$$Akurasi = \frac{\sum_i C_{ii}}{\sum_i \sum_j C_{ij}} \quad (2)$$

b) Presisi

Presisi adalah nilai dari berapa banyak data dengan kategori terklasifikasi benar dibandingkan dengan total keseluruhan data kategori terklasifikasi benar dan dirumuskan pada Persamaan (3).

$$Presisi_i = \frac{C_{ii}}{\sum_j C_{ij}} \quad (3)$$

c) Recall

Recall dilakukan untuk mengetahui perbandingan jumlah data kategori terklasifikasi benar oleh sistem dengan jumlah data kategori terklasifikasi benar dan salah dan dirumuskan pada Persamaan (4).

$$Recall_j = \frac{C_{jj}}{\sum_i C_{ij}} \quad (4)$$

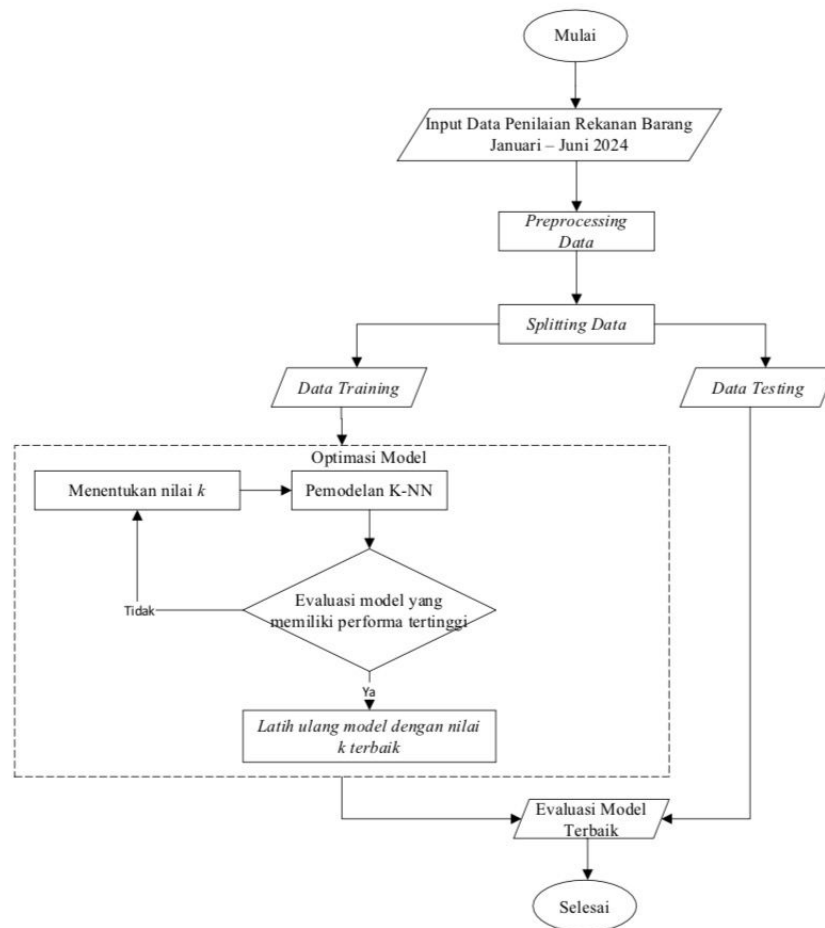
d) F1-Score

F1-Score merupakan perbandingan rata-rata presisi dan *recall* yang dibobotkan dan dirumuskan pada Persamaan (5).

$$F1\ Score = \frac{1}{4} \sum_{i=1}^4 \frac{2 \times Presisi_i \times Recall_j}{Presisi_i + Recall_j} \quad (5)$$

2.6 Alur Penelitian

Proses penelitian diawali dengan pengumpulan dan *preprocessing data* untuk memastikan kualitas data yang digunakan. Selanjutnya, data dibagi menjadi *data training* dan *data testing*, kemudian dilakukan pembangunan model menggunakan algoritma *K-Nearest Neighbors* (K-NN). Model yang dihasilkan dievaluasi untuk mengukur kinerjanya, dan apabila belum memenuhi kriteria yang ditetapkan, proses pelatihan dan evaluasi dilakukan kembali. Alur penelitian secara lengkap ditunjukkan pada Gambar 1.



Gambar 1. Alur Penelitian

3 HASIL DAN PEMBAHASAN

3.1 *Preprocessing Data*

Tahap *preprocessing data* dilakukan melalui proses transformasi data menggunakan teknik *Label Encoding* untuk mengubah variabel target ke dalam bentuk numerik. Pada tahap

ini, kategori “A”, “B”, “C”, “D”, dan “E” masing-masing dikonversi menjadi nilai 1, 2, 3, 4, dan 5. Transformasi ini bertujuan agar data kategorikal dapat diproses oleh algoritma *K-Nearest Neighbors* (K-NN) yang memerlukan data numerik dalam perhitungan jarak antar data.

3.2 *K-Nearest Neighbors* (K-NN)

3.2.1 Pemodelan *K-Nearest Neighbors* (K-NN)

Dalam proses pencarian nilai k terbaik dilakukan beberapa percobaan dengan menggunakan nilai $k = 1$ hingga $k = 10$ dan pembagian data latih dan data uji dengan proporsi 80:20. Berdasarkan hasil pemodelan pada data latih, diperoleh hasil yang menunjukkan variasi akurasi tergantung pada nilai k yang dipilih. Pada nilai $k = 1$, model berhasil memberikan akurasi sempurna sebesar 100% yang mengindikasikan adanya *overfitting*. Penurunan akurasi mulai terjadi ketika nilai k meningkat, seperti pada $k = 2$ yang menghasilkan akurasi sebesar 93.75%. Pada nilai k antara 3 hingga 8, model memberikan akurasi yang relatif stabil, yakni sekitar 97.92%, yang mengindikasikan bahwa model mulai mempertimbangkan lebih banyak tetangga dalam proses klasifikasi, sehingga keputusan yang diambil menjadi lebih *robust* dan kurang dipengaruhi oleh fluktuasi data. Namun, pada nilai $k = 10$, terdapat penurunan akurasi yang cukup signifikan menjadi 89.58%, yang mengindikasikan bahwa model mulai kehilangan sensitivitas terhadap pola-pola penting dalam data latih, seiring dengan semakin banyaknya tetangga yang dipertimbangkan. Oleh karena itu, berdasarkan hasil analisis pada data latih, proses pemodelan K-NN pada data uji akan menggunakan nilai k antara 3 hingga 8, karena nilai-nilai tersebut memberikan akurasi yang stabil dan menunjukkan kemampuan model untuk menggeneralisasi dengan baik, tanpa terpengaruh oleh *overfitting* yang terjadi pada nilai k yang lebih kecil atau kehilangan sensitivitas pada nilai k yang lebih besar. Berikut ini merupakan hasil pemodelan pada data latih yang disajikan pada Tabel 5.

Tabel 5. Akurasi Model di Setiap k pada Data Latih

Nilai k	Nilai Akurasi
1	100%
2	93.75%
3	97.92%
4	97.92%
5	97.92%
6	95.83%
7	97.92%
8	97.92%
9	97.92%
10	89.58%

Berdasarkan hasil analisis akurasi pada data latih, nilai k antara 3 hingga 8 menunjukkan performa yang stabil dengan akurasi sekitar 97.92%, yang mencerminkan bahwa model *K-Nearest Neighbors* (K-NN) dapat mengklasifikasikan data latih dengan baik. Namun pada akurasi model data uji, nilai k yang kecil, yaitu $k = 3$, menunjukkan penurunan akurasi yang signifikan menjadi 58.33%, yang mengindikasikan bahwa menjadi terlalu sensitif terhadap variasi kecil dalam data. Pada $k = 4$ hingga $k = 8$, nilai akurasi meningkat dan konsisten sebesar 91.67%, yang mengindikasikan model berada dalam keadaan seimbang antara sensitivitas terhadap data dikarenakan lebih banyak tetangga di sekitar data, sehingga keputusan klasifikasi menjadi lebih stabil dan tidak terpengaruh oleh *noise*.

Tabel 4 menampilkan hasil pemodelan pada data uji untuk menentukan nilai k terbaik, yaitu nilai k terbaik yang dipilih adalah $k = 6$. Pemilihan nilai $k = 6$ sebagai parameter optimal dalam model K-NN didasarkan pada pertimbangan keseimbangan antara akurasi, AUC, dan kecepatan komputasi. Meskipun nilai $k = 7, 8$, dan 9 memberikan hasil yang serupa dalam hal

akurasi dan AUC, nilai $k = 6$ dipilih karena memberikan keseimbangan yang optimal antara kompleksitas model dan performa. Pemilihan nilai k yang tepat dapat mempengaruhi kinerja model K-NN dalam klasifikasi (Ainurrohman & Wiyanti, 2023). Nilai k yang lebih kecil cenderung memberikan waktu komputasi yang lebih efisien tanpa mengorbankan akurasi secara signifikan dengan nilai $k = 6$, dengan mempertimbangkan *trade-off* antara *overfitting* dan *underfitting* (Puteri dkk., 2023). Oleh karena itu, penelitian ini akan menggunakan nilai $k = 6$ untuk memodelkan K-NN. Berikut merupakan hasil percobaan untuk menentukan nilai k yang paling baik disajikan dalam Tabel 6.

Tabel 6. Akurasi model pada data uji

Nilai k	Nilai Akurasi	Nilai AUC
1	83.33%	0.9
2	75%	0.8472
3	58.33%	0.912
4	91.67%	0.9565
5	91.67%	0.975
6	91.67%	0.9833
7	91.67%	0.9833
8	91.67%	0.9833
9	83.33%	0.9833
10	75%	0.9833

3.2.2 Perhitungan Jarak Euclidean

Pemilihan tetangga terdekat pada model K-NN dilakukan dengan menghitung jarak *euclidean* antara satu sampel uji dan seluruh observasi pada data latih berdasarkan lima variabel prediktor yang digunakan, yaitu kualitas barang (x_1), waktu pengiriman (x_2), report pengiriman (x_3), respon terhadap komplain (x_4), dan nilai akhir evaluasi (x_5). Selanjutnya, seluruh jarak diurutkan dari nilai terkecil hingga terbesar untuk menentukan peringkat kedekatan. Dalam penelitian ini ditetapkan $k = 6$, sehingga enam observasi dengan jarak terpendek ditandai sebagai tetangga terdekat (label “YA”) dan menjadi dasar pengambilan keputusan klasifikasi melalui majority voting, sedangkan observasi lainnya (label “TIDAK”) tidak dilibatkan dalam penetapan kelas. Berdasarkan komposisi kelas pada enam tetangga terdekat yang tercantum pada Tabel 7, kelas yang paling dominan ditetapkan sebagai prediksi akhir untuk data uji. Hasil perhitungan menunjukkan bahwa kelas prediksi kinerja rekanan mengikuti kelas mayoritas dari enam tetangga terdekat tersebut.

Tabel 7. Perhitungan Jarak Euclidean

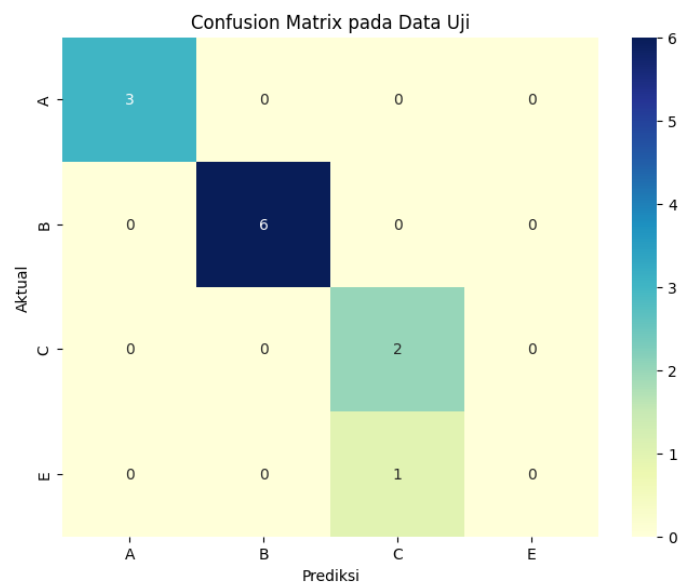
Nama Rekanan	Kelas						Jarak Euclid	Urutan	Tetangga
Rekanan Barang 5	89.2	90.7	87.45	87.45	88.7	2	0.988502	1	YA
Rekanan Barang 3	89.1	89.4	88.6	86.6	88.43	2	1.073676	2	YA
Rekanan Barang 4	94	93.8	90	90	92.25	1	1.287553	3	YA
Rekanan Barang 59	90	85	85	90	87.5	2	1.585405	4	YA
Rekanan Barang 60	85	84	86	81	85.5	2	3.235555	5	YA
Rekanan Barang 2	85	76.3	90	90	82.5	2	3.283054	6	TIDAK
Rekanan Barang 58	80	85	85	85	83.75	2	3.461491	7	TIDAK
Rekanan Barang 57	87.3	76.3	89.25	82.55	83.01	2	3.649169	8	TIDAK
Rekanan Barang 56	80	80	75	80	78.75	3	6.016013	9	TIDAK

3.3 Evaluasi Model

Evaluasi kinerja model dilakukan menggunakan *confusion matrix* serta metrik akurasi, presisi, *recall*, dan F1-Score. Meskipun penelitian ini menggunakan lima kategori kinerja

rekanan, yaitu A, B, C, D, dan E, data uji yang digunakan tidak mengandung sampel kelas D sehingga *confusion matrix* pada Gambar 2 hanya menampilkan hasil evaluasi untuk empat kelas, yaitu A, B, C, dan E, dengan total 12 data uji. Model menunjukkan performa yang sangat baik pada tiga kelas pertama, yaitu kelas B yang berhasil diklasifikasikan dengan benar untuk seluruh 6 data ujinya, kelas A untuk seluruh 3 data ujinya, dan kelas C untuk seluruh 2 data ujinya. Ketiga kelas tersebut mencapai nilai presisi dan recall sebesar 100% tanpa kesalahan klasifikasi. Sementara itu, kelas E mengalami kesalahan klasifikasi pada satu-satunya data ujinya yang diprediksi sebagai kelas C, sehingga menghasilkan satu false negative pada kelas E sekaligus satu false positive pada kelas C. Dengan demikian, dari total 12 data uji hanya terdapat satu kesalahan klasifikasi, sedangkan seluruh data pada kelas lainnya berhasil diprediksi dengan benar. Tidak ditemukan kesalahan klasifikasi silang pada kelas-kelas lain yang ditunjukkan oleh nilai nol pada seluruh elemen *off-diagonal* lainnya dalam *confusion matrix*.

Pola ini sejalan dengan temuan umum pada klasifikasi data dengan jumlah sampel per kelas yang tidak merata, di mana kelas dengan representasi data uji paling sedikit cenderung paling rentan terhadap kesalahan klasifikasi, sekalipun performa keseluruhan model tetap tinggi (Nuraeni dkk., 2024). Hal ini berkaitan dengan karakteristik algoritma K-NN yang menentukan kelas berdasarkan mayoritas suara dari k tetangga terdekat, ketika satu kelas hanya direpresentasikan oleh sedikit data dalam ruang fitur, peluang tetangga terdekatnya didominasi oleh data dari kelas lain menjadi lebih besar, terutama pada nilai k yang relatif besar seperti $k = 6$ yang digunakan pada penelitian ini. Dibandingkan dengan (Nuraeni dkk., 2024) yang menerapkan K-NN pada kasus klasifikasi kinerja karyawan dengan data tidak seimbang dan memperoleh akurasi 94% setelah penanganan menggunakan teknik SMOTE, penelitian ini memperoleh akurasi yang sedikit lebih rendah (91,67%) namun tanpa penanganan ketidakseimbangan kelas secara eksplisit menunjukkan bahwa penerapan teknik *oversampling* seperti SMOTE berpotensi turut meningkatkan kemampuan model dalam mengenali kelas dengan data terbatas seperti kelas E pada penelitian ini.



Gambar 2. Hasil *Confusion Matrix*

Berdasarkan hasil evaluasi model yang disajikan pada Tabel 8, algoritma *K-Nearest Neighbors* (K-NN) dengan nilai $k = 6$ menunjukkan akurasi sebesar 91.67%, yang mencerminkan bahwa model ini dapat memprediksi dengan tepat sebagian besar data uji. Nilai presisi sebesar 86.11% mengindikasikan bahwa model cukup baik dalam mengidentifikasi

kelas positif, meskipun ada sedikit kesalahan dalam prediksi kelas positif tersebut. *Recall* yang mencapai 91.67% menunjukkan bahwa model memiliki kemampuan yang baik dalam menangkap sebagian besar kelas positif yang ada pada data uji. Dengan F1-Score sebesar 88.33%, yang merupakan kombinasi dari presisi dan *recall*, model ini menunjukkan performa yang seimbang dan dapat diandalkan dalam klasifikasi data secara keseluruhan.

Tabel 8. Evaluasi Model K-NN Terbaik

Metrik Evaluasi	Performa
Akurasi	91.67%
Presisi	86.11%
<i>Recall</i>	91.67%
F1-Score	88.33%

Hasil akurasi sebesar 91,67% yang diperoleh pada penelitian ini sejalan dan kompetitif dibandingkan dengan penelitian terdahulu yang juga menerapkan algoritma *K-Nearest Neighbors* pada permasalahan klasifikasi serupa. Nuraeni dkk. (2024) yang menerapkan *K-NN* dengan teknik SMOTE pada kasus klasifikasi kinerja karyawan memperoleh akurasi sebesar 94%, sedangkan (Muhaimin dkk., 2024) pada klasifikasi prestasi akademik siswa menggunakan *K-NN* dengan *K-Fold Cross Validation* memperoleh akurasi tertinggi sebesar 91,39%. Dibandingkan kedua penelitian tersebut, akurasi penelitian ini berada pada rentang yang sebanding, sedikit di bawah (Nuraeni dkk., 2024) namun sedikit di atas (Muhaimin dkk., 2024). Perbedaan capaian (Nuraeni dkk., 2024) yang lebih tinggi diduga turut dipengaruhi oleh penerapan teknik SMOTE untuk menangani ketidakseimbangan kelas, yang tidak diterapkan secara eksplisit pada penelitian ini. Temuan ini memperkuat bahwa algoritma *K-Nearest Neighbors* konsisten menunjukkan performa yang baik untuk permasalahan klasifikasi berbasis data evaluasi atau penilaian, sekaligus membuka peluang peningkatan performa pada penelitian selanjutnya melalui penanganan ketidakseimbangan kelas.

Satu-satunya kesalahan klasifikasi pada data uji terjadi pada kelas E, yang diprediksi model sebagai kelas C. Hal ini kemungkinan disebabkan oleh kemiripan karakteristik skor antara rekanan berkategori E (0–59,9) dan rekanan dengan skor mendekati batas ambang bawah kategori C (70–79,9), khususnya apabila salah satu atau lebih variabel prediktor (x_1 – x_5) pada rekanan tersebut bernilai relatif tinggi meski skor akhirnya tergolong rendah. Karena *K-NN* menentukan kelas berdasarkan jarak *Euclidean* dan suara mayoritas dari k tetangga terdekat, rekanan dengan kombinasi nilai variabel yang berada di wilayah tumpang tindih antar kelas berisiko lebih besar terklasifikasi ke kelas tetangganya. Mengingat kelas E hanya direpresentasikan oleh satu data pada data uji, kesalahan ini juga sejalan dengan kecenderungan umum bahwa kelas dengan jumlah sampel paling sedikit lebih rentan terhadap kesalahan klasifikasi dibandingkan kelas dengan representasi data yang lebih banyak (Nuraeni dkk., 2024).

4 KESIMPULAN

Berdasarkan hasil penelitian, algoritma *K-Nearest Neighbors* (K-NN) berhasil diimplementasikan untuk klasifikasi kinerja rekanan penyedia barang di PT Krakatau Bandar Samudera dengan performa terbaik pada nilai $k = 6$ dan proporsi *data training-testing* sebesar 80:20. Model yang dihasilkan memperoleh akurasi sebesar 91,67%, presisi 86,11%, *recall* 91,67%, dan F1-Score 88,33%, yang menunjukkan bahwa K-NN memiliki kemampuan klasifikasi yang sangat baik, stabil, dan efektif. Satu-satunya kesalahan klasifikasi terjadi pada kelas ‘Tidak Sesuai Ekspetasi’ (E), yang merupakan kelas dengan jumlah data uji paling sedikit,

sehingga penanganan ketidakseimbangan kelas berpotensi meningkatkan performa model lebih lanjut.

Penelitian ini memiliki beberapa keterbatasan yang perlu menjadi perhatian pada penelitian selanjutnya. Pertama, data yang digunakan terbatas pada periode Januari–Juni 2024 dengan jumlah data uji yang relatif kecil, sehingga belum mencerminkan variasi kinerja rekanan secara menyeluruh sepanjang tahun. Kedua, nilai parameter k optimal dipilih dari rentang terbatas ($k = 1$ hingga 10) tanpa eksplorasi nilai k yang lebih luas. Ketiga, ketidakseimbangan jumlah data antar kelas, khususnya pada kelas E, belum ditangani secara eksplisit menggunakan teknik seperti SMOTE atau *class weighting*.

Hasil penelitian menunjukkan bahwa metode *K-Nearest Neighbors* (K-NN) dapat dimanfaatkan sebagai pendukung pengambilan keputusan yang lebih objektif dalam proses evaluasi kinerja rekanan penyedia barang. Untuk penelitian selanjutnya, disarankan melakukan perbandingan performa K-NN dengan algoritma klasifikasi lainnya, seperti *Naive Bayes*, *Random Forest*, dan *Decision Tree*. Selain itu, penelitian berikutnya dapat mengeksplorasi penggunaan metrik jarak selain *Euclidean*, seperti *Manhattan* dan *Minkowski*, serta menerapkan teknik penyeimbangan data dan optimasi parameter k untuk meningkatkan performa klasifikasi, khususnya pada kelas dengan jumlah data yang terbatas.

UCAPAN TERIMA KASIH

Penulis menyampaikan ucapan terima kasih kepada semua pihak yang telah memberikan dukungan, bimbingan, serta masukan selama pelaksanaan penelitian dan penyusunan artikel ini. Bantuan yang diberikan, baik dalam proses pengumpulan data, pengolahan data, maupun penyelesaian artikel, sangat berarti bagi penulis hingga penelitian ini dapat terselesaikan dengan baik.

DAFTAR PUSTAKA

- Ainurrohman, A., & Wiyanti, D. T. (2023). Analisis Performa Algoritma Decision Tree, Naive Bayes, K-Nearest Neighbor untuk Klasifikasi Zona Daerah Risiko Covid-19 di Indonesia. *Jurnal Teknologi Informasi Dan Ilmu Komputer*, 10(1), 115–122. <https://doi.org/10.25126/jtiik.2023105935>
- Baharuddin, M. M., Azis, H., & Hasanuddin, T. (2019). Analisis Performa Metode K-Nearest Neighbor Untuk Identifikasi Jenis Kaca. *ILKOM Jurnal Ilmiah*, 11(3), 269–274. <https://doi.org/10.33096/ilkom.v11i3.489.269-274>
- Dwi Fasnuari, H. A., Yuana, H., & Chulkamdi, M. T. (2022). Penerapan Algoritma K-Nearest Neighbor Untuk Klasifikasi Penyakit Diabetes Melitus. *Antivirus : Jurnal Ilmiah Teknik Informatika*, 16(2), 133–142. <https://doi.org/10.35457/antivirus.v16i2.2445>
- Hansen, & Hariyanto, S. (2023). Perbandingan Algoritma Data Mining Dalam Mengklasifikasi Penyakit Diabetes Menggunakan Model C4.5 Dan Naive Bayes. *Jurnal Algor*, 4(2), 1–10. <https://jurnal.buddhidharma.ac.id/index.php/algor/index>
- Hariyanto, M. (n.d.). Krakatau Bandar Samudera perkuat jaringan maritim nasional di IMW 2025. *Antara News*. Retrieved <https://www.antaranews.com/berita/4859873/krakatau-bandar-samudera-perkuat-jaringan-maritim-nasional-di-imw-2025>
- Muhaimin, A., Amin Hariyadi, M., & Imamudin, M. I. (2024). Klasifikasi Prestasi Akademik Siswa Berdasarkan Nilai Rapor dan Kedisiplinan dengan Metode K-Nearest Neighbor. *Jurnal Ilmu Komputer Dan Sistem Informasi (JIKOMSI)*, 7(1), 193–202. <https://doi.org/10.55338/jikoms.v7i1.2865>
- Novitadewi, N. M. A., Sugiartawan, P., & Fittryani, Y. P. (2023). Klasifikasi Data Penjualan Dengan Metode K-Nearest Neighbor Pada Pt. Terang Abadi Raya. *Jurnal Sistem Informasi Dan Komputer Terapan Indonesia (JSIKTI)*, 5(1), 11–20. <https://doi.org/10.33173/jsikti.173>

- Nugroho, A., & Religia, Y. (2021). Analisis Optimasi Algoritma Klasifikasi Naive Bayes menggunakan Genetic Algorithm dan Bagging. *Jurnal RESTI*, 5(3), 504–510. <https://doi.org/10.29207/resti.v5i3.3067>
- Nuraeni, F., Kurniadi, D., & Diazki, M. H. (2024). Algoritma K-Nearest Neighbor pada Kasus Dataset Imbalanced untuk Klasifikasi Kinerja Karyawan Perusahaan. *Jurnal Teknologi Informasi Dan Ilmu Komputer*, 11(3), 557–568. <https://doi.org/10.25126/jtiik.938144>
- Puteri, Q. A., Sagirani, T., & Lemantara, J. (2023). Perbandingan Algoritma Naïve Bayes dan K-Nearest Neighbor (KNN) untuk Mengetahui Keakuratan Diagnosa Penyakit Diabetes. *Jurnal Nasional Teknologi Dan Sistem Informasi*, 9(3), 247–254. <https://doi.org/10.25077/teknosi.v9i3.2023.247-254>
- Rahmat, A., Auliasari, K., & Pranoto, Y. A. (2020). IMPLEMENTASI METODE K-NEAREST NEIGHBOR (KNN) UNTUK SELEKSI CALON KARYAWAN BARU (Studi Kasus : BFI Finance Surabaya). In *Jurnal Mahasiswa Teknik Informatika* 4(2).
- Roihan, A., Sunarya, P. A., & Rafika, A. S. (2020). Pemanfaatan Machine Learning dalam Berbagai Bidang: Review paper. *IJCIT (Indonesian Journal on Computer and Information Technology)*, 5(1), 75–82. <https://doi.org/10.31294/ijcit.v5i1.7951>
- Wilujeng, D. T., Fatekurohman, M., & Tirta, I. M. (2023). Analisis Risiko Kredit Perbankan Menggunakan Algoritma K-Nearest Neighbor dan Nearest Weighted K-Nearest Neighbor. *Indonesian Journal of Applied Statistics*, 5(2), 142. <https://doi.org/10.13057/ijas.v5i2.58426>