

BAB II

TINJAUAN PUSTAKA

2.1. Landasan Teori

2.1.1 Waktu Tunggu

Waktu tunggu konsumen terhadap layanan jasa adalah titik awal dari interaksi antara konsumen dan penyedia layanan., waktu dianggap sebagai sumber daya yang harus dikelola dengan hati-hati. Semakin berharga waktu bagi konsumen, semakin negatif persepsinya jika dirasakan sebagai pemborosan (Julyanthry dkk., 2020).

Waktu tunggu memiliki empat aspek: obyektif, subyektif, kognitif, dan afektif. Waktu tunggu obyektif adalah waktu yang diukur dengan stopwatch sebelum layanan diberikan. Waktu tunggu subyektif adalah estimasi konsumen terhadap lamanya menunggu, yang bisa bergantung pada waktu yang diukur secara obyektif. Aspek kognitif waktu tunggu adalah penilaian konsumen terhadap lamanya menunggu, apakah masih dapat diterima atau tidak. Aspek afektif mencakup respons emosional seperti frustrasi atau kegembiraan saat menunggu (Kastella, 2019).

Selama puncak permintaan, perusahaan dapat meningkatkan kapasitasnya dengan menambah sumber daya. Saat permintaan menurun, mereka bisa menggoda konsumen dengan penawaran menarik. Reservasi juga bisa membantu meratakan beban permintaan, meskipun tidak bisa sepenuhnya mencegah penundaan. Waktu tunggu yang terlalu lama dapat mengakibatkan ketidakpuasan konsumen terhadap layanan (Fatihudin & Firmansyah, 2019).

2.1.2 Service Motor

Setiap sepeda motor pada akhirnya akan mengalami kondisi di mana berbagai bagian seperti mesin, transmisi, rangka, dan lainnya akan mengalami kelelahan dan keausan, menyebabkan penurunan kinerja seperti tenaga mesin yang menurun, akselerasi yang lambat, konsumsi bahan bakar yang boros, dan kemungkinan kerusakan yang dapat merembet ke komponen lainnya. Jika tidak ditangani melalui perawatan berkala, kondisi ini akan semakin memburuk dan membutuhkan biaya yang signifikan untuk memulihkan sepeda motor ke kondisi semula (Qiram, 2017). Menurut Hartanto, (2023), Servis motor adalah kegiatan perawatan berkala pada sepeda motor yang melibatkan:

1. Memeriksa setiap bagian sepeda motor untuk memastikan fungsinya.
2. Membersihkan bagian yang kotor untuk mencegah kerusakan sistem.
3. Menyetel bagian yang berubah agar sesuai dengan spesifikasi.
4. Memperbaiki atau mengganti komponen yang rusak atau aus.

Rangkaian kegiatan servis sepeda motor meliputi:

1. Bagian Mesin:
 - a. Memeriksa dan merawat baterai.
 - b. Mengganti oli pelumas mesin.
 - c. Membersihkan saringan udara.
 - d. Membersihkan saringan bahan bakar.
 - e. Memeriksa dan menyetel busi.
 - f. Membersihkan karburator.

- g. Menyetel katup.
- h. Menyetel campuran bahan bakar/putaran mesin.
- i. Menyetel kebebasan kopling.
- 2. Bagian Kelistrikan:
 - a. Memeriksa dan merawat baterai.
 - b. Memeriksa fungsi kelistrikan seperti lampu dan indikator.
- 3. Bagian Chasis:
 - a. Memeriksa dan menyetel gerak bebas rem.
 - b. Memeriksa dan merawat rantai roda.
 - c. Memeriksa poros kemudi.
 - d. Memeriksa kondisi ban dan tekanan angin.
 - e. Memeriksa dan mengencangkan baut-baut pengikat.

2.1.3 Bengkel Motor

Sebuah bengkel adalah lokasi di mana seorang teknisi melakukan pekerjaannya untuk memberikan layanan perbaikan dan perawatan kendaraan. Bengkel umum kendaraan bermotor adalah tempat yang berfungsi untuk memperbaiki dan merawat kendaraan bermotor agar tetap memenuhi persyaratan teknis dan dapat digunakan dengan aman di jalan. Ini sesuai dengan tuntutan yang diatur dalam Peraturan Pemerintah Nomor 44 Tahun 1993 tentang Kendaraan dan Pengemudi, khususnya pasal 126, 127, 128, dan 129, yang menegaskan bahwa setiap kendaraan bermotor harus memenuhi standar teknis dan kelayakan yang ditetapkan (Sofi & Dharmawan, 2022).

Bengkel merupakan usaha wirausaha kecil dan menengah yang bergerak dalam

bidang jasa perbaikan, baik untuk sepeda motor maupun mobil. Bengkel sepeda motor, khususnya, adalah usaha yang bertujuan untuk memperbaiki sepeda motor agar kembali berfungsi dengan baik sesuai keinginan pemilik atau spesifikasi aslinya (Donggeari dkk., 2023).

2.1.4 Data Mining

Data mining adalah proses eksplorasi dan pengungkapan pola, hubungan, dan informasi yang tersembunyi dalam kumpulan data besar menggunakan teknik dan algoritma terkait statistik dan sistem basis data. Ini merupakan tahap kunci dalam proses "penambangan data dari *database*" (KDD), yang mencakup serangkaian langkah seperti analisis data, manajemen data, pra-pemrosesan data, pembuatan model dan inferensi, pengukuran kepentingan, pengelompokan, pemrosesan hasil temuan, visualisasi, dan pembaruan secara online (Putra dkk., 2023).

Tujuan utama dari data mining adalah menggali dan mengumpulkan informasi yang berharga dari berbagai sumber data, termasuk data mart dan data warehouse, untuk keperluan analisis. Proses ini memanfaatkan berbagai algoritma dan teknik, seperti analisis asosiasi, model jaringan saraf tiruan, dan pengklasifikasi, untuk mengidentifikasi pola-pola dan anomali yang relevan dalam volume data yang besar (Jollyta dkk., 2020).

Data mining memiliki banyak aplikasi dalam dunia bisnis, termasuk optimisasi pasar, klasifikasi email, dan deteksi anomali dalam data. Meskipun proses data mining dapat bervariasi tergantung pada proyek dan teknik yang digunakan, langkah-langkah umumnya mencakup penetapan tujuan, pengumpulan data, penerapan algoritma data *mining*, dan evaluasi hasilnya (Urva dkk., 2023).

2.1.5 K-Nearest Neighbors

Metode *K-Nearest Neighbor* adalah metode non-parametrik yang dapat digunakan untuk klasifikasi berdasarkan keterdekatan dengan tetangga terdekat dan juga untuk regresi. Algoritma *K-Nearest Neighbor* bekerja dengan mencari jarak antara data yang akan dievaluasi (data latih) dan sekumpulan tetangga terdekat k (*neighbor*) dalam data baru (data uji). Saat data baru yang tidak diketahui masuk untuk diklasifikasikan, kategori data baru tersebut ditentukan berdasarkan mayoritas kategori dari tetangga terdekat k dalam data uji. Karakteristik data yang akan diklasifikasikan diekstraksi dan dibandingkan dengan karakteristik setiap data kategori yang sudah dikenal dalam data uji, lalu tetangga terdekat k diambil untuk menghitung mayoritas kategori data (Liantoni, 2015).

Penentuan nilai k dalam algoritma *K-Nearest Neighbor* dapat dilakukan dengan mencari nilai k terdekat $k_1, k_2, \dots, k_s$. Semakin banyak data yang ada, semakin kecil nilai k yang dipilih, namun jika dimensi data lebih besar, nilai k yang dipilih harus lebih tinggi. Lebih baik menggunakan angka ganjil untuk nilai k seperti $k = 1, 3, 5$, dst. Nilai k harus memenuhi syarat $k < N$, dimana N adalah jumlah dataset latih, karena nilai k digunakan untuk mencari mayoritas kelas/label pada data latih dan tidak boleh melebihi jumlah dataset latih. Terdapat berbagai cara untuk mencari tetangga terdekat atau jarak dalam algoritma K-NN, termasuk Jarak *Euclidean*, Jarak *Manhattan*, Jarak *Cosine*, Jarak Korelasi, dan Jarak *Hamming*. Jarak antara dua tetangga terdekat berdasarkan nilai kemiripan dapat dihitung menggunakan jarak *Euclidean* (Rachma, 2022).

$$Dist(X, Y) = \sqrt{\sum_{i=1}^D (X_i - Y_i)^2}$$

..... (2.1)

Keterangan:

$Dist(X, Y)$: jarak antar objek (Euclidean Distancing)

X_i : sampel data

Y_i : data uji

D : dimensi data

i : variabel data

2.1.6 Confusion Matrix

Confusion Matrix adalah alat evaluasi yang penting dalam pembelajaran mesin, terutama dalam masalah klasifikasi. Ini adalah tabel yang digunakan untuk mengevaluasi kinerja suatu model klasifikasi dengan membandingkan prediksi model dengan nilai sebenarnya dari data uji (Purwanto & Nugroho, 2023).

Confusion Matrix terdiri dari empat sel utama:

1. *True Positive* (TP): Ini adalah jumlah kasus di mana model dengan benar memprediksi kelas positif.
2. *True Negative* (TN): Ini adalah jumlah kasus di mana model dengan benar memprediksi kelas negatif.
3. *False Positive* (FP): Ini adalah jumlah kasus di mana model salah memprediksi kelas positif ketika sebenarnya kelasnya negatif. Juga dikenal sebagai kesalahan tipe I.
4. *False Negative* (FN): Ini adalah jumlah kasus di mana model salah memprediksi kelas negatif ketika sebenarnya kelasnya positif. Juga dikenal sebagai kesalahan tipe II.

Dari nilai-nilai di atas, berbagai metrik evaluasi dapat dihitung, seperti akurasi, presisi, *recall* (sensitivitas), dan *F1-score*. Berikut adalah rumus untuk menghitung

parameter yang ada di *Confusion Matrix*:

1. Accuracy

Akurasi adalah metode pengujian berdasarkan seberapa dekat nilai prediksi dengan nilai sebenarnya atau aslinya. Dengan mengetahui jumlah data yang terklasifikasi dengan benar, maka akurasinya dapat diketahui. Untuk menghitung nilai akurasi, Anda dapat menggunakan rumus berikut:

$$\text{Accuracy} = \frac{TP+TN}{\text{Total}}$$

..... (2.2)

Keterangan:

- a. *True Positive* :(TP) jumlah prediksi benar untuk kelas *positive*.
- b. *True Negative* :(TN) jumlah prediksi salah untuk kelas *positive*.
- c. Total: jumlah semua data aktual

2. Precision

Presisi adalah metode pengujian yang membandingkan jumlah informasi yang relevan yang diperoleh sistem dengan jumlah semua informasi yang diambil oleh sistem, terlepas dari apakah itu terkait atau tidak. Untuk menghitung nilai presisi, Anda dapat menggunakan rumus berikut:

$$\text{precision} = \frac{TP}{TP+FP}$$

.....(2.3)

Keterangan:

- a. *True Positive* (TP): jumlah prediksi benar untuk kelas *positive*.
- b. *False Positive* (FP) jumlah prediksi salah untuk kelas *positive*.

3. *Recall*

Ini memberikan gambaran yang komprehensif tentang kinerja model dalam mengklasifikasikan data. *Confusion Matrix* sangat berguna dalam memahami di mana model cenderung melakukan kesalahan dan dapat membantu dalam menyesuaikan model atau menilai kelayakan model untuk keperluan tertentu.

2.1.7 **Google Colab**

Google Colab (singkatan dari *Collaboratory*) adalah layanan *Cloud Computing* yang disediakan oleh *Google*. Ini memungkinkan pengguna untuk menulis dan mengeksekusi kode *Python* di *browser* web, menggunakan sumber daya komputasi *Google* secara gratis. *Colab* didukung oleh *Google Cloud Platform* dan terutama digunakan untuk pengembangan dan penelitian di bidang pembelajaran mesin, ilmu data, dan pengembangan perangkat lunak (Pane & Saputra, 2020). Fitur utama dari *Google Colab* meliputi:

1. *Python Environment Colab* menyediakan lingkungan *Python* yang lengkap dengan akses ke berbagai pustaka populer seperti *TensorFlow*, *PyTorch*, dan *scikit-learn*.
2. Interaktivitas Pengguna dapat menulis dan mengeksekusi kode *Python* secara interaktif dalam sel-sel yang dapat diedit. Mereka juga dapat menambahkan teks, gambar, grafik, dan catatan menggunakan *Markdown*.
3. Kolaborasi *Colab* memungkinkan beberapa pengguna untuk bekerja secara bersama-sama pada notebook yang sama secara real-time. Ini memfasilitasi kolaborasi tim dalam proyek-proyek data science dan pembelajaran mesin.
4. Sumber Daya Komputasi *Colab* menyediakan akses ke sumber daya komputasi

Google, termasuk CPU, GPU, dan TPU. Pengguna dapat memilih untuk menggunakan sumber daya ini secara gratis atau memilih opsi pembayaran untuk sumber daya yang lebih kuat.

5. Penyimpanan dan Integrasi *Notebook Colab* dapat disimpan di *Google Drive* pengguna dan diimpor atau diekspor dalam format Jupyter. Ini memudahkan untuk berbagi dan menyimpan proyek-proyek Colab. *Google Colab* telah menjadi alat populer di komunitas ilmu data dan pembelajaran mesin karena kemudahannya dalam penggunaan, akses ke sumber daya komputasi, dan fitur kolaboratifnya.

2.2. Kajian Penelitian

Penelitian sebelumnya telah memanfaatkan beberapa pendekatan dalam melakukan klasifikasi menggunakan metode *K-Nearest Neighbors* (KNN), menguji hasilnya dengan confusion matrix, dan mengkaji masalah waktu tunggu. Dalam konteks penelitian ini, peneliti akan memperluas pemahaman yang relevan dengan judul penelitian sebelumnya yang membahas penerapan metode *K-Nearest Neighbors* dan confusion matrix. Berikut adalah rangkuman dari beberapa penelitian sebelumnya yang diidentifikasi oleh peneliti sebagai relevan untuk penelitian mereka saat ini.

Tabel 2.1 Rujukan Jurnal

No	Tahun	Penulis	Judul	Jurnal
1	2021	Cholil, Saifur Rohman Handayani, Titis Prathivi, Rastri Ardianita, Tria	Implementasi Algoritma Klasifikasi <i>K-Nearest Neighbor</i> (KNN) Untuk Klasifikasi Seleksi Penerima Beasiswa	IJCIT (Indonesian Journal on Computer and Information Technology)

No	Tahun	Penulis	Judul	Jurnal
2	2020	Hasdyna, Novia Dinata, Rozzi Kesuma	<i>Analisis Matthew Correlation Coefficient</i> pada K-Nearest Neighbor dalam Klasifikasi Ikan Hias	INFORMAL: Informatics Journal
3	2020	Yudhana, Anton Sunardi, Sunardi Hartanta, Agus Jaka Sri	Algoritma K-Nn Dengan Untuk Prediksi Hasil Penggergajian Kayu Sengon	Transmisi
4	2018	Prasetio, Rizki Tri Rismayadi, Ali Akbar Anshori, Iedam Fardian	Implementasi Algoritma Genetika pada k-nearest neighbours untuk Klasifikasi Kerusakan Tulang Belakang	Jurnal Informatika
5	2021	Rezkika, Febbyola Betha Nurina Irawan, Agung Yuda Susilo	KLASIFIKASI MASA TUNGGU ALUMNI UNTUK MENDAPATKAN PEKERJAAN MENGGUNAKAN ALGORITMA C4.5	Progresif: Jurnal Ilmiah Komputer
6	2020	Kurniawan, Yogiek Barokah, Tiyyssa Indah	Klasifikasi Penentuan Pengajuan Kartu Kredit Menggunakan K-Nearest Neighbor	Jurnal Ilmiah Matrik
7	2023	Anggraeni, Novia Dwi Octaviano, Alvino	Penerapan Data Mining Untuk Klasifikasi Persediaan Barang Prediksi Persediaan Barang Menggunakan Metode Safety Stock & ROP (Studi Kasus : PT . Macro Jaya Agung)	OKTAL : Jurnal Ilmu Komputer dan Science
8	2015	Megawati Hakim, Lukman Irbantoro, Dolly	Penurunan Waktu Tunggu Pelayanan Obat Rawat Jalan Instalasi Farmasi Rumah Sakit Baptis Batu	Jurnal Kedokteran Brawijaya
9	2020	Zhang, Shichao	Cost-sensitive KNN classification	Neurocomputing
10	2021	Isnain, Auliya Rahman Supriyanto,	Implementation of K- Nearest Neighbor (K-NN) Algorithm For Public	IJCCS (Indonesian Journal of Computing and Cybernetics Systems)

Penelitian pertama yang dilakukan oleh Cholil dkk., (2021) yang berjudul “Implementasi Algoritma Klasifikasi *K-Nearest Neighbor* (KNN) Untuk Klasifikasi Seleksi Penerima Beasiswa”. Tujuan dari penelitian ini adalah untuk memperbaiki proses seleksi penerimaan beasiswa di SMA melalui penerapan algoritma *K-Nearest Neighbor* (KNN), sehingga memastikan bahwa penerima beasiswa adalah orang yang tepat sasaran. Hasil dari penelitian ini adalah terseleksinya 30 orang dari 89 data yang telah dilakukan klasifikasi menggunakan algoritma KNN. Pengujian sistem menunjukkan bahwa algoritma KNN memiliki tingkat akurasi sebesar 90.5%, yang menunjukkan kemampuannya dalam mengklasifikasikan seleksi penerimaan beasiswa dengan baik. Dengan demikian, hasil penelitian ini membuktikan bahwa algoritma KNN dapat digunakan secara efektif untuk memilih penerima beasiswa yang paling pantas berdasarkan data yang tersedia.

Penelitian kedua yang berjudul “*Analisis Matthew Correlation Coefficient* pada *K-Nearest Neighbor* dalam Klasifikasi Ikan Hias” yang ditulis oleh Hasdyna & Dinata, (2020). Tujuan dari penelitian ini adalah untuk mengukur kinerja algoritma *K-Nearest Neighbor* (K-NN) dengan menggunakan *Matthew Correlation Coefficient* (MCC). Data yang digunakan dalam penelitian ini adalah data ikan hias yang terdiri dari 3 kelas yaitu *Premium*, *Medium*, dan *Low*. Hasil analisis koefisien korelasi *Matthew* pada K-NN dengan menggunakan didapatkan nilai MCC tertinggi pada kelas *Medium* sebesar 0.786542. Nilai MCC tertinggi kedua terdapat pada kelas *Premium* sebesar 0.567434. Sedangkan nilai MCC terendah terdapat pada kelas *Low* sebesar 0.435269. Secara keseluruhan, nilai MCC adalah secara statistik sebesar 0.596415. Ini

menunjukkan bahwa algoritma K-NN mampu mengklasifikasikan data dengan baik, meskipun terdapat variasi dalam kinerja klasifikasi antara kelas *Premium*, *Medium*, dan *Low*.

Penelitian ketiga yang berjudul “Algoritma K-Nn Dengan Untuk Prediksi Hasil Penggajian Kayu Sengon” oleh Yudhana dkk., (2020). Penelitian ini bertujuan untuk memperbaiki efisiensi proses penggajian kayu Sengon dengan menerapkan algoritma *K-Nearest Neighbor* (K-NN) menggunakan metode *Euclidean Distance*. Awalnya, perhitungan manual dilakukan dengan *Ms Excel*, namun kemudian dikembangkan aplikasi berbasis *web* dengan PHP, *MySQL*, dan framework *Laravel*. Data training dan testing digunakan untuk menguji aplikasi, dengan variasi nilai *k* (1-5). Hasil penelitian menunjukkan bahwa aplikasi dapat memprediksi hasil penggajian dengan akurasi sebesar 70%, meskipun terdapat tingkat kesalahan sebesar 30% jika dibandingkan dengan hasil riil di lapangan.

Penelitian keempat yang berjudul “Implementasi Algoritma Genetika pada *K-Nearest Neighbours* untuk Klasifikasi Kerusakan Tulang Belakang” oleh Prasetyo dkk., (2018). Penelitian ini bertujuan untuk meningkatkan akurasi klasifikasi gangguan tulang belakang dengan mengimplementasikan algoritma genetika pada algoritma *K-Nearest Neighbors* (KNN). Penggunaan *Computer Aided Diagnosis* (CAD) System diusulkan untuk membantu radiologis dalam mendeteksi kelainan pada tulang belakang secara lebih optimal. Dataset terdiri dari tiga kelas klasifikasi penyakit tulang belakang, yaitu *herniated disk*, *spondylolisthesis*, dan kelas normal, diambil berdasarkan ekstraksi citra MRI. Penelitian dilakukan melalui lima eksperimen dengan variasi pembagian data training dan data testing menggunakan *split validation*. Metode yang diusulkan menggabungkan algoritma genetika untuk

seleksi fitur dan optimasi parameter algoritma *k-nearest neighbors*.

Hasil penelitian menunjukkan peningkatan signifikan dalam akurasi klasifikasi, dengan rata-rata akurasi sebesar 93% dari lima eksperimen. Ini merupakan peningkatan yang signifikan dibandingkan dengan penggunaan algoritma *K-Nearest Neighbors* saja, yang hanya menghasilkan rata-rata akurasi sebesar 82.54%. Dengan demikian, penelitian ini menunjukkan bahwa penggunaan algoritma genetika pada algoritma *K-Nearest Neighbors* dapat meningkatkan akurasi dalam klasifikasi gangguan tulang belakang, berpotensi memberikan kontribusi penting dalam bidang diagnostik medis.

Penelitian kelima yang berjudul “Klasifikasi Masa Tunggu Alumni Untuk Mendapatkan Pekerjaan Berdasarkan Kompetensi Menggunakan Algoritma C4.5 (Studi Kasus : Fasilkom Unsika)” oleh Rezkika dkk., (2021). Tujuan dari penelitian ini adalah untuk memprediksi waktu tunggu alumni dalam mendapatkan pekerjaan menggunakan algoritma decision tree C4.5, serta membandingkannya dengan fitur forward selection. Pelacakan status pekerjaan alumni merupakan indikator penting untuk mengevaluasi kualitas lulusan perguruan tinggi. Teknologi informasi, seperti data mining, digunakan sebagai alat untuk menghasilkan pengetahuan yang diperlukan dalam pelacakan tersebut. Data yang diproses menggunakan algoritma C4.5 dengan bantuan perangkat lunak *RapidMiner Studio*.

Hasil penelitian menunjukkan bahwa algoritma decision tree C4.5 dengan fitur *forward selection* memberikan performa terbaik, dengan akurasi 80,37%, presisi 79,56%, *recall* 81,34%, f-measure 80,40%, dan AUC 0,914 yang masuk dalam kategori klasifikasi yang sangat baik. Dengan demikian, algoritma C4.5 berbasis

Forward Selection terbukti meningkatkan tingkat akurasi secara signifikan, dibandingkan dengan algoritma decision tree C4.5 tunggal. Terdapat peningkatan akurasi sebesar 25,93%, menunjukkan bahwa algoritma ini dapat menjadi metode yang efektif dalam memprediksi waktu tunggu alumni dalam mendapatkan pekerjaan.

Penelitian keenam yang berjudul “Klasifikasi Penentuan Pengajuan Kartu Kredit Menggunakan K-Nearest Neighbor” oleh Kurniawan & Barokah, (2020). Tujuan dari penelitian ini adalah untuk memudahkan bank atau analis dalam menentukan kategori kartu kredit yang sesuai untuk nasabah bank. Masalah yang dihadapi adalah kesulitan dalam menentukan kategori kartu kredit yang sesuai dengan nasabah bank yang telah mendaftar. Penelitian ini menggunakan metode K-Nearest Neighbor untuk mengklasifikasikan calon nasabah dalam pengajuan kartu kredit sesuai dengan kategori nasabah, dengan menggunakan data nasabah di Bank BNI Syariah Surabaya. Metode K-Nearest Neighbor digunakan untuk mencari pola pada data nasabah sehingga variabel yang ditetapkan sebagai faktor pendukung dalam bentuk jenis kelamin, status rumah, status, jumlah tanggungan (anak), profesi, dan pendapatan tahunan dapat ditentukan. Hasil dari penelitian ini menunjukkan bahwa rata-rata nilai presisi sebesar 92%, nilai *recall* sebesar 83%, dan nilai akurasi sebesar 93%.

Penelitian ketujuh yang berjudul “Penerapan Data Mining Untuk Klasifikasi Persediaan Barang Prediksi Persediaan Barang Menggunakan Metode *Safety Stock* & ROP (Studi Kasus: PT. Macro Jaya Agung)” oleh Anggraeni & Octaviano, (2023). Penelitian ini bertujuan untuk membantu PT. Macro Jaya Agung dalam

memprediksi minat pelanggan dan mengelola persediaan barang dagang secara efektif. Dengan menggunakan algoritma *K-Nearest Neighbor* dan C4.5 serta metode perhitungan *safety stock* dan ROP, hasil penelitian menunjukkan tingkat akurasi yang cukup baik, yaitu 88.89% untuk *K-Nearest Neighbor* dan 66.67% untuk C4.5 menggunakan confusion matrix, serta 86.00% dan 71.81% menggunakan cross validation. Hasil penelitian menyarankan *safety stock* sebanyak 131,889 pcs dan ROP sebanyak 186 pcs untuk tabung besi 2.2L (0.25 M3). Dengan demikian, penelitian ini membantu PT. Macro Jaya Agung dalam mengelola persediaan barang dagang dengan lebih efektif, mengurangi penumpukan barang, dan menjaga arus keuangan perusahaan agar berjalan dengan lancar.

Penelitian kedelapan yang berjudul “Penurunan Waktu Tunggu Pelayanan Obat Rawat Jalan Instalasi Farmasi Rumah Sakit Baptis Batu” oleh Megawati dkk., (2015). Tujuan penelitian ini adalah untuk mengevaluasi efek intervensi terhadap waktu tunggu dalam layanan obat rawat jalan serta kepuasan pengunjung. Metode penelitian yang digunakan adalah pre dan post tes dengan pendekatan studi eksperimental pada 119 sampel pengunjung selama 2 bulan. Data dikumpulkan dengan mengukur waktu tunggu sebelum dan sesudah intervensi, serta kepuasan pengunjung menggunakan skala Likert. Analisis dilakukan dengan menggunakan uji T berpasangan dan Wilcoxon untuk membandingkan waktu tunggu dan kepuasan pelanggan sebelum dan sesudah intervensi. Hasil penelitian menunjukkan perbaikan yang signifikan dalam waktu tunggu, dengan rerata waktu tunggu untuk pelayanan obat jadi dan obat racikan yang menurun secara signifikan ($p < 0,05$).

Perbaikan waktu tunggu juga terlihat berdasarkan jumlah obat dalam satu resep. Selain itu, kepuasan pelanggan terhadap waktu tunggu juga meningkat secara signifikan setelah intervensi ($p < 0,05$). Dengan demikian, hasil penelitian ini menunjukkan bahwa intervensi yang dilakukan berhasil memperbaiki waktu tunggu dalam layanan obat rawat jalan serta meningkatkan kepuasan pengunjung.

Penelitian kesembilan yang berjudul “*Cost-sensitive KNN classification*” oleh Zhang, (2020). Penelitian ini bertujuan untuk mengembangkan metode klasifikasi KNN yang sensitif terhadap biaya untuk menangani masalah distribusi kelas yang tidak seimbang dalam data klinis dan mengurangi biaya kesalahan klasifikasi. Dua *classifier* baru, yaitu *Direct-CS-KNN* dan *Distance-CS-KNN*, dirancang untuk mengatasi masalah ini dengan efektif. Beberapa strategi peningkatan, seperti smoothing, pemilihan nilai K minimum, dan pemilihan fitur, juga digunakan bersama dengan KNN *classifier*. Eksperimen dilakukan dengan menggunakan dataset dunia nyata dan menunjukkan bahwa classifier KNN yang sensitif terhadap biaya secara signifikan mengurangi biaya kesalahan klasifikasi dibandingkan dengan metode yang sudah ada, bahkan melebihi kinerja CS-C4.5 pada beberapa dataset dari UCI. Dalam penelitian ini, kami menguraikan konsep pembelajaran sensitif terhadap biaya, mengusulkan dua classifier baru, dan menunjukkan efisiensi serta keunggulan mereka dalam eksperimen dengan dataset dunia nyata.

Penelitian kesepuluh yang berjudul “*Implementation of K-Nearest Neighbor (K-NN) Algorithm For Public Sentiment Analysis of Online Learning*” oleh Isnain dkk., (2021). Penelitian ini bertujuan untuk menerapkan algoritma KNN (K-Nearest Neighbor) dalam melakukan analisis sentimen pengguna Twitter terkait kebijakan

pemerintah tentang Pembelajaran *Online*. Data Tweet sebanyak 1825 data tweet Indonesia dikumpulkan dari 1 Februari 2020 hingga 30 September 2020 menggunakan perpustakaan *Python*, *Tweepy*. Pembobotan kata menggunakan TF-IDF, akan diklasifikasikan ke dalam dua kelas nilai sentimen, positif dan negatif. Setelah pengujian dengan K sebesar 20, hasil

akurasi tertinggi diperoleh ketika $K = 10$ dengan nilai akurasi sebesar 84,65%, dengan presisi 87%, *recall* 86%, f- measure 87%, dan tingkat kesalahan 0,12%. Diperoleh juga kecenderungan opini publik terhadap pembelajaran online cenderung positif.

Berdasarkan hasil beberapa penelitian yang dijadikan acuan dalam kajian pustaka dan teori pada penelitian ini, dapat disimpulkan bahwa metode *k-nearest neighbors* dapat digunakan sebagai salah satu referensi dalam proses pengambilan keputusan untuk menentukan waktu tunggu. Hal ini didasarkan pada tingkat akurasi dan presisi yang baik, seperti yang terbukti dari hasil-hasil penelitian sebelumnya. Penelitian yang dilakukan oleh Anggraeni & Octaviano, (2023) dengan judul "Penerapan Data *Mining* untuk Klasifikasi Persediaan Barang: Prediksi Persediaan Barang Menggunakan Metode Safety Stock & ROP (Studi Kasus: PT. Macro Jaya Agung)" menunjukkan bahwa metode *K-Nearest Neighbors* menghasilkan kinerja yang lebih baik dibandingkan dengan algoritma *Decision Tree* atau pohon keputusan C45.

Selain itu, penelitian yang dilakukan oleh Rezkika dkk., (2021) yang berjudul "Klasifikasi Masa Tunggu Alumni Untuk Mendapatkan Pekerjaan Berdasarkan Kompetensi Menggunakan Algoritma C4.5 (Studi Kasus : Fasilkom Unsika)", juga

menunjukkan bahwa waktu tunggu dapat diprediksi melalui proses klasifikasi. Hasil penelitian terakhir yang dilakukan oleh Megawati dkk., (2015) dari Universitas Brawijaya tentang "Penurunan Waktu Tunggu Pelayanan Obat Rawat Jalan Instalasi Farmasi Rumah Sakit Baptis Batu" juga menghasilkan hasil yang baik, yang berkontribusi pada peningkatan pelayanan jasa.

Oleh karena itu, judul skripsi ini menggarisbawahi pendekatan penelitian yang belum pernah dieksplorasi sebelumnya dalam konteks waktu tunggu dalam proses layanan servis motor. Ini diambil dengan merujuk pada studi sebelumnya yang memusatkan pada klasifikasi waktu tunggu. Judul skripsi ini adalah "Implementasi *K-Nearest Neighbors* untuk Prediksi Waktu Tunggu Service Motor di *Develop Tech*", dengan tujuan memanfaatkan teknik *K-Nearest Neighbors* untuk memprediksi waktu tunggu dalam layanan motor.