

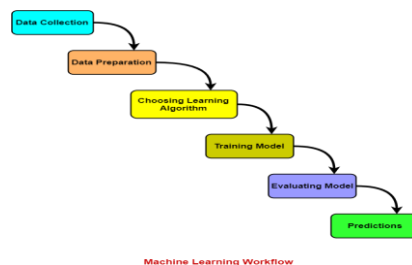
BAB II

LANDASAN TEORI

2.1 Machine Learning

Penerapan *machine learning* dipahami sebagai rangkaian proses analitis yang disusun secara sistematis untuk mengubah data mentah menjadi informasi yang bermakna. Proses ini dimulai dari pemahaman terhadap data hingga menghasilkan keluaran berupa model yang dapat digunakan untuk klasifikasi atau prediksi. Literatur metodologi analitik menekankan bahwa kejelasan alur kerja menjadi faktor penentu keberhasilan implementasi *machine learning* karena setiap tahap saling bergantung dan tidak dapat dipisahkan secara parsial [1].

Kerangka proses analitis tersebut relevan dengan penelitian inventory apotek yang memanfaatkan data historis penjualan obat. Data penjualan tidak serta-merta dapat digunakan sebagai masukan algoritma tanpa melalui tahapan pengolahan yang tepat. Setiap keputusan metodologis pada tahap awal akan memengaruhi kualitas hasil klasifikasi pada tahap akhir. Oleh karena itu, penerapan *machine learning* menuntut pendekatan yang runtut, terencana, dan berbasis kaidah ilmiah.



Gambar 2.1 Alur Umum Penerapan Machine Learning Secara End-to-End

Sumber : https://www.gatevidyalay.com/machine-learning-workflow-process-steps/?utm_source

Gambar ini menampilkan alur umum penerapan machine learning dari tahap perumusan masalah hingga pemakaian model. Proses dimulai saat kebutuhan bisnis atau pertanyaan penelitian diterjemahkan menjadi tujuan dan metrik keberhasilan yang jelas. Berikutnya data dikumpulkan dari sumber relevan, lalu dipersiapkan melalui pembersihan, integrasi, penghilangan duplikasi, serta penanganan nilai hilang dan outlier agar kualitas data memadai. Setelah data rapi, dilakukan transformasi dan/atau rekayasa fitur, kemudian dataset dibagi menjadi train–validasi–uji untuk mencegah kebocoran informasi. Tahap selanjutnya adalah memilih algoritma dan melakukan pelatihan serta penyetelan hiperparameter, biasanya secara iteratif. Kinerja dievaluasi memakai metrik yang sesuai (misalnya akurasi, F1, MAE) dan analisis error untuk memastikan model tidak overfitting dan tetap adil. Model terpilih kemudian dideploy ke aplikasi atau pipeline, disertai monitoring performa, drift data, dan loop umpan balik untuk retraining saat kondisi berubah. Sepanjang alur, dokumentasi eksperimen, versioning data-model, dan pengujian reproducibility memastikan integritas ilmiah serta memudahkan audit, pemeliharaan, dan peningkatan berkelanjutan di lingkungan produksi.

2.1.1 Pengumpulan Data

Pengumpulan data merupakan tahap awal yang menentukan keberhasilan seluruh proses *machine learning*. Data berperan sebagai representasi fenomena nyata yang ingin dipelajari oleh model. Kualitas, kelengkapan, dan relevansi data menjadi faktor krusial karena algoritma hanya dapat belajar dari informasi yang tersedia. Studi pengembangan model prediksi berbasis data kesehatan menunjukkan bahwa pemahaman mendalam terhadap karakteristik data sangat penting untuk menghindari kesalahan interpretasi pada tahap lanjutan [2].

Pemahaman data mencakup identifikasi variabel, tipe data, serta keterkaitan antarvariabel dalam konteks permasalahan penelitian. Pada penelitian inventory apotek, data penjualan obat mencerminkan pola permintaan yang terbentuk dari aktivitas transaksi harian. Pemahaman ini memungkinkan peneliti menentukan variabel yang relevan sebagai masukan model klasifikasi kebutuhan stok obat.

2.1.2 Data Preprocessing dan Transformasi Data

Data preprocessing bertujuan menyiapkan data agar layak digunakan oleh algoritma *machine learning*. Proses ini mencakup pembersihan data dari nilai kosong, duplikasi, dan inkonsistensi yang berpotensi menurunkan kualitas model. Literatur *machine learning* secara konsisten menegaskan bahwa kesalahan pada tahap *preprocessing* dapat menghasilkan model yang bias dan tidak stabil [3].

Transformasi data dilakukan untuk menyesuaikan format dan skala variabel sehingga algoritma dapat memproses data secara optimal. Teknik normalisasi sering digunakan untuk menyetarakan skala variabel agar tidak mendominasi proses pembelajaran model. Studi empiris menunjukkan bahwa normalisasi berkontribusi terhadap peningkatan kinerja algoritma klasifikasi berbasis jarak, termasuk *Support Vector Machine* [4].

2.1.3 Data latih dan data Uji

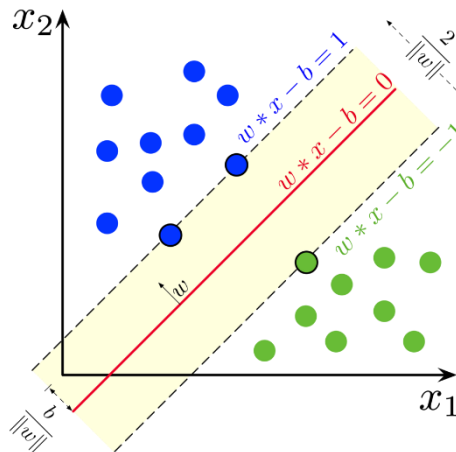
Pembagian data menjadi data latih dan data uji merupakan prosedur standar dalam penerapan *machine learning*. Data latih digunakan untuk membangun model, sedangkan data uji berfungsi menilai kemampuan generalisasi model terhadap data baru. Literatur metodologi *machine learning* menyebutkan bahwa pembagian data merupakan langkah penting untuk mencegah *overfitting* dan memastikan validitas hasil penelitian [5].

Pelatihan model dilakukan dengan mengoptimalkan parameter algoritma berdasarkan data latih. Proses ini memungkinkan model mempelajari pola hubungan antara variabel masukan dan kelas keluaran. Penelitian di berbagai bidang kesehatan menunjukkan bahwa pelatihan model yang terkontrol menghasilkan model prediktif yang lebih andal dan konsisten

2.2 Algoritma Support Vector Machine (SVM)

Algoritma *Support Vector Machine (SVM)* merupakan metode *machine learning* yang dirancang untuk menyelesaikan permasalahan klasifikasi dengan membangun batas keputusan yang optimal antar kelas data. Prinsip utama yang mendasari algoritma ini adalah pencarian *hyperplane* yang mampu memisahkan data dari kelas berbeda dengan margin maksimum. Margin maksimum dipahami sebagai jarak terjauh antara *hyperplane* dan titik data terdekat dari masing-masing kelas, yang dikenal sebagai *support vectors*. Pendekatan ini menjadikan *SVM* memiliki kemampuan generalisasi yang kuat ketika diterapkan pada data baru yang belum pernah dilatih sebelumnya.

Kemampuan generalisasi tersebut menjadi aspek penting dalam konteks prediksi kebutuhan stok obat. Data penjualan obat di apotek umumnya memiliki variasi pola dan jumlah data yang terbatas. Algoritma *SVM* mampu membangun model klasifikasi yang stabil karena tidak seluruh data memengaruhi posisi *hyperplane*, melainkan hanya data-data kritis yang berada di sekitar batas keputusan. Karakteristik ini membedakan *SVM* dari metode klasifikasi lain yang sensitif terhadap perubahan distribusi data.



Gambar 2.2 Ilustrasi Hyperplane dan Margin pada SVM Linear

Sumber : https://en.wikipedia.org/wiki/Support_vector_machine?utm_source

[#/media/File:SVM_margin.png](#)

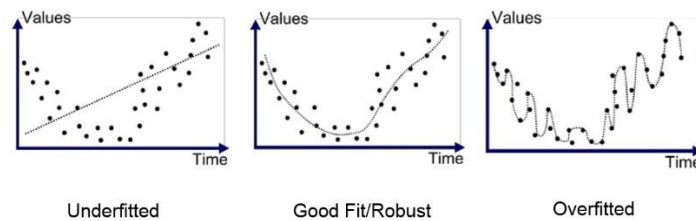
Konsep inti Support Vector Machine (SVM) pada klasifikasi biner: sebuah garis pemisah (*hyperplane*) ditempatkan di ruang fitur untuk memisahkan dua kelompok titik. Dua garis putus-putus paralel di kiri dan kanan merupakan batas margin, yakni jarak terdekat dari *hyperplane* ke sampel tiap kelas. Titik yang menempel pada batas tersebut disebut support vectors; titik-titik inilah yang mengunci posisi *hyperplane* sehingga titik lain tidak banyak memengaruhi keputusan. Tujuan pelatihan SVM adalah memaksimalkan lebar margin karena margin lebih besar umumnya meningkatkan kemampuan generalisasi. Secara geometri, lebar margin berbanding terbalik dengan norma vektor normal w , sehingga optimasi mendorong $\|w\|$ kecil. Bila data tidak sepenuhnya terpisah, soft-margin memperbolehkan pelanggaran batas (slack) dengan penalti C , menyeimbangkan akurasi pelatihan dan ketangguhan pada data baru. Untuk kasus linear-separable, *hyperplane* optimal berada di tengah dua batas margin. Jika pola nonlinier, SVM dapat memakai kernel untuk memetakan fitur ke ruang lebih tinggi tanpa menghitung koordinatnya secara eksplisit lebih efisien.

2.2.1 Fungsi Kernel dalam Support Vector Machine

Data empiris sering kali tidak dapat dipisahkan secara linier pada ruang fitur awal. Algoritma *SVM* mengatasi permasalahan tersebut melalui penggunaan fungsi *kernel* yang memetakan data ke dalam ruang berdimensi lebih tinggi. Transformasi ini memungkinkan pemisahan kelas yang sebelumnya tumpang tindih menjadi lebih jelas tanpa harus menghitung koordinat ruang fitur secara eksplisit. Pendekatan *kernel trick* menjadikan *SVM* fleksibel dan efisien dalam menangani pola data yang kompleks [6].

Fleksibilitas fungsi *kernel* memberikan keunggulan tersendiri bagi *SVM* dalam permasalahan klasifikasi berbasis data historis. Berbagai penelitian menunjukkan bahwa pemilihan fungsi *kernel* yang sesuai berpengaruh terhadap kinerja model klasifikasi. Kemampuan ini relevan dengan konteks inventory apotek, di mana pola permintaan obat tidak selalu mengikuti hubungan linier yang sederhana.

Stabilitas model menjadi salah satu keunggulan utama *SVM*. Algoritma ini dirancang untuk meminimalkan kesalahan klasifikasi sekaligus menjaga kompleksitas model tetap terkendali. Literatur *machine learning* menunjukkan bahwa *SVM* memiliki ketahanan yang baik terhadap *overfitting*, terutama ketika diterapkan pada dataset dengan ukuran kecil hingga menengah [7]. Ketahanan ini muncul karena model hanya bergantung pada *support vectors* sebagai elemen kunci pembentuk keputusan. Karakteristik tersebut menjadikan *SVM* sesuai untuk penelitian prediksi kebutuhan stok obat yang mengandalkan data penjualan historis dalam periode terbatas. Pendekatan ini memungkinkan apotek memperoleh model klasifikasi yang andal tanpa memerlukan volume data yang sangat besar.



Gambar 2.3 Perbandingan Overfitting dan Generalisasi Model yang Baik

Sumber : https://julienharbulot.com/overfitting.html?utm_source

Gambar ini membandingkan model yang overfitting dengan model yang mampu melakukan generalisasi baik melalui pola kurva error/loss pada data training dan validation/test. Ketika pelatihan dimulai, kedua kurva biasanya turun bersama karena model belajar sinyal utama. Pada titik 'best fit', loss training dan loss validation memiliki bentuk serupa dan selisihnya kecil, sehingga prediksi pada data baru tetap stabil. Overfitting ditunjukkan saat loss training terus menurun, tetapi loss validation berhenti membaik lalu meningkat; kondisi ini menandakan model terlalu kompleks, menangkap noise, dan menghasilkan generalization gap yang besar. Sebaliknya, underfitting terjadi bila kedua loss sama-sama tinggi karena kapasitas model terlalu rendah sehingga pola data tidak terwakili. Secara konseptual, gambar juga mengilustrasikan trade-off bias–variance: meningkatkan kompleksitas menurunkan bias namun menaikkan variance, sehingga diperlukan regularisasi, early stopping, atau validasi silang untuk memilih kompleksitas optimal. Dalam praktik, zona terbaik ditentukan dengan memonitor metrik pada validasi dan memilih epoch atau parameter ketika kinerja validasi mencapai nilai terbaik.

2.3 Teknik Evaluasi Algoritma *Machine Learning*

Teknik evaluasi algoritma *machine learning* berfungsi sebagai instrumen utama untuk menilai kualitas dan keandalan model klasifikasi yang dibangun. Evaluasi tidak hanya berperan menunjukkan tingkat ketepatan prediksi, tetapi juga memastikan bahwa model mampu bekerja secara konsisten ketika dihadapkan pada data baru. Literatur evaluasi algoritma terawasi menegaskan bahwa evaluasi yang tidak memadai berpotensi menghasilkan interpretasi kinerja yang keliru dan menyesatkan pengambilan keputusan.

Kebutuhan evaluasi menjadi semakin penting pada penelitian terapan yang menggunakan data historis sebagai dasar analisis. Model klasifikasi yang tampak baik pada data latih belum tentu memiliki kemampuan generalisasi yang memadai. Oleh karena itu, teknik evaluasi digunakan sebagai mekanisme objektif untuk menguji kesesuaian antara hasil prediksi model dan kondisi aktual yang tercermin pada data uji.

2.3.1 Metrik Evaluasi Kinerja Model Klasifikasi

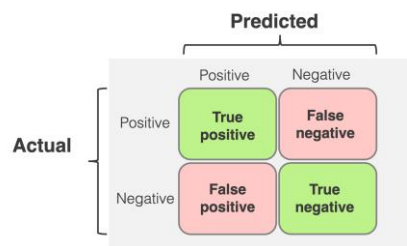
Akurasi digunakan untuk mengukur proporsi prediksi yang benar terhadap keseluruhan data uji. Metrik ini memberikan gambaran umum mengenai kinerja model klasifikasi. Literatur evaluasi *machine learning* menyebutkan bahwa akurasi sering dijadikan indikator awal dalam penilaian model karena kemudahannya untuk dipahami dan dihitung [10].

Keterbatasan akurasi muncul ketika data memiliki distribusi kelas yang tidak seimbang. Kondisi tersebut menyebabkan akurasi tinggi tidak selalu mencerminkan kemampuan model dalam mengenali seluruh kelas secara

proporsional. Oleh karena itu, akurasi perlu dikombinasikan dengan metrik evaluasi lain agar penilaian kinerja model menjadi lebih komprehensif.

Presisi dan *recall* memberikan perspektif yang lebih spesifik mengenai kinerja model klasifikasi. Presisi mengukur ketepatan model dalam memprediksi suatu kelas, sedangkan *recall* menunjukkan kemampuan model dalam mengenali seluruh data yang benar-benar termasuk dalam kelas tersebut. Penelitian prediktif di bidang kesehatan menunjukkan bahwa kedua metrik ini penting ketika kesalahan klasifikasi memiliki implikasi praktis yang signifikan [11]

Nilai *F1-score* digunakan sebagai ukuran harmonis yang menggabungkan presisi dan *recall*. Metrik ini memberikan representasi keseimbangan antara kedua ukuran tersebut dan sering digunakan pada dataset dengan distribusi kelas yang tidak merata. Literatur menunjukkan bahwa *F1-score* mampu memberikan gambaran kinerja model yang lebih stabil dibandingkan akurasi semata [12].



Gambar 2.4 Confusion Matrix Dasar Hitung Metrik Evaluasi Model

Sumber : https://www.evidentlyai.com/classification-metrics/confusion-matrix?utm_source

Gambar ini menampilkan confusion matrix biner 2x2 yang memetakan kelas aktual pada baris dan prediksi model pada kolom, sehingga empat keluaran dasar True Positive (TP), False Positive (FP), True Negative (TN), dan False

Negative (FN)—terlihat eksplisit. TP menunjukkan kasus positif yang diprediksi benar, TN kasus negatif yang diprediksi benar, FP adalah alarm palsu ketika negatif diprediksi positif, sedangkan FN adalah miss ketika positif diprediksi negatif. Dari empat sel ini, metrik evaluasi diturunkan secara konsisten: $akurasi = (TP+TN)/(TP+TN+FP+FN)$ menggambarkan ketepatan umum; $precision = TP/(TP+FP)$ menilai kemurnian prediksi positif; $recall/sensitivity = TP/(TP+FN)$ menilai kemampuan menangkap kasus positif; $specificity = TN/(TN+FP)$ menilai kemampuan menolak negatif; dan $F1 = 2 \cdot (precision \cdot recall) / (precision + recall)$ menyeimbangkan precision–recall. Gambar juga membantu membaca trade-off kesalahan: pada diagnosis medis, FN biasanya lebih kritis; pada deteksi spam, FP dapat lebih mengganggu. Dengan demikian, confusion matrix adalah fondasi auditabel untuk memilih metrik yang paling selaras dengan risiko keputusan. Jika probabilitas dipakai, ambang keputusan dapat diubah untuk menggeser FP–FN, lalu dievaluasi dengan ROC-AUC atau PR-AUC tambahan.

2.4 Alat Bantu Pemrograman dan Tools Pendukung

Alat bantu pemrograman dan *tools* pendukung memiliki peran penting dalam mendukung penerapan metode *machine learning* pada penelitian ini. Aplikasi Orange digunakan sebagai lingkungan *visual programming* yang memungkinkan proses pengolahan data, pemodelan menggunakan algoritma *Support Vector Machine* (SVM), serta evaluasi model dilakukan secara terstruktur dan sistematis. Pemanfaatan Orange memudahkan peneliti dalam membangun alur analisis tanpa memerlukan pengkodean yang kompleks, sehingga fokus

penelitian dapat diarahkan pada analisis hasil prediksi kebutuhan stok obat di apotek secara lebih efektif dan akurat.

2.4.1 Aplikasi Orange

Aplikasi Orange merupakan perangkat lunak *open-source* yang dirancang untuk mendukung analisis data dan penerapan *machine learning* melalui pendekatan *visual programming*. Karakteristik utama Orange terletak pada penggunaan *widget* yang merepresentasikan setiap tahapan analisis, mulai dari pemuatan data, *data preprocessing*, pemodelan, hingga evaluasi hasil. Pendekatan visual ini memungkinkan proses analisis dilakukan secara terstruktur dan transparan, sehingga setiap langkah dapat ditelusuri dan dipahami dengan mudah. Literatur menegaskan bahwa lingkungan *visual analytics* berperan penting dalam menjembatani konsep teoritis *machine learning* dengan implementasi praktis, khususnya bagi penelitian terapan yang tidak berfokus pada pengembangan kode tingkat lanjut [14]



Gambar 2.5 Aplikasi Orange

Sumber : <https://x.com/OrangeDataMiner/photo>

2.5. Pengelolaan Inventory dalam Apotek

Pengelolaan *inventory* obat menempati posisi strategis dalam operasional apotek karena berkaitan langsung dengan kesinambungan pelayanan kesehatan dan efisiensi sumber daya. Stok obat yang dikelola secara tepat memungkinkan apotek memenuhi kebutuhan pasien secara berkelanjutan tanpa menimbulkan pemborosan akibat kedaluwarsa atau penumpukan. Literatur farmasi menegaskan bahwa *inventory management* bukan sekadar aktivitas administratif, melainkan fungsi manajerial yang memerlukan perencanaan dan pengendalian berbasis data. Pandangan ini menempatkan inventory sebagai instrumen penghubung antara kebutuhan pasien dan kapasitas operasional apotek.

Peran strategis inventory juga tercermin dari dampaknya terhadap kinerja operasional. Praktik pengelolaan stok yang tidak terstruktur sering memicu ketidakseimbangan antara permintaan dan ketersediaan obat. Kondisi tersebut berdampak pada peningkatan biaya penyimpanan dan risiko hilangnya kepercayaan pelanggan. Penelitian lintas sektor menunjukkan bahwa pengelolaan inventory yang efektif berkontribusi signifikan terhadap peningkatan efisiensi organisasi secara keseluruhan.



Gambar 2.6 Diagram Alur Dasar Pengelolaan Inventory Farmasi Harian

Sumber : <https://share.google/images/R7HdkNAFbV7LVwQvf>

Gambar menunjukkan diagram siklus pengelolaan inventory farmasi, disusun melingkar agar menegaskan bahwa prosesnya berulang dan saling terkait. Siklus dimulai dari perencanaan kebutuhan (berdasarkan pola konsumsi dan target layanan), lalu berlanjut ke pengadaan. Setelah produk datang, tahap penerimaan menekankan verifikasi jumlah, kondisi kemasan, nomor batch, dan tanggal kedaluwarsa, sebelum obat masuk ke penyimpanan dengan pengaturan lokasi, keamanan, serta rotasi stok (FIFO/FEFO). Dari gudang, obat didistribusikan ke unit/pasien, dan setiap pengeluaran serta penerimaan dicatat pada kartu stok atau sistem agar saldo selalu mutakhir. Bagian pengendalian persediaan diilustrasikan sebagai pemeriksaan rutin (*cycle count/stock opname*), pemantauan stok minimum–maksimum, serta pemicu pemesanan ulang untuk mencegah stockout. Di sisi lain, pencatatan dan pelaporan menutup siklus melalui rekap pemakaian, sisa stok, dan temuan selisih. Terakhir, diagram memasukkan monitoring evaluasi dan penarikan/pemusnahan untuk item rusak atau kedaluwarsa, sehingga kepatuhan mutu dan keselamatan pasien terjaga. Keseluruhan alur membantu transparansi, akuntabilitas, dan audit kepatuhan terhadap standar penyimpanan serta distribusi obat nasional.

2.5.1 Karakteristik Inventory Obat

Inventory obat memiliki karakteristik khusus yang membedakannya dari persediaan pada sektor lain. Obat bersifat sensitif terhadap waktu karena memiliki masa kedaluwarsa serta implikasi langsung terhadap keselamatan pasien. Karakteristik ini menuntut pengelolaan yang lebih hati-hati dan presisi. Studi empiris pada fasilitas kesehatan menunjukkan bahwa kegagalan memahami

karakteristik inventory obat sering menjadi penyebab utama inefisiensi pengelolaan stok.

Kondisi tersebut memperkuat argumen bahwa teori inventory umum perlu diadaptasi ke dalam konteks farmasi. Adaptasi ini mencakup penyesuaian kebijakan penyimpanan, pengendalian jumlah stok, serta prioritas pengadaan berdasarkan tingkat kebutuhan dan risiko. Dengan demikian, konsep inventory obat tidak dapat dilepaskan dari konteks layanan kesehatan yang menuntut akurasi dan tanggung jawab tinggi.

2.5.2 Stok Obat dan Pola Permintaan

Permintaan obat bersifat dinamis dan dipengaruhi oleh berbagai faktor, seperti pola penyakit, perilaku masyarakat, serta karakteristik wilayah layanan. Variabilitas ini menimbulkan tantangan tersendiri dalam pengelolaan inventory apotek. Penelitian di berbagai fasilitas kesehatan menunjukkan bahwa fluktuasi permintaan merupakan penyebab utama terjadinya kelebihan atau kekurangan stok obat [17]. Pemahaman terhadap pola permintaan menjadi prasyarat penting dalam menyusun strategi pengadaan yang efektif.

Pendekatan tradisional yang mengandalkan pengalaman subjektif sering kali tidak mampu menangkap kompleksitas pola permintaan. Literatur manajemen farmasi menekankan pentingnya pemanfaatan data historis penjualan sebagai dasar analisis kebutuhan stok. Pendekatan ini memungkinkan identifikasi pola berulang yang dapat digunakan sebagai dasar perencanaan inventory yang lebih akurat [18].

2.6 Penelitian Terdahulu

Tabel penelitian terdahulu disusun untuk memberikan gambaran ringkas mengenai penelitian-penelitian sebelumnya yang relevan dan memiliki kesamaan metodologis, khususnya dalam penggunaan algoritma *Support Vector Machine* (SVM). Penyajian tabel ini bertujuan memudahkan perbandingan aspek peneliti, fokus kajian, data yang digunakan, serta hasil yang diperoleh dari masing-masing penelitian. Melalui perbandingan tersebut, posisi penelitian ini dapat diidentifikasi secara jelas dalam peta keilmuan, sekaligus menegaskan kontribusi dan kelebihanannya dibandingkan penelitian terdahulu yang telah ada.

Tabel 2.1 Tabel Penelitian Terdahulu

No	Peneliti	Judul	Tahun	Data dan Algoritma	Hasil
1	Abdullah & Abdulazeez	<i>Machine Learning Applications based on SVM Classification: A Review</i>	2021	Data sekunder multi-domain; algoritma <i>Support Vector Machine</i> untuk prediksi dan klasifikasi	SVM memiliki kemampuan prediksi yang stabil dan akurat pada berbagai jenis data numerik dan kategorikal
2	Zhang et al.	<i>Prediction of Surface Settlement in Shield-Tunneling Construction Process Using PCA-PSO-RVM Machine Learning</i>	2023	Data numerik konstruksi; PCA dan SVM untuk prediksi penurunan permukaan	SVM mampu memprediksi nilai penurunan tanah dengan tingkat kesalahan yang rendah

3	Devihosur et al.	<i>Estimating the Energy Demand and Carbon Emission Reduction Potential of Singapore's Future Road Transport Sector</i>	2024	Data transportasi dan energi; SVM untuk prediksi permintaan energi	SVM efektif dalam memprediksi kebutuhan energi jangka menengah dengan akurasi tinggi
4	Huang et al.	<i>Investigating Molecular Mechanisms Using Single-Cell Transcriptomics</i>	2025	Data biologis berdimensi tinggi; SVM untuk prediksi pola ekspresi gen	SVM mampu memprediksi hubungan molekuler kompleks pada data berskala besar
5	Gao et al.	<i>Consistent Quality Oriented Rate Control in HEVC</i>	2022	Data kompresi video; SVM untuk prediksi kualitas bit-rate	Model SVM menghasilkan prediksi kualitas video yang konsisten dan efisien
6	Penelitian Ini	<i>Analisis Inventory Apotek di Silangkitang Menggunakan Metode Support Vector Machine (SVM) untuk Prediksi Kebutuhan Stok Obat</i>	2025	Data historis penjualan dan stok obat; algoritma <i>Support Vector Machine</i> dengan aplikasi Orange	SVM digunakan untuk memprediksi kebutuhan stok obat secara akurat sebagai dasar pengambilan keputusan inventory

Tabel ini menegaskan bahwa sebagian besar penelitian terdahulu terkait inventory obat masih didominasi oleh pendekatan deskriptif seperti analisis ABC dan matriks. Penelitian yang secara eksplisit mengadopsi *Support Vector Machine* masih terbatas dan umumnya berada pada konteks non-inventory atau bersifat tinjauan umum. Oleh karena itu, penelitian ini memiliki kelebihan dengan mengintegrasikan metode SVM secara langsung pada prediksi kebutuhan stok

obat apotek skala lokal, sekaligus memanfaatkan alat bantu visual Orange untuk meningkatkan keterterapan hasil penelitian.

Aspek kebaruan penelitian ini terletak pada penerapan algoritma *Support Vector Machine* secara langsung untuk memprediksi kebutuhan stok obat pada apotek skala lokal, yang selama ini masih didominasi oleh pendekatan deskriptif seperti analisis ABC dan matriks klasifikasi. Penelitian terdahulu umumnya berhenti pada tahap pengelompokan obat berdasarkan nilai konsumsi atau frekuensi penggunaan, tanpa menyediakan mekanisme prediktif yang mampu mengantisipasi kebutuhan stok di masa mendatang. Penelitian ini melampaui keterbatasan tersebut dengan memanfaatkan kemampuan SVM dalam membangun batas keputusan nonlinier berdasarkan data historis penjualan obat. Keunggulan penelitian juga terlihat pada integrasi algoritma SVM dengan alat bantu visual Orange, yang memungkinkan proses pemodelan, pengujian, dan evaluasi dilakukan secara transparan serta mudah direplikasi tanpa pengkodean kompleks. Hasil penelitian diharapkan tidak hanya meningkatkan akurasi prediksi stok, tetapi juga memberikan dasar pengambilan keputusan inventory yang lebih proaktif, adaptif, dan aplikatif bagi pengelola apotek.

2.7 Flowchart Kerangka Penelitian

Flowchart kerangka penelitian disusun untuk menggambarkan alur penelitian secara sistematis dan terintegrasi, sehingga setiap tahapan dapat dipahami sebagai satu kesatuan proses yang saling berkaitan. Representasi visual ini berfungsi sebagai peta konseptual yang menjelaskan bagaimana penelitian bergerak dari identifikasi masalah hingga evaluasi hasil model. Literatur metodologi menegaskan bahwa visualisasi alur penelitian membantu memperjelas

logika internal penelitian dan memperkuat konsistensi antara tujuan, metode, dan hasil yang diharapkan .

Tahap awal dalam flowchart penelitian diawali dengan perumusan masalah dan penentuan tujuan penelitian. Permasalahan yang diangkat berangkat dari kebutuhan pengelolaan inventory obat di apotek yang menuntut ketepatan dalam memprediksi kebutuhan stok. Penentuan tujuan penelitian menjadi landasan arah seluruh tahapan berikutnya, karena setiap proses analisis dirancang untuk menjawab permasalahan yang telah dirumuskan secara jelas. Pendekatan sistematis dalam perumusan tujuan penelitian dianggap penting untuk menjaga koherensi metodologis sepanjang proses penelitian.

Tahapan berikutnya pada flowchart adalah pengumpulan data. Data yang digunakan berupa data historis penjualan obat yang merepresentasikan pola permintaan masa lalu. Data historis dipandang sebagai sumber informasi utama dalam penelitian prediktif karena mengandung pola laten yang dapat dimanfaatkan oleh algoritma *machine learning*. Literatur menunjukkan bahwa kualitas data pada tahap awal sangat menentukan keberhasilan model pada tahap selanjutnya, sehingga proses pengumpulan data harus dilakukan secara cermat dan terstruktur [20]

Setelah data terkumpul, flowchart penelitian mengarah pada tahap pengolahan data awal. Tahap ini mencakup pemeriksaan kelengkapan data, identifikasi data hilang, serta penyesuaian format data agar sesuai dengan kebutuhan analisis. Pengolahan data awal bertujuan memastikan bahwa data yang digunakan benar-benar siap untuk diproses oleh algoritma *machine learning*.

Penelitian-penelitian terdahulu menunjukkan bahwa pengolahan data yang tidak memadai dapat menurunkan akurasi dan keandalan model prediktif [21].

Tahap selanjutnya dalam flowchart adalah proses *preprocessing* data. Proses ini mencakup normalisasi, transformasi variabel, dan seleksi atribut yang relevan. *Preprocessing* berperan penting dalam meningkatkan kualitas data sehingga pola yang ada dapat dikenali secara optimal oleh algoritma. Literatur *machine learning* menegaskan bahwa *preprocessing* merupakan salah satu faktor kunci dalam membangun model yang stabil dan mampu melakukan generalisasi dengan baik [22].

Flowchart kemudian mengarahkan penelitian pada tahap pemodelan. Tahap ini melibatkan penerapan algoritma *Support Vector Machine* sebagai metode utama untuk melakukan klasifikasi dan prediksi kebutuhan stok obat. Pemilihan algoritma ini didasarkan pada kemampuannya dalam menangani data berdimensi tinggi dan membentuk batas keputusan yang optimal. Penelitian-penelitian terdahulu menunjukkan bahwa struktur pemodelan yang jelas dalam alur penelitian membantu memastikan bahwa algoritma diterapkan secara konsisten dengan tujuan analisis [21]

Setelah model dibangun, flowchart penelitian berlanjut ke tahap pengujian model. Tahap ini bertujuan untuk menilai kinerja model menggunakan data uji yang terpisah dari data latih. Pengujian model memberikan gambaran mengenai kemampuan model dalam memprediksi data yang belum pernah dilihat sebelumnya. Literatur metodologi menekankan bahwa pemisahan data latih dan data uji merupakan langkah penting untuk menghindari bias evaluasi dan memastikan objektivitas penilaian kinerja model.

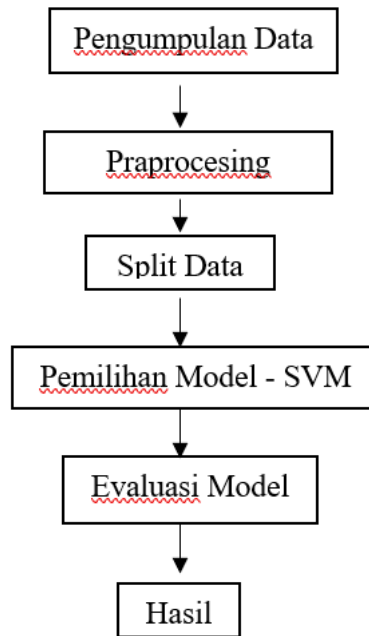
Tahap evaluasi model menjadi bagian krusial dalam flowchart kerangka penelitian. Evaluasi dilakukan dengan menggunakan metrik evaluasi yang relevan untuk menilai kinerja klasifikasi. Hasil evaluasi digunakan sebagai dasar untuk menilai apakah model telah memenuhi kriteria keandalan yang ditetapkan. Penelitian di berbagai bidang menunjukkan bahwa evaluasi model yang sistematis membantu peneliti memahami kekuatan dan keterbatasan model secara komprehensif [23]

Flowchart penelitian juga menggambarkan kemungkinan iterasi dalam proses analisis. Apabila hasil evaluasi menunjukkan bahwa kinerja model belum optimal, alur penelitian dapat kembali pada tahap *preprocessing* atau pemodelan untuk melakukan penyesuaian. Pendekatan iteratif ini mencerminkan karakteristik penelitian *machine learning* yang bersifat adaptif dan berorientasi pada perbaikan berkelanjutan. Literatur menyebutkan bahwa fleksibilitas alur penelitian merupakan keunggulan pendekatan berbasis data dibandingkan metode analisis statis.

Tahap akhir dalam flowchart kerangka penelitian adalah interpretasi hasil dan penarikan kesimpulan. Hasil prediksi kebutuhan stok obat diinterpretasikan dalam konteks pengelolaan inventory apotek. Interpretasi ini bertujuan mengaitkan temuan empiris dengan permasalahan praktis yang menjadi fokus penelitian. Literatur metodologi menegaskan bahwa interpretasi hasil yang kontekstual penting agar penelitian tidak berhenti pada capaian teknis, tetapi memberikan kontribusi nyata bagi pengambilan Keputusan.

Flowchart kerangka penelitian juga berfungsi sebagai alat komunikasi ilmiah. Representasi visual alur penelitian memudahkan pembaca memahami

struktur penelitian tanpa harus menelusuri seluruh uraian metodologis secara detail. Penelitian-penelitian terdahulu menunjukkan bahwa visualisasi alur kerja meningkatkan transparansi dan keterbacaan penelitian, khususnya pada studi yang melibatkan proses analitis yang kompleks.



Gambar 2.8 Flowchart Kerangka Penelitian

Gambar ini menggambarkan alur kerja dalam pengolahan data menggunakan metode Support Vector Machine (SVM). Proses dimulai dengan tahap Pengumpulan Data, di mana data yang dibutuhkan dikumpulkan untuk analisis lebih lanjut. Setelah itu, data akan mengalami tahap *Praprocessing*, yang bertujuan untuk membersihkan dan mempersiapkan data agar siap digunakan dalam model. Langkah berikutnya adalah Split Data, di mana data dibagi menjadi dua bagian: satu untuk pelatihan dan satu lagi untuk pengujian model. Setelah pembagian data, dilakukan Pemilihan Model - SVM, yang merupakan proses pemilihan model yang tepat menggunakan algoritma Support Vector Machine. Selanjutnya, dilakukan Evaluasi Model, untuk menilai seberapa baik kinerja

model berdasarkan data pengujian. Akhirnya, hasil evaluasi diperoleh pada langkah Hasil, yang memberikan wawasan tentang akurasi dan efektivitas model yang telah diterapkan. Gambar ini secara keseluruhan menjelaskan alur dari pengolahan data hingga evaluasi model SVM.