

BAB II

TINJAUAN PUSTAKA

2.1. Landasan Teori

2.1.1. Prediksi

Prediksi adalah proses memperkirakan suatu peristiwa di masa depan berdasarkan data masa lalu dan kini. Tujuannya untuk mengurangi kesalahan antara perkiraan dan kejadian sebenarnya. Meskipun tidak selalu akurat prediksi berusaha memberikan estimasi yang mendekati kenyataan (Mustika et al., 2021). Dalam bisnis prediksi membantu menentukan persediaan barang optimal serta mengelola stok guna mengurangi kesalahan perencanaan. Selain itu, prediksi memungkinkan ekstraksi wawasan dari data besar, menjadikan data mining sebagai elemen penting dalam analisis prediktif (Alfani W.P.R., Rozi, & Sukmana, 2021).

Keberadaan peramal dan ahli nجوم di era kuno serta abad pertengahan mencerminkan rasa ingin tahu manusia tentang masa depan. Seiring waktu, ketertarikan ini berkembang, salah satunya melalui rubrik astrologi yang mulai disindikasikan sejak tahun 1973 (Luthfiarta, Febriyanto, Lestiawan, & Wicaksono, 2020). Tujuan utama peramalan kebijakan adalah memperoleh informasi mengenai perubahan yang akan terjadi di masa depan serta dampaknya terhadap implementasi kebijakan dan konsekuensi yang mungkin timbul (Kafil, 2019).

2.1.2. Permintaan Produk

Permintaan produk merujuk pada kebutuhan konsumen terhadap barang-barang yang disediakan oleh perusahaan. Barang-barang ini umumnya disimpan terlebih dahulu di gudang sebagai persediaan, sebelum akhirnya didistribusikan atau dijual sesuai dengan permintaan pasar di masa mendatang (Wau, 2022).

Monitoring permintaan produk dalam suatu unit usaha sangat penting untuk memperoleh data yang akurat mengenai kebutuhan pasar terhadap barang yang disediakan. Manajemen permintaan menjadi aktivitas krusial dalam dunia bisnis, karena perusahaan harus mampu mengelola ketersediaan produk secara proporsional agar sesuai dengan tingkat permintaan. Hal ini bertujuan untuk memastikan terpenuhinya kebutuhan konsumen, sekaligus menghindari risiko kehabisan atau penumpukan stok yang dapat merugikan perusahaan (S Pasaribu, 2021).

Permintaan produk merupakan komponen penting dalam aktivitas perdagangan, karena mencerminkan seberapa besar kebutuhan pasar terhadap barang yang ditawarkan perusahaan. Setiap produk yang disediakan memiliki nilai ekonomi yang dapat dipasarkan, dan tujuan utamanya adalah untuk memenuhi permintaan konsumen secara tepat, sehingga dapat mendukung pencapaian target pendapatan yang diharapkan oleh perusahaan (Andarista, Fitriansyah, & Ramliyana, 2022).

Data atau informasi mengenai permintaan produk memainkan peran penting dalam pengelolaan distribusi barang. Ketidaktepatan dalam pencatatan data dapat menyebabkan ketidakseimbangan antara produk yang tersedia dan kebutuhan pasar yang berujung pada penumpukan barang tidak laku atau kekurangan produk yang justru dibutuhkan. Untuk menghindari hal tersebut, perusahaan perlu menerapkan sistem informasi yang mendukung pencatatan dan analisis permintaan secara akurat. Teknologi ini membantu meningkatkan efisiensi kerja serta memberikan kemudahan bagi pengguna dalam mengambil keputusan strategis terkait permintaan produk yang sesuai dengan kebutuhan pasar (Fatmawati, 2019).

2.1.3. *Vending machine*

Vending machine telah ada sejak abad pertama, ketika seorang ilmuwan dari Alexandria menciptakan mekanisme otomatis untuk mendistribusikan barang. Mesin ini mulai berkembang secara modern pada tahun 1880 di London sebagai alat penjual kartu pos, kemudian diperkenalkan di Amerika Serikat pada tahun 1888 oleh Thomas Adams Gum Company untuk menjual permen karet. Saat ini, *vending machine* telah menjadi perangkat otomatis yang mampu mengeluarkan berbagai jenis barang setelah pembayaran dilakukan, seperti makanan ringan, minuman, rokok, tiket, hingga produk bernilai tinggi seperti emas dan permata. Dengan perkembangan teknologi, metode pembayaran pun semakin beragam, tidak hanya menggunakan uang tunai tetapi juga kartu kredit dan dompet digital, menjadikannya semakin praktis dan mudah digunakan (Wiyanti, Soedjoko, & Safaatullah, 2020).

Keberadaan *vending machine* membawa berbagai manfaat, baik bagi konsumen maupun pemilik usaha. Bagi konsumen, *vending machine* memberikan kemudahan dalam bertransaksi karena dapat diakses kapan saja tanpa perlu antri atau berinteraksi langsung dengan penjual. Selain itu, bagi pelaku bisnis, penggunaan *vending machine* dapat mengurangi biaya operasional yang biasanya diperlukan dalam menjalankan toko fisik, seperti biaya tenaga kerja dan sewa tempat. Dengan sistem otomatisasi yang terus berkembang, *vending machine* menjadi solusi bisnis yang semakin diminati, terutama di era modern yang mengutamakan efisiensi dan kenyamanan dalam berbagai sektor industri (Sujana, Hanipah, Agustina, S, & Aulia, 2019).

2.1.4. *Data Processing*

Data pre-processing adalah proses mengubah data ke dalam format yang lebih optimal agar dapat diproses secara lebih cepat dan efisien dalam data mining, *machine learning*, maupun berbagai tugas *data science* lainnya (Varma et al., 2023). Proses ini bertujuan untuk memastikan bahwa data yang digunakan bersih, terstruktur, dan siap untuk diproses lebih lanjut. Adapun beberapa tahapan data pre-processing dalam penelitian ini adalah sebagai berikut:

1. *Date Time Index*

Date time index merupakan proses awal dalam pengolahan data deret waktu (*time series*) yang memastikan setiap titik data memiliki indeks waktu yang tepat, konsisten, dan berurutan. Penataan indeks waktu ini penting untuk menjamin bahwa model prediktif seperti LSTM mampu memahami urutan kronologis peristiwa yang terjadi dalam *dataset* (Khan et al., 2024).

2. *Handling Missing Values*

Handling missing values adalah tahapan krusial dalam proses pembersihan data yang bertujuan untuk mengatasi ketidaklengkapan informasi, khususnya pada atribut yang memengaruhi pemodelan seperti kategori produk. Teknik yang umum digunakan meliputi penghapusan entri kosong, imputasi nilai berdasarkan rata-rata, median, atau interpolasi berbasis waktu (El-Bakry et al., 2021).

3. Ekstraksi Fitur Waktu

Ekstraksi fitur waktu merupakan proses pemecahan informasi waktu ke dalam komponen-komponen seperti bulan (*month*) dan hari (*day*), dengan tujuan untuk membantu model mengenali pola musiman atau fluktuasi yang terjadi pada periode waktu tertentu (Chernikov et al., 2022).

4. Normalisasi Data

Normalisasi adalah teknik penskalaan yang mengubah nilai-nilai fitur numerik ke dalam rentang yang sama, biasanya antara 0 dan 1, dengan tujuan memberikan kontribusi yang setara dari setiap fitur terhadap performa model (Tran et al., 2021). Metode Min-Max Scaling adalah bentuk normalisasi yang paling umum, mengubah setiap nilai dalam sebuah kolom secara linear. Dimana Min-Max ditunjukkan pada Persamaan 2.1.

$$X' = \frac{X - X_{min}}{X_{max} - X_{min}} \quad (2.1)$$

Di mana X_{scaled} adalah nilai setelah normalisasi, X adalah nilai asli data, X_{min} adalah nilai minimum dalam kolom, dan X_{max} adalah nilai maksimum dalam kolom tersebut.

2.1.5. *Exploratory Data Analysis (EDA)*

Exploratory Data Analysis (EDA) merupakan tahap awal dalam proses analisis data yang bertujuan untuk memperoleh pemahaman menyeluruh terhadap struktur, karakteristik, serta pola yang terkandung dalam data sebelum dilakukan pemodelan lanjutan. EDA membantu peneliti dalam mendeteksi distribusi, outlier, pola musiman, serta hubungan antar variabel melalui pendekatan statistik deskriptif dan visualisasi (Pearson et al., 2020).

1. *Arsitektur Exploratory Data Analysis (EDA)*

Arsitektur Exploratory Data Analysis (EDA) pada penelitian ini dirancang sebagai kerangka analisis terstruktur yang digunakan untuk mengeksplorasi dan memahami karakteristik data penjualan sebelum tahap pemodelan dilakukan. Pendekatan EDA menekankan keterkaitan antar tahapan analisis guna

memperoleh pemahaman data yang menyeluruh, sehingga proses pemodelan dapat disesuaikan dengan kondisi data yang sebenarnya. Melalui pendekatan statistik dan visualisasi, EDA berperan dalam mengungkap struktur, pola, serta hubungan yang terdapat dalam data penjualan *vending machine*, selanjutnya menjadi dasar dalam perancangan model prediksi permintaan (Wickham et al., 2019).

Secara konseptual, arsitektur EDA disusun dalam tahapan yang saling terintegrasi, dimulai dari input data dan pemeriksaan struktur data untuk memastikan kesesuaian format dan kelengkapan variabel. Tahapan berikutnya meliputi analisis statistik deskriptif untuk menggambarkan karakteristik data, analisis pola waktu untuk mengidentifikasi tren dan pola musiman, serta analisis hubungan antar variabel guna mengetahui keterkaitan antara faktor waktu dan volume penjualan. Proses EDA juga mencakup deteksi *outlier* untuk mengidentifikasi data ekstrem, serta penentuan fitur relevan yang akan digunakan pada tahap pemodelan. Dengan demikian, arsitektur EDA memastikan bahwa data yang digunakan dalam pemodelan prediksi berbasis *Long Short-Term Memory* (LSTM) telah dipahami secara komprehensif dan selaras dengan tujuan penelitian. (Hyndman & Athanasopoulos, 2021).

2. Statistik Deskriptif

Statistik deskriptif digunakan untuk memberikan gambaran umum mengenai data penjualan. Parameter statistik yang digunakan meliputi nilai rata-rata (*mean*), nilai tengah (*median*), nilai minimum, nilai maksimum, dan simpangan baku (*standard deviation*) (Bruce et al., 2020). Secara matematis, statistik deskriptif dirumuskan sebagai berikut.

- *Mean* (rata - rata)

$$\mu = \frac{1}{n} \sum_{i=1}^n x_i \quad (2.2)$$

Keterangan:

μ = nilai rata – rata data

x_i = nilai data ke- i

n = jumlah total data

- *Standard Deviation*

$$\sigma = \sqrt{\frac{1}{n} \sum_{i=1}^n (x_i - \mu)^2} \quad (2.3)$$

Keterangan:

σ = simpangan baku

μ = nilai rata – rata data

x_i = nilai data ke- i

n = jumlah total data

Statistik ini digunakan untuk memahami sebaran dan variasi data penjualan pada kategori *food* dan *drink*.

3. Analisis Pola Waktu (*Time Based Analysis*)

Analisis pola penjualan berdasarkan waktu, data dikelompokkan (*aggregation*) berdasarkan interval harian dan bulanan. Proses ini bertujuan untuk mengidentifikasi tren jangka panjang serta pola musiman yang muncul pada periode tertentu (Hyndman & Athanasopoulos, 2021). Secara umum, agregasi data dilakukan dengan.

$$X_t^{(agg)} = \sum_{i=1}^k x_i \quad (2.4)$$

Keterangan:

$X_t^{(agg)}$ = total penjumlahan pada interval waktu ke- t

x_i = nilai penjumlahan individual

k = jumlah data dalam satu interval waktu

4. Analisis Hubungan Antar Variabel

Hubungan antar variabel dianalisis menggunakan koefisien korelasi Pearson untuk mengukur tingkat hubungan linier antara dua variabel (Schober et al., 2021). Persamaan korelasi Pearson dirumuskan sebagai berikut.

$$\gamma = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum (x_i - \bar{x})^2 \sum (y_i - \bar{y})^2}} \quad (2.5)$$

Keterangan:

γ = koefisien korelasi Pearson

x_i, y_i = pasangan data ke- i

$\bar{x}, \bar{y}$ = nilai rata – rata variabel x dan y

Analisis ini digunakan untuk mengetahui keterkaitan antara variabel waktu dan volume penjualan produk.

5. Deteksi *Outlier*

Deteksi *outlier* dilakukan menggunakan metode *Interquartile Range* (IQR) untuk mengidentifikasi nilai ekstrem yang dapat memengaruhi performa model (Aggarwal, 2021). Nilai IQR dihitung sebagai.

$$IQR = Q_3 - Q_1 \quad (2.6)$$

Keterangan:

Q_1 = kuartil pertama

Q_3 = kuartil ketiga

Data dikategorikan sebagai *outlier* jika memenuhi kondisi.

$$x < Q_1 - 1.5 \times IQR \text{ atau } x > Q_3 + 1.5 \times IQR \quad (2.7)$$

6. Visualisasi Data

Visualisasi menjadi alat penting dalam EDA karena mampu menyajikan pola dan tren yang tidak mudah ditangkap oleh statistik deskriptif saja. Grafik penjualan bulanan, dapat menunjukkan puncak permintaan di waktu tertentu, sementara grafik harian dapat mengungkap fluktuasi stok harian yang harus dikelola secara tepat. Wawasan yang diperoleh dari proses EDA menjadi landasan dalam penyusunan model prediksi permintaan yang akurat (Waskom, 2021).

2.1.6. *Deep learning*

Deep learning merupakan cabang dari kecerdasan buatan (*Artificial Intelligence*) dan *machine learning* yang dikembangkan berdasarkan konsep jaringan saraf tiruan (*neural network*) dengan banyak lapisan (*multiple layers*). Teknologi ini dirancang untuk meningkatkan akurasi dalam berbagai tugas, seperti deteksi objek, pengenalan suara, penerjemahan bahasa, dan lain sebagainya. Berbeda dengan teknik *machine learning* tradisional, *Deep learning* mampu mengenali pola dan mengekstraksi fitur dari data secara otomatis, tanpa aturan eksplisit atau pengetahuan khusus dari manusia. Kemampuannya dalam memproses data kompleks menjadikannya metode yang efektif dalam berbagai aplikasi berbasis kecerdasan buatan (Yousefi, Alaghband, & Garibay, 2019).

Deep learning adalah sekumpulan algoritma dalam machine learning yang dirancang untuk belajar pada berbagai tingkat abstraksi dengan memanfaatkan *Artificial Neural Network* (ANN), sebuah struktur yang terinspirasi dari cara kerja otak manusia. ANN terdiri dari tiga atau lebih lapisan yang memungkinkan model memproses data besar, belajar dari pengalaman, dan mengenali pola kompleks

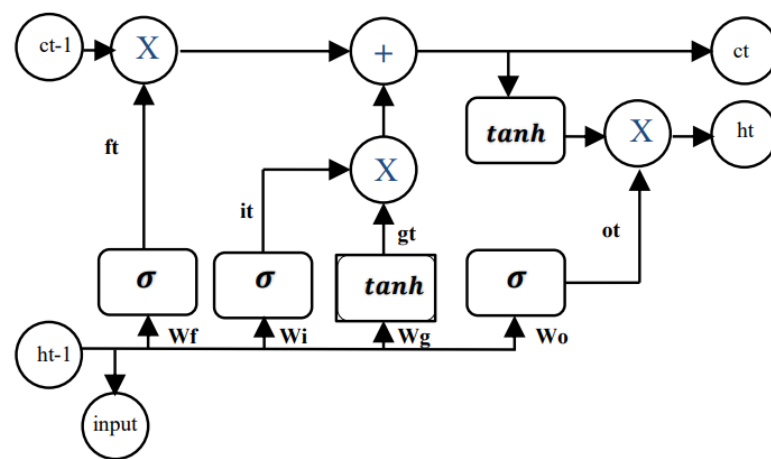
secara otomatis (Muhammad Haris Diponegoro, Sri Suning Kusumawardani, & Indriana Hidayah, 2021). Setiap tingkat dalam jaringan ini merepresentasikan konsep yang berbeda, memungkinkan sistem *deep learning* untuk mengekstraksi fitur mendalam dari data dan menyelesaikan masalah yang sulit diatasi dengan algoritma machine learning tradisional. Salah satu penerapannya adalah dalam permainan catur, di mana kecerdasan buatan (AI) mampu menganalisis jutaan langkah dari data pertandingan sebelumnya untuk menentukan strategi terbaik dalam hitungan detik. Kemampuan ini menunjukkan bahwa *deep learning* dapat "belajar" secara mandiri dan mengembangkan solusi optimal dari pola yang ditemukan (Raup, Ridwan, Khoeriyah, Supiana, & Zaqiah, 2022).

2.1.7. Long Short Term Memory (LSTM)

Long Short Term Memory (LSTM) adalah salah satu jenis jaringan saraf tiruan yang dirancang untuk menyimpan, menganalisis, serta memprediksi data dalam jangka waktu yang panjang. Keunggulan utama dari LSTM terletak pada kemampuannya dalam mempertahankan informasi sekuensial dalam jangka panjang, yang sulit dicapai dengan metode pemrosesan data konvensional (Wiranda & Sadikin, 2019).

LSTM mengatasi masalah ini dengan menggunakan mekanisme khusus berupa *gates* (gerbang) yang mengatur aliran informasi, sehingga model dapat mempertahankan atau melupakan data sesuai kebutuhan, dan meningkatkan kinerja dalam analisis data berurutan (Cholissodin & Soebroto, 2021). Struktur gerbang ini meliputi *input gate* untuk mengatur informasi baru yang masuk, *forget gate* untuk menentukan informasi mana yang harus dilupakan, serta *output gate* untuk mengontrol informasi yang akan diteruskan sebagai output (Larasati & Primandari,

2021). LSTM memiliki beberapa lapisan tersembunyi (*hidden layers*) yang berfungsi untuk mengelola aliran informasi. Saat data mengalir melalui jaringan, setiap sel memori bertugas menyimpan informasi yang dianggap relevan dan membuang informasi yang tidak penting. (Hasiholan, Cholissodin, & Yudistira, 2022).



Gambar 2. 1 Arsitektur LSTM

Sumber: (Hasiholan et al., 2022)

Pada Gambar 2.2 merupakan arsitektur LSTM (Long Short-Term Memory) yang terdiri dari serangkaian sel memori yang memiliki struktur unik untuk menyaring informasi melalui mekanisme gerbang (*gate*). Setiap gerbang memiliki persamaan matematis yang digunakan untuk mengontrol bagaimana informasi disimpan, diperbarui, atau dibuang dalam proses pembelajaran (Bastian Sianturi, Cholissodin, & Yudistira, 2023). Dalam model LSTM, terdapat empat jenis gerbang utama yang berperan dalam mengatur aliran data, yaitu *forget gate*, *input gate*, *cell gate*, dan *output gate*. Berikut ini adalah tahapan proses penerapan metode LSTM:

1. *Forget Gate*

Proses dimulai dengan forget gate, bertugas memutuskan informasi apa dari cell state yang perlu dilupakan (Karno, 2020). Proses ini dihitung menggunakan fungsi aktivasi sigmoid (σ), menerima input berupa status tersembunyi sebelumnya (h_{t-1}) dan input saat ini (x_t). Fungsi sigmoid menghasilkan nilai antara 0 (menghapus informasi) dan 1 (menyimpan informasi). Dimana Forget Gate ditunjukkan pada Persamaan 2.2.

$$f_t = \sigma (W_f \cdot [h_{t-1}, x_t] + b_f) \quad (2.8)$$

Keterangan:

W_f = Bobot dari *forget gate*

h_{t-1} = State sebelumnya atau state pada waktu t-1

x_t = Input pada waktu t.

σ = Fungsi aktivasi *sigmoid*

2. *Input Gate*

Input gate berfungsi untuk memfilter dan menambah informasi baru ke dalam cell state. Dalam prosesnya, digunakan fungsi aktivasi, yaitu fungsi sigmoid bertugas memutuskan seberapa banyak informasi yang akan diperbarui, dan fungsi tanh membuat kandidatvektor nilai baru untuk dimasukkan ke memori (Akbar, Santoso, & Warsito, 2023). Input gate ditunjukkan pada persamaan 2.3.

$$i_t = \sigma (W_i \cdot [h_{t-1}, x_t] + b_i) \quad (2.9)$$

Keterangan:

W_i = Bobot dari input *gate*

h_{t-1} = State sebelumnya atau state pada waktu t-1

x_t = Input pada waktu t.

σ = Fungsi aktivasi *sigmoid*

3. Cell Gate

Cell state berfungsi untuk menyimpan informasi jangka panjang yang relevan. Pembaruan nilai dalam cell state dilakukan dengan memanfaatkan informasi dari input gate dan forget gate. Setelah informasi yang perlu dilupakan diputuskan oleh forget gate, informasi baru yang dipilih oleh input gate akan digabungkan dengan memori yang ada untuk memperbarui cell state (Selle, Yudistira, & Dewi, 2022). Cell gate ditunjukkan pada persamaan 2.4.

$$\bar{C}_t = \tanh (W_c \cdot [h_{t-1}, x_t] + b_c) \quad (2.10)$$

Keterangan:

$\bar{C}_t$ = *Intermediate cell state*

W_c = Bobot dari *cell state*

h_{t-1} = *State* sebelumnya atau *state* pada waktu t-1

x_t = Input pada waktu t.

4. Output Gate

Berfungsi untuk memilih bagian dari memory cell yang akan dihasilkan, dengan memanfaatkan fungsi aktivasi sigmoid. Nilai dari memory cell diproses melalui fungsi aktivasi tanh. Kemudian, kedua nilai dari gatetersebut dikalikan untuk menghasilkan keluaran akhir (h_t) (Nazhiroh, Fitria, & Permana, 2024). Output gate ditunjukkan pada persamaan 2.5.

$$o_t = \sigma (W_o \cdot [h_{t-1}, x_t] + b_o) \quad (2.11)$$

Keterangan:

W_o = Bobot dari *output gate*

h_{t-1} = *State* sebelumnya atau *state* pada waktu t-1

x_t = Input pada waktu t.

σ = Fungsi aktivasi *sigmoid*

2.1.8. Evaluation

Untuk mengukur hasil performa dari algoritma yang digunakan untuk peramalan, terdapat beberapa jenis metrik pengukuran yang umum digunakan. Contoh metrik pengukuran yang menggambarkan persentase dari kesalahan peramalan adalah *Mean Absolute Error* (MAE), *Mean Squared Error* (MSE), dan *Root Mean Squared Error* (RMSE).

1. Mean Absolute Error (MAE)

Mean Absolute Error (MAE) merupakan salah satu metrik yang digunakan untuk menilai tingkat akurasi suatu model peramalan. MAE menghitung rata-rata selisih absolut antara nilai prediksi dan nilai aktual, sehingga menunjukkan seberapa besar kesalahan rata-rata dalam peramalan. Semakin kecil nilai MAE, semakin akurat model dalam memprediksi data (Suryanto, 2019). MAE dapat dihitung menggunakan persamaan 2.6.

$$MAE = \frac{1}{n} \sum_{i=1}^n |y_i - \hat{y}_i| \quad (2.12)$$

Keterangan:

y_i = Nilai aktual

$\hat{y}_i$ = Hasil prediksi

n = Jumlah data pelatihan

2. Mean Squared Error (MSE)

Mean Squared Error (MSE) dihitung dengan mengambil rata-rata dari kuadrat selisih antara nilai prediksi yang dihasilkan oleh model dan nilai aktual dari data yang diamati. Metrik ini memberikan gambaran tentang seberapa besar kesalahan prediksi, dimana nilai MSE yang lebih kecil menunjukkan bahwa model memiliki akurasi yang lebih baik dalam melakukan estimasi (Nuha, 2023).

Fungsi loss function yang digunakan dalam training adalah Mean Squared Error (MSE), yang menghitung rata-rata selisih kuadrat antara nilai aktual dan prediksi. MSE ditunjukkan pada persamaan 2.7.

$$MSE = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2 \quad (2.13)$$

Keterangan:

y_i = Nilai aktual

$\hat{y}_i$ = Hasil prediksi

n = Jumlah data pelatihan

3. *Root Mean Square Error* (RMSE)

Root Mean Square Error (RMSE) adalah metrik yang digunakan untuk mengukur tingkat kesalahan dalam hasil prediksi. Semakin kecil nilai RMSE (mendekati nol), semakin akurat model dalam melakukan peramalan (Suprayogi, Trimaijon, & Mahyudin, 2014). RMSE berfungsi sebagai metode alternatif dalam mengevaluasi teknik peramalan dengan menghitung rata-rata akar kuadrat dari jumlah kesalahan prediksi. Metrik ini mudah diimplementasi serta penggunaannya yang luas dalam berbagai studi yang berkaitan dengan prediksi dan peramalan (Sanjaya & Heksaputra, 2020). RMSE ditunjukkan pada persamaan 2.8.

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2} \quad (2.14)$$

Keterangan:

y_i = Nilai aktual

$\hat{y}_i$ = Hasil prediksi

n = Jumlah data pelatihan

2.2. Penelitian Terkait

Berbagai penelitian telah dilakukan untuk mengembangkan metode prediksi dalam analisis data penjualan dan manajemen stok. Penelitian-penelitian ini memberikan wawasan tentang efektivitas berbagai algoritma, termasuk metode berbasis *deep learning* seperti LSTM dalam menangani data deret waktu. Tabel 2. 1 merangkum beberapa penelitian terkait yang relevan dengan penelitian ini.

Tabel 2. 1 Penelitian Terkait

No	Nama Peneliti/Journal		Hasil Penelitian
1	Peneliti (tahun)	(Prayetno, Purba, Wirawan, & Sweet, 2025)	Penelitian ini menyoroti efektivitas metode berbasis data dalam meningkatkan efisiensi manajemen inventaris, menunjukkan bahwa metode LSTM secara efektif memprediksi stok suku cadang dengan akurasi yang signifikan, mencapai MAE 12%, MSE 2%, dan RMSE 15%, sehingga berhasil menangkap pola musiman dan tren dalam data penjualan, meningkatkan perkiraan kebutuhan stok. Penelitian ini memiliki keterbatasan berupa variabel input yang terbatas pada data penjualan, tidak adanya perbandingan dengan algoritma lain.
	Judul	Purchasing Prediction Using Machine Learning Algorithms for Optimizing Inventory Management	
	Algoritma/ Metode	Long Short-Term Memory (LSTM)	
2	Peneliti (tahun)	(Rivaldo, Maesaroh, Informatika, & Buana, 2024)	Dataset yang digunakan dalam penelitian ini berisi data penjualan historis dari perusahaan e-commerce selama enam bulan, yaitu dari April 2024 hingga September 2024. Data ini mencakup catatan transaksi harian, penelitian ini fokus pada analisis satu produk dari total 18 produk.
	Judul	Comparison of Long Short Term Memory (LSTM) and LightGBM Algorithms to Improve Inventory Stock Efficiency through Forecasting	

No	Nama Peneliti/Journal		Hasil Penelitian
	Algoritma/ Metode	LSTM dan LightGBM	Penelitian menunjukkan bahwa LSTM secara signifikan lebih unggul dibandingkan LightGBM dalam memprediksi barang keluar, dibuktikan dengan nilai MAE sebesar 82.20, RMSE sebesar 113.46, dan MAPE sebesar 30.15%, yang lebih rendah dibandingkan LightGBM.
3	Peneliti (tahun)	(Gunawan, Ricky, Karen Valerine, M Bairaja Dimiliu, 2024)	Dataset yang digunakan dalam penelitian ini diperoleh dari penyedia dataset publik Kaggle. Dataset tersebut berisi data penjualan toko furnitur yang mencakup periode dari Januari 2014 hingga Februari 2017. Secara keseluruhan, dataset ini terdiri dari 2.121 catatan dengan 21 kolom.
	Judul	Analisis Prediksi Penjualan Toko Furnitur Dengan Metode Long Short Term Memory (LSTM)	
	Algoritma/ Metode	LSTM	Penelitian ini menunjukkan bahwa LSTM efektif dalam memprediksi penjualan furnitur, membantu optimalisasi manajemen penjualan dan inventaris. Model mencapai RMSE 39,27%, MAE 32,74%, dan MAPE 42,28%, mengindikasikan akurasi prediksi yang dapat diterima.
4	Peneliti (tahun)	(Mehmood et al., 2023)	Penelitian ini menggunakan dataset penjualan dari 20 mesin penjual minuman dingin di Manchester selama periode 1 Januari–31 Mei 2023, dan menyimpulkan bahwa penerapan prediksi deret waktu dapat meningkatkan efisiensi restocking.
	Judul	<i>Vending machine</i> Product Demand Prediction using Machine Learning Algorithms	
	Algoritma/ Metode	XGBoost, FB Prophet, ARIMA, dan Support Vector Regresi,	Penelitian ini memiliki keterbatasan karena hanya memanfaatkan data dalam periode lima bulan dan ruang lingkup yang terbatas pada satu lokasi, sehingga generalisasi hasilnya masih

No	Nama Peneliti/Journal		Hasil Penelitian
			terbatas. Selain itu, meskipun beberapa algoritma diuji, hasil evaluasi menunjukkan bahwa Support Vector Regression (SVR) masih memiliki performa yang kurang optimal dengan nilai RMSE 25.08, MAE 11.5, dan MAPE 26.7% sehingga akurasi prediksi tidak sebaik pendekatan yang lainnya.
5	Peneliti (tahun)	(Mulyana, YkhSanur, FikriYadi, Sahroni, & Sumarsono, 2023)	Penerapan model prediksi stok menggunakan Backpropagasi memerlukan data pelatihan dengan Mean Absolute Error (MAE) untuk menilai kehilangan dataset. Dataset terdiri dari berbagai jenis produk dan data stok yang dikumpulkan dari BabyQu selama tahun 2022 untuk tujuan pelatihan dengan mencapai tingkat akurasi 85,72% dengan kehilangan data 24,36.
	Judul	Penerapan Metode Neural Network Dengan Struktur Backpropagation Untuk Memprediksi Kebutuhan Stok Pada Toko Umkm Perlengkapan Bayi Babyqu	
	Algoritma/ Metode	Artificial Neural Network	
6	Peneliti (tahun)	(Thimoty Tombeng & Ardian, 2021)	Penelitian ini berfokus pada pengembangan model penjualan di industri ritel menggunakan jaringan Memori Jangka Pendek Panjang (LSTM), dengan data transaksi dari supermarket di Taiwan. Model ini menganalisis berbagai faktor yang mempengaruhi penjualan, dengan Kesalahan Persentase Mutlak Rata-rata (MAPE) sebesar 7,2% untuk bir dan 6,1% untuk susu segar, yang mengindikasikan efektivitas pendekatan <i>deep learning</i> dalam peramalan penjualan.
	Judul	Market Sales Prediction Using <i>Deep learning</i> Approach	
	Algoritma/ Metode	Long Short Term Memory (LSTM)	

No	Nama Peneliti/Journal		Hasil Penelitian
7	Peneliti (tahun)	(Anshory, Priyandari, & Yuniaristanto, 2020)	Dataset yang digunakan dalam penelitian ini mencakup data penjualan dan pembelian di Apotek Suganda dari 1 Juni 2019 hingga 30 September 2019. Data tersebut meliputi informasi produk farmasi, pemasok, serta transaksi pesanan dan penjualan, yang penting dalam memperkirakan permintaan.
	Judul	Peramalan Penjualan Sediaan Farmasi Menggunakan Long Short-term Memory: Studi Kasus pada Apotik Suganda	
	Algoritma/ Metode	Long Short-term Memory (LSTM)	Hasil penelitian menunjukkan bahwa LSTM lebih unggul dibandingkan metode peramalan tradisional dalam memprediksi penjualan farmasi, dengan MAPE terendah sebesar 4.6151.
8	Peneliti (tahun)	(Abbasimehr, Shabani, & Yousefi, 2020)	Dataset yang digunakan dalam penelitian ini terdiri dari data kuantitas penjualan bulanan untuk produk tertentu dari perusahaan furnitur, mencakup periode dari 2007 hingga 2017, dengan total data 132 bulan.
	Judul	An optimized model using LSTM network for demand forecasting	
	Algoritma/ Metode	LSTM	Penelitian ini menyimpulkan bahwa metode peramalan LSTM multi-layer yang diusulkan mengungguli metode tradisional, mencapai nilai MAPE dan RMSE terendah dalam fase pelatihan dan pengujian. Metode ini mencatat uji MAPE 0,1023 dan RMSE 2172,4, peringkat pertama di antara semua metode yang diuji.
9	Peneliti (tahun)	(Ashari & Sadikin, 2020)	Dataset yang digunakan dalam penelitian ini berisi 603 titik data penjualan harian obat "X" dari PT. Metiska Farma selama periode 2017 hingga 2019. Setiap data mencakup dua atribut utama, yaitu tanggal transaksi dan nilai penjualan dalam rupiah.
	Judul	Prediksi Data Transaksi Penjualan	

No	Nama Peneliti/Journal		Hasil Penelitian
10		Time Series Menggunakan Regresi LSTM	Model terbaik mencapai RMSE 286.465.424 pada pelatihan dan 187.013.430 pada pengujian, sementara MAPE masing-masing sebesar 78,7% dan 30,9% menggunakan metode LSTM dalam memprediksi data penjualan.
	Algoritma/ Metode	LSTM	
	Peneliti (tahun)	(Wiranda & Sadikin, 2019)	Penelitian ini menyimpulkan bahwa LSTM efektif untuk memprediksi penjualan di industri farmasi, khususnya untuk PT. Peternakan Metika. Hasil membuktikan bahwa LSTM mencapai akurasi prediksi dengan RMSE 13.762.154,00 dan MAPE 12%. Temuan menunjukkan bahwa metode peramalan manual tradisional kurang efektif dibandingkan dengan teknik pembelajaran mesin
	Judul	Penerapan Long Short Term Memory Pada Data Time Series Untuk Memprediksi Penjualan Produk Pt. Metiska Farma	
	Algoritma/ Metode	LSTM	

Tabel 2.1 menyajikan penelitian terdahulu yang telah mengeksplorasi berbagai metode prediksi penyediaan stok produk, seperti LSTM, ANN, SVR, LightGBM, XGBoost, serta pendekatan lainnya. Namun banyak dari penelitian tersebut yang belum melakukan evaluasi yang komprehensif, sehingga masih terdapat kekurangan dalam analisis performa model. Oleh karena itu, penelitian ini berfokus pada penggunaan LSTM sebagai metode utama dengan pendekatan evaluasi yang lebih mendalam untuk meningkatkan akurasi dan keandalan model prediksi. Berdasarkan temuan penelitian sebelumnya, LSTM telah terbukti memiliki kemampuan yang baik dalam menangani data deret waktu, sehingga dipilih sebagai model yang akan dikembangkan lebih lanjut dalam konteks prediksi permintaan produk pada *vending machine*.