

Prediksi Ketersediaan Obat Penyakit Hipertensi Berbasis Data Historis Menggunakan Algoritma *Long Short-Term Memory* (LSTM)

1st Reihaini Fikria Bunga Oktaviani
Departemen Sistem Informasi
Fakultas Rekayasa Industri
Telkom University
Bandung, Indonesia
reihainifbo@student.telkomuniversity.ac.id

2nd Oktariani Nurul Pratiwi
Departemen Sistem Informasi
Fakultas Rekayasa Industri
Telkom University
Bandung, Indonesia
onurulp@telkomuniversity.ac.id

3rd Alfian Akbar Gozali
Departemen Rekayasa Perangkat Lunak
Fakultas Ilmu Terapan
Telkom University
Bandung, Indonesia
alfian@telkomuniversity.ac.id

Manajemen ketersediaan obat merupakan aspek penting dalam menjamin pelayanan kesehatan yang optimal. Namun, sering terkendala oleh ketidakseimbangan stok obat yang menyebabkan kerugian finansial, risiko obat kadaluarsa adanya ketidakpuasan pasien. Hal ini selaras dengan isu global terkait akses obat untuk penyakit kronis seperti Hipertensi. Penelitian ini bertujuan untuk mengembangkan model prediksi ketersediaan stok obat dengan pendekatan *deep learning*, menggunakan algoritma *Long Short-Term Memory* (LSTM), guna meningkatkan efektivitas pengelolaan persediaan obat di fasilitas kesehatan, khususnya untuk jenis-jenis obat yang sangat dibutuhkan. Penelitian ini menggunakan metodologi KDD (*Knowledge Discovery in Database*) dengan data transaksi obat historis dari Januari 2021 hingga Desember 2024 untuk diproses dan dimodelkan. Setelah itu, dari hasil pemodelan dilakukan evaluasi menggunakan metrik MAPE dan RMSE yang hasilnya menunjukkan kinerja bervariasi; model mampu memprediksi beberapa *brand* obat dengan akurasi yang baik sebanyak 5 *brand* (MAPE < 30%) dan menghadapi tantangan signifikan pada *brand* lain dengan *error* yang tinggi (MAPE > 50%). Meskipun demikian, model ini berpotensi meningkatkan efisiensi stok obat dan kepuasan pasien untuk obat-obatan yang terprediksi baik.

Kata Kunci: Ketersediaan Obat, Prediksi, KDD, *Deep Learning*, LSTM, Hipertensi

I. PENDAHULUAN

Manajemen ketersediaan obat adalah kunci utama dalam memastikan pelayanan kesehatan yang optimal. Pengelolaan obat yang optimal sangat dipengaruhi oleh ketersediaan sumber daya dalam sistem, dengan tujuan utama memastikan penyediaan obat yang berkualitas, terdistribusi secara merata, dan sesuai dengan jenis serta jumlah yang dibutuhkan oleh masyarakat di fasilitas kesehatan [1]. Namun, ada keadaan di mana stok obat di pelayanan kesehatan tidak seimbang, termasuk kelebihan dan kekurangan, yang dapat berdampak negatif dan merugikan pasien serta fasilitas kesehatan itu sendiri [2]. Tidak tersedianya obat yang dibutuhkan tentu saja mengecewakan pasien yang sangat membutuhkan obat tersebut. Sementara itu, melimpahnya obat yang kurang dibutuhkan akan menyebabkan kerugian finansial akibat obat yang kadaluarsa karena terlalu lama disimpan di gudang [3].

Isu ini juga relevan dengan isu kesehatan global, terutama dalam konteks Universal Health Coverage (UHC) yang menekankan akses adil terhadap obat berkualitas dan terjangkau. Banyak negara, termasuk Indonesia, masih menghadapi kendala besar, seperti kenaikan harga obat baru dan kelangkaan obat esensial, khusus untuk penyakit kronis seperti penyakit Hipertensi [4][5]. Masyarakat yang sakit dengan kondisi kronis sangat rentan, karena membutuhkan akses obat yang konsisten dan jangka panjang.

Untuk mengatasi masalah ketersediaan obat dibutuhkan sistem informasi terkomputerisasi yang mampu memprediksi tingkat stok secara optimal. Ini penting untuk mencegah baik kelebihan stok obat maupun kekurangan stok obat di fasilitas kesehatan. Prediksi merupakan proses sistematis untuk memperkirakan kejadian yang berpotensi terjadi di masa depan dengan mengandalkan data historis dan kondisi terkini sebagai dasar analisis. Tujuan utama dari proses ini adalah meminimalkan selisih antara hasil aktual dan hasil prediksi [6]. Salah satu metode canggih yang banyak digunakan proses prediksi adalah *deep learning*, karena kemampuannya dalam mengenali pola kompleks dari kumpulan data berukuran besar, sehingga dapat menghasilkan estimasi yang lebih akurat.

Deep learning menawarkan beragam algoritma yang dapat digunakan untuk prediksi di berbagai sektor, termasuk teknologi, kesehatan, dan industri lainnya. Di antara algoritma paling populer yang digunakan untuk prediksi adalah *Long Short-Term Memory*. Algoritma ini dirancang secara khusus untuk menangani data sekuensial dan mampu mengidentifikasi hubungan temporal antar data yang tersusun berdasarkan urutan waktu. LSTM memiliki keunggulan dalam mempertahankan informasi untuk periode waktu yang lebih lama [7].

Beberapa penelitian terbaru telah menunjukkan potensi besar algoritma LSTM untuk memprediksi ketersediaan obat dan menunjukkan hasil yang menjanjikan dalam meningkatkan manajemen obat. Misalnya, sebuah penelitian yang menerapkan LSTM untuk memprediksi ketersediaan obat di Apotek Suganda menunjukkan hasil yang baik dengan tingkat kesalahan prediksi yang rendah sekitar 4% menggunakan metrik *Mean Absolute Percentage Error* (MAPE) [9]. Penelitian lain mengevaluasi 30 model obat

yang dilatih menggunakan lingkungan berbasis GPU dan menemukan bahwa hanya 5 model yang memiliki akurasi di bawah 70%, menunjukkan kinerja keseluruhan yang kuat dari model LSTM di bidang kesehatan [10]. Secara keseluruhan, algoritma LSTM yang disajikan dalam kedua penelitian ini menunjukkan kinerja yang sangat baik dalam memprediksi ketersediaan obat dan memiliki potensi besar untuk prediksi yang lebih luas. Akurasi yang lebih baik dapat dicapai melalui penggunaan dataset yang lebih komprehensif dan berkualitas tinggi.

II. KAJIAN TEORI

A. Ketersediaan Obat

Obat merupakan zat atau senyawa yang dapat memengaruhi proses biologis dalam tubuh dan digunakan untuk tujuan pencegahan, pengobatan, diagnosis penyakit, atau untuk menimbulkan kondisi fisiologis tertentu [11]. Dalam konteks fasilitas pelayanan kesehatan, obat termasuk aset lancar yang memiliki peran penting dalam keberlangsungan intervensi medis, mengingat lebih dari 90% layanan kesehatan melibatkan penggunaan obat. Oleh karena itu, ketersediaan obat menjadi indikator penting yang mencerminkan kualitas pelayanan kesehatan.

Terjadinya kekosongan obat, kehabisan stok, maupun penumpukan obat yang berlebihan tidak hanya menimbulkan dampak medis bagi pasien, tetapi juga menimbulkan konsekuensi ekonomi bagi institusi pelayanan kesehatan. Situasi ini menuntut adanya sistem pengelolaan obat yang efisien dan efektif [12]. Permasalahan tersebut kerap kali disebabkan oleh perencanaan yang kurang tepat serta lemahnya manajemen logistik, yang dapat mengakibatkan distribusi obat tidak akurat apabila tidak mempertimbangkan data tren historis maupun kebutuhan penyakit spesifik di tiap wilayah [13]. Oleh karena itu, perencanaan dan pengelolaan ketersediaan obat yang berbasis data menjadi sangat penting guna mencegah terjadinya kekurangan maupun kelebihan stok di fasilitas kesehatan.

B. Prediksi

Prediksi merupakan suatu proses sistematis yang digunakan untuk memperkirakan kejadian yang paling mungkin terjadi di masa depan dengan memanfaatkan informasi historis dan kondisi saat ini. Dalam konteks pengambilan keputusan, prediksi berperan penting dalam mengidentifikasi tren baru, memperkirakan kondisi mendatang, serta menyediakan dasar pertimbangan yang lebih objektif dan terukur [6].

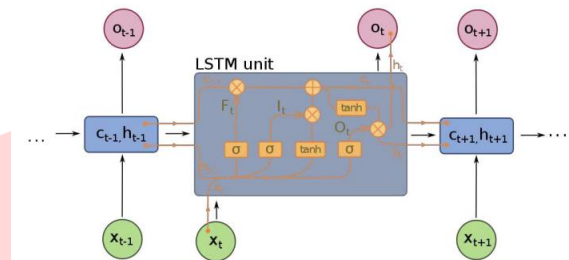
Salah satu pendekatan modern yang banyak digunakan dalam proses prediksi adalah *deep learning*. Teknologi ini memiliki kemampuan untuk mengenali pola yang kompleks dari kumpulan data berukuran besar dan menghasilkan prediksi dengan tingkat akurasi yang tinggi. Melalui pelatihan algoritma, model *deep learning* dapat dioptimalkan untuk menyelesaikan berbagai permasalahan spesifik sesuai dengan karakteristik data yang dianalisis [14].

C. Algoritma Long Short-Term Memory (LSTM)

Algoritma *Long Short-Term Memory* (LSTM) pertama kali diperkenalkan oleh Hochreiter dan Schmidhuber pada tahun 1997. LSTM merupakan pengembangan dari arsitektur *Recurrent Neural Network* (RNN), yang dirancang untuk

memproses data berurutan dengan mempertimbangkan informasi dari langkah waktu sebelumnya. RNN sendiri menggunakan *multi-layer* dan mekanisme perulangan dalam menangani data *time series*, sehingga termasuk dalam kategori algoritma *deep learning* [15].

Keunggulan utama LSTM terletak pada kemampuannya dalam mengatasi masalah ketergantungan jangka panjang (*long term dependencies*), yaitu kemampuan untuk mempertahankan informasi penting dari langkah-langkah sebelumnya dalam rangkaian data yang panjang [8].



GAMBAR 1
Arsitektur LSTM [16]

Gambar tersebut menggambarkan arsitektur LSTM yang dirancang untuk memproses data *time series*. Struktur ini terdiri atas unit-unit neuron yang berperan sebagai *cell state* (c_t) serta tiga jenis *gate* utama yang mengatur aliran informasi. Ketiga *gate* tersebut berfungsi untuk memproses dan menyaring informasi berasal dari *input* saat ini (X_t), *cell state* sebelumnya (c_{t-1}) dan *hidden state* sebelumnya (h_{t-1}). Mekanisme ini secara langsung memengaruhi pembentukan *hidden state output* (h_t) yang merupakan representasi akhir dari unit LSTM pada waktu tertentu. Interaksi antar komponen ini memungkinkan LSTM untuk mempertahankan informasi penting sekaligus menyaring informasi yang tidak relevan dalam rangkaian data yang panjang [16]. Berikut penjelasan untuk setiap *gate* pada satu sel memori LSTM [8].

1. Forget Gate (f_t)

Forget gate berfungsi untuk menentukan seberapa banyak informasi dari *cell state* sebelumnya yang akan dipertahankan atau dihapus dalam memori sel. *Gate* ini adalah lapisan *sigmoid* yang mengambil *output* pada h_{t-1} dan input X_t kemudian menerapkannya pada fungsi aktivasi *sigmoid*. *Output* dari *gate* ini berada pada rentang 0 atau 1. Jika $f_t = 0$ maka informasi sebelumnya akan dilupakan, sementara jika $f_t = 1$ maka informasi sebelumnya dipertahankan.

Rumus dari f_t adalah:

$$f_t = \sigma(W_f S_{t-1} + W_f X_t) \quad (1)$$

Dengan:

W_f : Bobot dari *forget gate*.

S_{t-1} : *State* sebelumnya atau *state* pada waktu $t-1$.

X_t : *Input* pada waktu t (sekarang)

σ : Fungsi aktivasi *sigmoid*.

2. Input Gate (i_t)

Input gate berfungsi untuk mengatur sejauh mana informasi baru akan disimpan ke dalam *cell state*. Mekanisme ini mencegah penyimpanan data yang tidak relevan dengan cara menyaring informasi melalui fungsi aktivasi *sigmoid*. Nilai *output* dari *Gate* ini berperan

mengambil berada dalam rentang 0 atau 1, yang merepresentasikan tingkat kontribusi informasi baru terhadap pembaruan *cell state*.

Rumus dari i_t adalah:

$$i_t = \sigma(W_i S_{t-1} + W_i X_t) \quad (2)$$

Dengan:

W_i : Bobot dari *input gate*.

S_{t-1} : *State* sebelumnya atau *state* pada waktu t-1.

X_t : *Input* pada waktu t (sekarang)

σ : Fungsi aktivasi sigmoid.

3. Output Gate (o_t)

Output gate berfungsi dalam menentukan seberapa besar informasi dari memori sel yang akan diteruskan sebagai *output*. *Gate* ini mengontrol seberapa banyak *state* yang lewat ke *output* dan bekerja dengan cara yang sama dengan *gate* lainnya.

Rumus dari o_t adalah:

$$o_t = \sigma(W_o S_{t-1} + W_o X_t) \quad (3)$$

Dengan,

W_o : Bobot dari *output gate*.

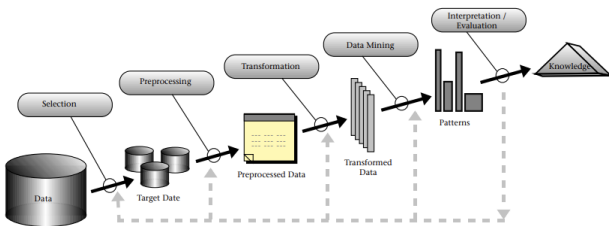
S_{t-1} : *State* sebelumnya atau *state* pada waktu t-1.

X_t : *Input* pada waktu t (sekarang)

σ : Fungsi aktivasi sigmoid.

III. METODOLOGI PENELITIAN

Metodologi yang digunakan pada penelitian ini adalah KDD (*Knowledge Discovery*).



GAMBAR 2
Tahapan KDD [17]

A. Selection

Tahap awal ini melibatkan penentuan data yang relevan dan fokus pada subset atribut yang akan digunakan untuk penemuan pengetahuan.

B. Pre-processing

Tahap ini berfokus pada peningkatan kualitas data seperti melakukan *data cleaning*, penghapusan *missing values*, dan *outlier*. Tujuannya adalah untuk memastikan data bersih dan dapat diandalkan.

C. Transformation

Tahap ini data diubah ke dalam format yang lebih tepat dan terstruktur agar dapat diproses secara optimal dalam tahap *data mining*. Tujuannya untuk memastikan bahwa data berada dalam bentuk yang sesuai dengan kebutuhan analisis dan karakteristik algoritma yang digunakan.

D. Data Mining

Tahap ini adalah tahapan penting dengan menerapkan algoritma yang paling sesuai dengan karakteristik data dan tujuan analisis, seperti pada penelitian ini yang menerapkan

algoritma LSTM, proses ini mencakup penyetelan parameter model guna mengoptimalkan kinerjanya.

E. Interpretation / Evaluation

Tahap ini merupakan hasil dari *data mining* untuk menilai kualitas, keakuratan, dan relevansi pengetahuan yang ditemukan berdasarkan metrik evaluasi tertentu.

IV. HASIL DAN PEMBAHASAN

A. Selection

Tahap awal ini melibatkan penentuan data yang relevan dan fokus pada subset atribut yang akan digunakan untuk penemuan pengetahuan.

Data yang digunakan dalam penelitian ini adalah data historis ketersediaan obat dari Yayasan Kesehatan Swasta Kota Bandung, yang mencatat informasi terkait aktivitas pemberian resep obat kepada pasien dalam rentang waktu dari Januari tahun 2021 hingga Desember tahun 2024. Dataset ini memiliki 13 atribut, meliputi nomor resep, kode apotek, tanggal transaksi, ID pasien, jenis kelamin, umur, ICD10, nama penyakit, *brand*, *generic*, sediaan, jumlah, dan kemasan. Namun, setelah melakukan eksplorasi dan analisis awal, penelitian ini memutuskan hanya menggunakan atribut tanggal resep, ICD10, *brand*, dan jumlah dalam proses pengelolaan dan pemodelan data lebih lanjut karena relevansinya yang paling tinggi dengan tujuan prediksi ketersediaan obat.

TABEL 1
Atribut Data

Tanggal	ICD10	Brand	Jumlah
2021-10-27	I10	Lifezar tab. 50 mg	30
2022-08-09	G62	Neurohax 5000 tab.	14
2023-01-02	K04.0	Methylprednisolone tab. 4 mg	10
2023-12-27	E11	Metformin tab. 500 mg	40
2024-07-19	N18	Mecobalamin tab. 500 mg	84

B. Pre-processing

Tahap ini berfokus pada peningkatan kualitas data. Tujuannya adalah untuk memastikan data bersih dan dapat diandalkan.

Proses *pre-processing* data dilanjutkan dengan standarisasi penamaan kolom setelah seluruh data berhasil digabungkan. Anomali penulisan nama kolom, seperti penggunaan kapitalisasi yang tidak seragam (misalnya, 'Tanggal' dengan 'brand'), ditemukan di beberapa kolom.

Dalam konteks penelitian ketersediaan obat ini, kolom "*brand*" mempresentasikan jenis atau nama dagang obat yang menjadi titik krusial karena secara langsung memengaruhi validitas hasil analisis. Sehingga perlu adanya pembersihan data (*data cleaning*) pada kolom ini untuk memastikan konsistensi penulisan nama-nama *brand* ke dalam format huruf kecil (*lowercase*).

TABEL 2
Konversi Nama Brand

Brand	Brand_new
Lifesar tab. 50 mg	lifezartab50mg
Candepress tab. 8 mg	candepresstab8mg
Cefadroxil kaps. 500 mg (Ber)	cefadroxilkaps500mgber
Kolkatriol tab. 0,25 m	kolkatrioltab025mg
Amlodipine tab. 5 mg (Kim)	amlodipinetab5mgkim

Selain proses tersebut, tahapan *pre-processing* juga mencakup identifikasi dan penanganan data duplikat. Pengecekan ini dilakukan untuk memastikan bahwa setiap entri dalam dataset bersifat unik dan tidak berulang. Data duplikat ini umumnya muncul akibat pencatatan ganda, kesalahan *input*, atau proses penggabungan data dari berbagai sumber tanpa normalisasi yang memadai. Apabila tidak ditangani, keberadaan duplikasi ini dapat menyebabkan pembelajaran model menjadi tidak akurat karena bobot informasi menjadi tidak seimbang.

Untuk memastikan relevansi data dengan fokus penelitian, dilakukan pemfilteran dataset secara spesifik untuk kasus penyakit Hipertensi (I10), sehingga hanya data obat yang berhubungan dengan diagnosis ini yang akan diproses lebih lanjut. Selain itu, pada tahap ini, analisis lebih lanjut akan difokuskan pada 50 *brand* obat teratas (*top 50 drugs*), kriteria ini dipilih karena merefleksikan obat-obat yang paling aktif dan relevan dalam operasional fasilitas kesehatan, sehingga hasil pemodelan akan lebih berdampak secara praktis.

Kemudian, dilakukan tahapan *handling missing value* dan penghapusan *outlier* untuk memastikan kualitas dan kontinuitas data. Proses ini krusial untuk menjamin setiap *brand* memiliki *entry* data untuk setiap tanggal dalam rentang observasi, bahkan jika tidak ada transaksi. Selanjutnya, transaksi akan dijumlahkan per tanggal untuk setiap *brand*. Proses pembersihan ini menjadi krusial karena tanpa penghapusan *outlier* dan penataan ulang struktur data, model LSTM tidak akan mampu bekerja secara optimal. LSTM menuntut data yang kontinu dan terstruktur dengan baik agar dapat mempelajari pola temporal secara akurat.

C. Transformation

Tahap ini data ditransformasikan ke dalam format yang lebih sesuai agar dapat diproses secara efektif dalam tahap *data mining*.

Analisis lebih lanjut pada penelitian ini akan difokuskan pada 50 *brand* obat teratas (*top 50 drugs*). Pemilihan “50 *brand* obat teratas” didasarkan pada hasil observasi awal terhadap distribusi jumlah transaksi obat. Ditemukan bahwa sebagian besar *brand* hanya memiliki satu kali transaksi (jumlah = 1), di mana model LSTM memiliki persyaratan minimum terhadap panjang urutan waktu (*timesteps*) agar mampu menangkap pola temporal. Data dengan hanya satu transaksi tidak memenuhi kebutuhan ini, karena perhitungan *timesteps* pada LSTM setidaknya memerlukan dua titik data atau lebih (*timesteps* > 1) untuk membangun dependensi waktu. Oleh karena itu, *brand* obat dengan jumlah transaksi yang sangat minim tidak dapat diuji secara efektif oleh model ini.

Adapun pemilihan 50 *brand* obat teratas dilakukan dengan mengurutkan data berdasarkan frekuensi transaksi tertinggi selama periode observasi. Kriteria ini dipilih karena

dianggap mewakili obat-obat yang paling sering digunakan dalam praktik klinis dan operasional fasilitas kesehatan. Dengan demikian, model yang dibangun akan lebih fokus pada entitas yang relevan dan berdampak secara praktis, serta memungkinkan pembelajaran pola temporal yang lebih kuat dan representatif.

Selanjutnya, untuk mempersiapkan data *brand* obat ini agar dapat diproses oleh model LSTM, dilakukan tahap *label encoding*. Proses ini mengubah representasi label kategorikal (nama *brand* dalam format teks, contoh: ‘cebextab’) menjadi nilai numerik diskrit (angka integer). Hal ini penting karena sebagian besar algoritma *deep learning* tidak dapat memproses *input* berupa teks secara langsung. Dengan *label encoding*, setiap nama *brand* yang unik akan diberikan sebuah angka integer unik (misalnya, ‘cebextab’ menjadi ‘0’, lalu ‘amlodipine’ menjadi ‘1’, dan seterusnya), yang kemudian disimpan dalam kolom baru bernama ‘label *brand*’.

Tabel berikut menunjukkan 50 *brand* obat teratas beserta nilai *label encoding* yang sesuai.

TABEL 3
Dataset 50 Brand Teratas dan Hasil Label Encoding

Label Brand	Brand
0	Adalat tab. 30 mg, oros
1	Allopurinol tab. 100 mg (Hex)
2	Amcor tab. 10 mg
3	Amcor tab. 5 mg
4	Amlodipine tab. 10 mg (Pro)
5	Amlodipine tab. 5 mg (Dex)
6	Amlodipine tab. 5 mg (Ind)
...	...
49	Voltadex emulgel 1%/20 gr

D. Data Mining

Tahapan ini merupakan proses inti dari penelitian yang berfungsi untuk membangun model prediksi menggunakan algoritma LSTM. Tahap pemodelan dimulai dengan persiapan data historis transaksi obat yang telah melalui proses pra-pemrosesan di mana data di filter untuk mendapatkan 50 *brand* obat teratas yang paling relevan. Kemudian, data dibagi menjadi data latih (80%) dan data uji (20%) secara berbasis waktu, memastikan model dievaluasi pada data yang belum pernah dilihat sebelumnya.

Selanjutnya, normalisasi data diterapkan pada kolom ‘jumlah’ menggunakan MinMaxScaler baik untuk data latih maupun data uji dengan tujuan menskalakan fitur-fitur pada rentang yang seragam. Data yang telah dinormalisasi diubah menjadi sekuens input dengan panjang *timesteps* 7, mewakili 7 langkah waktu historis, dengan nilai transaksi berikutnya sebagai target prediksi. Sekuens ini di-*reshape* dalam bentuk tiga dimensi yang sesuai untuk *input layer* mode.

Arsitektur model prediksi dibangun secara sekuensial menggunakan algoritma LSTM, yang terbukti efektif untuk data *time series*. Model ini terdiri dari beberapa *layer* yang dirancang untuk menangkap pola temporal dan menghasilkan prediksi.

Proses pelatihan model dilakukan menggunakan fungsi *model.fit()* dengan data pelatihan berupa pasangan *input* (X_train) dan label (y_train). Pelatihan dijalankan selama 100 *epoch* dengan ukuran *batch* sebesar 32, di mana pada

setiap iterasi model memproses 32 sampel sebelum memperbarui bobot. Setelah pelatihan, model memproyeksikan hasil pembelajaran menjadi *output* prediksi berupa jumlah transaksi. Selama proses pelatihan, digunakan juga data validasi (*X_test* dan *y_test*) yang berfungsi untuk memantau performa model di luar data pelatihan dan mendeteksi potensi *overfitting* secara dini.

Output dari *layer* LSTM terakhir tidak langsung mewakili prediksi, namun terlebih dahulu diolah oleh *layer Dense* sebagai lapisan akhir yang bertugas mengubah representasi hasil pembelajaran menjadi nilai prediksi.

Berikut susunan *layer* secara lengkap.

TABEL 4
Susunan Layer

Layer (Type)	Output Shape
lstm (LSTM)	(None, 7, 128)
dropout (Dropout)	(None, 7, 128)
lstm_1	(None, 64)
dropout_1 (Dropout)	(None, 64)
dense (Dense)	(None, 1)

Arsitektur ini bekerja dengan cara memproses urutan data transaksi dalam bentuk sekuens historis, di mana setiap langkah waktu mengandung nilai numerik hasil normalisasi dari jumlah transaksi per *brand* obat. Proses pembelajaran dilakukan langsung terhadap pola-pola temporal yang muncul dalam data historis tersebut.

E. Evaluation

Penelitian ini menggunakan dua metrik evaluasi, yaitu *Mean Absolute Percentage Error* (MAPE) dan *Root Mean Square Error* (RMSE), untuk menilai kinerja model prediksi [18]. Tabel 5 menunjukkan hasil evaluasi kinerja model LSTM dengan hasil yang *brand* yang baik.

TABEL 5
Evaluasi Kinerja Model Baik

Label Brand	Brand	MAPE (%)	RMSE
1	Allopurinol tab. 100 mg (Hex)	19.2%	57.71
4	Amlodipine tab. 10 mg (Pro)	14.2%	101.49
5	Amlodipine tab. 5 mg (Dex)	16.9%	125.54
6	Amlodipine tab. 5 mg (Ind)	19.7%	67.9
21	Candesartan tab. 8 mg. (Dex)	17.4%	128.1

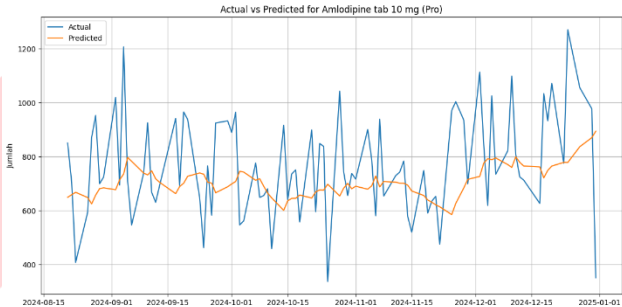
Lalu, Tabel 6 untuk hasil evaluasi kinerja model LSTM dengan hasil *brand* yang buruk.

TABEL 6
Evaluasi Kinerja Model Buruk

Label Brand	Brand	MAPE (%)	RMSE
24	Esvat tab. 10 mg	91.9%	30.1
27	Hexavask tab. 10 mg	90.8%	63.23
37	Oxyvit tab.	91.7%	31.78

Berdasarkan hasil evaluasi menggunakan metrik MAPE dan RMSE yang ditunjukkan tabel-tabel diatas, terlihat adanya variasi performa model antar *brand* obat yang berbeda. Secara spesifik, model menunjukkan kinerja yang baik pada beberapa *brand* saja, di antaranya *brand* 5, *brand* 6, *brand* 21, dan *brand* 39 yang mencapai nilai MAPE kurang dari 30%. Sebaliknya, terdapat beberapa *brand* yang menunjukkan kinerja model kurang optimal, yaitu *brand* 24, 27, dan 37 dengan nilai MAPE yang tinggi hingga 90%.

Berikut gambar yang menunjukkan perbandingan antara nilai aktual dan nilai hasil prediksi model LSTM untuk salah satu *brand* obat dengan kinerja terbaik, yaitu *Amlodipine tab. 10 mg (Pro)*.



GAMBAR 3
Visualiasi Perbandingan Nilai Prediksi dan Nilai Aktual dari Brand Amlodipine tab. 10 mg (Pro)

Garis prediksi model berhasil menangkap pergerakan umum data, termasuk tren kenaikan volume menuju akhir tahun 2024. Selain itu, model ini cukup responsif terhadap perubahan arah data, garis prediksi berusaha untuk naik atau turun sesuai dengan data aktual, tidak terlalu datar. Ini mengindikasikan model belajar sebagian dari dinamika pola. Jarak antara garis aktual dan prediksi terlihat moderat (cukup baik) karena *magnitude* kesalahan tidak terlalu ekstrem, namun konsisten terjadi akibat efek penghalusan model.

Meskipun model merespons arah perubahan, model masih kesulitan dalam menangkap *amplitude* sebenarnya dari puncak-puncak tertinggi dan lembah-lembah terendah. Prediksi cenderung lebih halus, tidak mencapai nilai setinggi puncak aktualnya

F. Implikasi Hasil Penelitian

Berikut adalah nilai prediksi untuk 6 bulan ke depan dari 5 *brand* yang memiliki hasil evaluasi yang baik.

TABEL 7
Prediksi untuk 6 Bulan ke Depan

Tanggal	Brand 1	Brand 4	Brand 5	Brand 6	Brand 21
2025-01	4621	616	635	1417	408
2025-02	4196	598	534	1542	388
2025-03	4646	596	518	1548	388
2025-04	4496	595	515	1548	388
2025-05	4646	595	514	1548	388
2025-06	4496	595	514	1548	388

Pada Tabel 8, hasil prediksi untuk 6 bulan ke depan (Januari hingga Juni 2025) menunjukkan pola yang stabil dan mendatar setelah kenaikan awal bulan. Hal ini mengindikasikan bahwa kecenderungan model untuk

menghasilkan prediksi yang konstan merupakan karakteristik yang umum jika tidak ada informasi dinamis baru yang memicu fluktuasi.

V. KESIMPULAN

Penelitian ini bertujuan untuk memprediksi ketersediaan stok obat di Yayasan Kesehatan Swasta Kota Bandung dari rentang Januari tahun 2021 hingga Desember tahun 2024. Hasil evaluasi menggunakan metrik MAPE dan RMSE menunjukkan kinerja model yang bervariasi. Model berhasil memprediksi kebutuhan stok 5 *brand* obat dengan akurasi baik (MAPE dibawah 30%). Hasil prediksi yang memiliki banyak model stabil ini memiliki dampak penting untuk perencanaan strategis ketersediaan obat, di mana prediksi memberikan estimasi jumlah yang dapat diandalkan. Namun, ada keterbatasan dalam kemampuan model untuk memprediksikan fluktuasi yang ada pada data karena terlihat masih adanya *brand* obat dengan akurasi yang rendah pada evaluasi. Meskipun nilai MAPE menunjukkan kategori yang cukup dan nilai rata-rata RMSE yang baik, perlu dicatat bahwa masih terdapat sebagian besar *brand* obat yang menunjukkan nilai MAPE di atas 50% bahkan mencapai hingga 91.8%. Artinya, model secara signifikan masih kesulitan dalam memprediksi sebagai besar *brand* obat.

Saran untuk penelitian selanjutnya yaitu investigasi lebih lanjut faktor-faktor penyebab tingginya *error* pada *brand* tertentu, seperti eksplorasi fitur data tambahan, optimasi *hyperparameter* yang lebih mendalam, atau penerapan arsitektur model yang lebih kompleks untuk menangani data yang lebih bervariasi, guna mencapai akurasi yang lebih merata di seluruh jenis obat.

DAFTAR PUSTAKA

- [1] Y. H. Irawan, N. A. Rostikarina, and Y. Rahmawati, "Kajian Literatur Pengelolaan Obat Di Rumah Sakit," *ASSYIFA J. Ilmu Kesehat.*, vol. 2, no. 2, pp. 336–342, 2024, doi: 10.62085/ajk.v2i2.100.
- [2] I. P. San, S. B. Andi, and K. A. Muh, "Pengelolaan Kebutuhan Logistik Farmasi pada Instalasi Farmasi RS Islam Faisal Makassar Pharmaceutical Logistics Management of The Pharmacy Installation , Faisal Islamic Hospital Makassar," *Promot. J. Kesehat. Masy.*, vol. 10, no. 02, pp. 78–85, 2020.
- [3] E. D. Madyatmadja, L. Kusumawati, S. P. Jamil, W. Kusumawardhana, S. Informasi, and U. B. Nusantara, "Infotech: journal of technology information," *Raden Ario Damar*, vol. 7, no. 1, pp. 55–62, 2021.
- [4] World Health Organization, "Access to Medicines and Health Products." Accessed: May 25, 2025. [Online]. Available: <https://www.who.int/our-work/access-to-medicines-and-health-products>
- [5] D. Ferdinal, I. Nursukmi, and R. R. Putra, "Prediksi Obat Kronis Penyakit Diabetes Melitus Menggunakan Metode Monte Carlo," *J. Komput. Teknol. Inf. dan Sist. Inf.*, vol. 3, no. 1, pp. 665–672, 2024, doi: 10.62712/juktisi.v3i1.182.
- [6] L. Ainiyah and M. Bansori, "Prediksi Jumlah Kasus COVID-19 Menggunakan Metode Autoregressive Integrated Moving Average," *J. Sains Dasar*, vol. 10, no. 2, pp. 62–68, 2021.
- [7] J. Nurhakiki *et al.*, "Studi Kepustakaan: Pengenalan 4 Algoritma Pada Pembelajaran Deep Learning Berserta Implikasinya," *J. Pendidik. Berkarakter*, no. 1, pp. 270–281, 2024, [Online]. Available: <https://doi.org/10.51903/pendekar.v2i1.598>
- [8] L. Wiranda and M. Sadikin, "Penerapan Long Short Term Memory pada Data Time Series untuk Memprediksi Penjualan Produk PT. Metiska Farma," *J. Nas. Pendidik. Tek. Inform.*, vol. 8, no. 3, pp. 184–196, 2019.
- [9] M. I. Anshory, Y. Priyandari, and Y. Yuniaristanto, "Peramalan Penjualan Sediaan Farmasi Menggunakan Long Short-term Memory: Studi Kasus pada Apotik Suganda," *Performa Media Ilm. Tek. Ind.*, vol. 19, no. 2, pp. 159–174, 2020, doi: 10.20961/performa.19.2.45962.
- [10] M. Khadafi and O. Yusrianti, "Sistem Prediksi Kebutuhan Stock Obat Menggunakan Metode Long Short Term Memory (LSTM) Time Series Berbasis Deep Learning A-26," vol. 8, no. 1, pp. 26–33, 2025.
- [11] W. L. Prabowo, "Teori Tentang Pengetahuan Peresepan Obat," *J. Med. hutama*, vol. 02, no. 04, pp. 402–406, 2021.
- [12] T. M. Ramzi, R. A. Dakhi, A. Sirait, D. Nababan, and E. Sembiring, "Analisis Manajemen Logistik Obat Di Instalasi Farmasi Rumah Sakit Umum Haji Medan," *J. Kesehat. Masy.*, vol. 7, no. 3, pp. 16838–16852, 2023, [Online]. Available: <https://ejournal.unsrat.ac.id/index.php/jikmu/article/view/7853>
- [13] A. M. Ulfa and D. Chalidyanto, "Evaluasi Proses Manajemen Logistik Obat di UPTD Puskesmas Kabupaten Sampang," *Media Gizi Kesmas*, vol. 10, no. 2, p. 196, 2021, doi: 10.20473/mgk.v10i2.2021.196-204.
- [14] Muhammad Haris Diponegoro, Sri Suning Kusumawardani, and Indriana Hidayah, "Tinjauan Pustaka Sistematis: Implementasi Metode Deep Learning pada Prediksi Kinerja Murid," *J. Nas. Tek. Elektro dan Teknol. Inf.*, vol. 10, no. 2, pp. 131–138, 2021, doi: 10.22146/jnteti.v10i2.1417.
- [15] F. Yanti, B. Nurina Sari, and S. Defiyanti, "Implementasi Algoritma Lstm Pada Peramalan Stok Obat (Studi Kasus: Puskesmas Beber)," *J. Mhs. Tek. Inform.*, vol. 8, no. 4, pp. 6082–6089, 2024.
- [16] M. K. Wisyaldin, G. M. Luciana, and H. Pariaman, "Pendekatan LSTM untuk Memprediksi Kondisi Motor 10 kV pada PLTU Batubara," *Kilat*, vol. 9, no. 2, pp. 311–318, 2020, [Online]. Available: <http://jurnal.itpln.ac.id/kilat/article/view/997%0Ahttp://jurnal.itpln.ac.id/kilat/article/download/997/775>
- [17] U. Fayyad, G. Piatesky-Shapiro, and P. Smyth, "Data Mining and Knowledge Discovery in Databases," *Commun. ACM*, vol. 39, no. 11, pp. 24–26, 1996, doi: 10.1145/240455.240463.
- [18] M. M. Azman, "Analisa perbandingan nilai akurasi moving average dan exponential smoothing untuk sistem peramalan pendapatan pada perusahaan XYZ," *J. Sist. dan Inform.*, vol. 13, no. 2, pp. 36–45, 2019.