

Comparison of the Accuracy Levels of Naive Bayes, Random Forest, and Long Short-Term Memory (LSTM) Methods in Predicting Gold Jewelry Sales

Muhammad Arfianto Pandu W^{*1}, Rujianto Eko Saputro², Purwadi³, Umdah Aulia Rohmah⁴

^{1,2,3}Magister Ilmu Komputer, Fakultas Ilmu Komputer, Universitas Amikom Purwokerto, Indonesia

⁴Ekonomi Syariah, Fakultas Ekonomi dan Bisnis Islam, Universitas Islam Negeri Prof.K.H.Saifuddin Zuhri, Purwokerto, Indonesia

Email: ¹muhammad.arfianto92@gmail.com

Received : Jul 18, 2025; Revised : Jul 28, 2025; Accepted : Aug 3, 2025; Published : Feb 15, 2026

Abstract

Gold has long been recognized as a safe haven asset, especially during economic uncertainty. Accurate prediction of gold jewelry sales is essential for inventory management and business strategy, particularly in high-demand regions such as Imogiri. This study aims to compare the accuracy levels of three machine learning methods—Naïve Bayes, Random Forest, and Long Short-Term Memory (LSTM)—in predicting gold jewelry sales using historical transaction data from Toko Emas Parimas. The dataset comprises 4,595 records from January 2022 to December 2024. The research employs data preprocessing, including data cleaning, feature transformation, and normalization, followed by classification into sales categories. Two data-splitting schemes (80:20 and 70:30) were implemented to evaluate model generalization. The models were trained and tested using performance metrics such as accuracy, precision, recall, and F1-score. The results show that Random Forest achieved perfect classification with an accuracy of 1.00 in both schemes, outperforming the other models. Naïve Bayes also performed well with accuracy up to 0.98, while LSTM showed moderate results with accuracy ranging from 0.82 to 0.88. These findings indicate that Random Forest is the most reliable model for sales prediction of gold jewelry, especially for static classification tasks. The study provides practical insights for retailers and decision-makers in selecting suitable analytical models, and it highlights the importance of aligning analytical methods with data characteristics to improve decision support systems in retail.

Keywords : *Gold Jewelry, Long Short-Term Memory (LSTM), Method Comparison, Model Accuracy, Naive Bayes, Random Forest, Sales Prediction*

This work is an open access article and licensed under a Creative Commons Attribution-Non Commercial 4.0 International License



1. PENDAHULUAN

Emas dikenal sebagai instrumen investasi dengan nilai intrinsik tinggi, yang berfungsi sebagai lindung nilai terhadap inflasi dan ketidakpastian ekonomi global[1][2]. Stabilitas nilai emas dan korelasinya yang rendah terhadap aset berisiko[3] [4] menjadikannya pilihan utama dalam diversifikasi portofolio [5] dan pelestarian kekayaan jangka panjang[6][7]. Namun, harga dan penjualan emas sangat dipengaruhi oleh faktor-faktor ekonomi makro seperti nilai tukar, suku bunga, dan ketegangan geopolitik[8] [9] Oleh karena itu, kemampuan memprediksi penjualan emas secara akurat menjadi krusial bagi investor dan pelaku pasar dalam menyusun strategi investasi berbasis data[10]. Tren permintaan emas yang meningkat juga terlihat di daerah seperti Imogiri, di mana Toko Parimas mencerminkan tingginya dinamika pasar. Hal ini menuntut strategi pengelolaan stok yang cermat dari segi jumlah dan berat emas yang dijual.

Untuk mendukung pengambilan keputusan yang lebih akurat, Toko Parimas memiliki kebutuhan untuk mengetahui prediksi penjualan emas di masa mendatang. Prediksi ini meliputi estimasi jumlah item emas dan total gram yang akan terjual dalam periode tertentu. Dengan adanya informasi prediktif yang akurat, pihak toko dapat mengelola persediaan secara lebih efisien, meminimalkan risiko kekurangan atau kelebihan stok, serta meningkatkan kepuasan pelanggan. Dalam konteks ini, pemanfaatan teknologi

machine learning menjadi solusi potensial untuk menghasilkan model prediksi yang andal dan berbasis data historis penjualan toko.

Seiring berkembangnya teknologi analisis data, metode prediksi berbasis machine learning dan deep learning telah banyak digunakan untuk mengolah data historis dan mengidentifikasi pola dalam penjualan emas. Di antara metode yang sering digunakan adalah Naive Bayes, Random Forest, dan Long Short-Term Memory (LSTM). Naive Bayes merupakan algoritma berbasis probabilistik yang sederhana namun cepat dalam proses klasifikasi[11][12][13]. Sementara itu, Random Forest sebagai metode ensemble berbasis decision tree menawarkan kemampuan generalisasi yang baik dalam menangani data kompleks[14][15][16][17]. Di sisi lain, LSTM, sebagai model jaringan saraf tiruan yang dirancang khusus untuk menangani data time series, memiliki keunggulan dalam mengenali pola temporal dan ketergantungan jangka panjang dalam data[18][19][20].

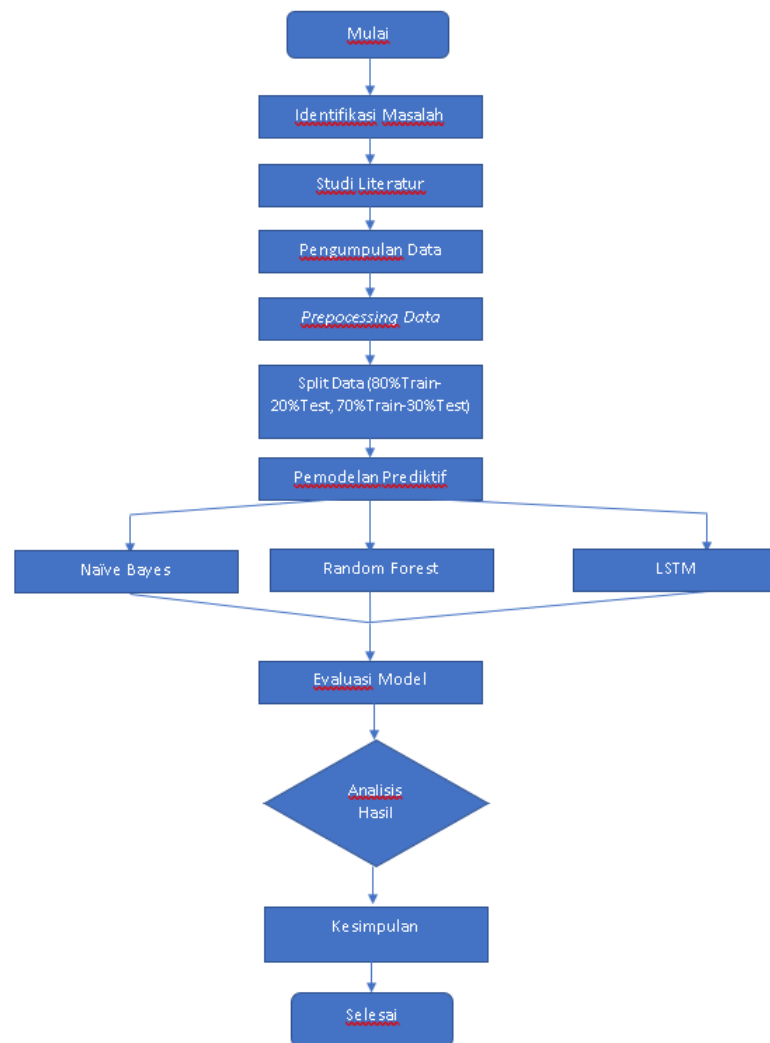
Prediksi penjualan emas dan harga emas memiliki peran penting dalam merencanakan investasi dan strategi bisnis di sektor pertambangan, terutama dalam menghadapi volatilitas pasar yang tinggi[21]. Sejumlah penelitian sebelumnya telah dilakukan untuk meramalkan perilaku penjualan emas menggunakan pendekatan algoritma machine learning. Di antara metode yang telah digunakan, Naive Bayes dan Random Forest menjadi dua pendekatan yang cukup sering diaplikasikan karena keunggulannya masing-masing dalam konteks klasifikasi dan peramalan. Sebagai metode ensemble yang menggabungkan banyak pohon keputusan, Random Forest mampu menangani data yang lebih variatif dan kompleks, serta mengurangi risiko overfitting. Keunggulannya yang terletak pada kemampuan untuk menangani interaksi antar variabel yang lebih rumit membuatnya efektif dalam meramalkan pembelian berkelanjutan, termasuk dalam sektor penjualan perhiasan emas. Dengan tingkat akurasi yang tinggi, Random Forest telah menunjukkan bahwa model ini dapat diandalkan untuk mengklasifikasikan pelanggan yang berpotensi melakukan repeat order, serta memberikan wawasan penting untuk strategi pemasaran yang lebih terarah[22]. Naive Bayes, yang berbasis pada prinsip probabilistik, dikenal efektif dalam mengidentifikasi peluang produk yang diminati konsumen. Beberapa studi menunjukkan bahwa algoritma ini dapat membantu perusahaan dalam mengoptimalkan strategi pemasaran dan memperkuat posisi pasar, terutama karena kesederhanaannya dalam menangani data kategori dan kecepatan komputasi yang tinggi [23][24]. Sedangkan Long Short-Term Memory (LSTM) menonjol sebagai salah satu model yang paling andal dalam menangani data deret waktu yang kompleks dan bersifat dinamis. LSTM dirancang khusus untuk mengatasi keterbatasan pada jaringan saraf tiruan konvensional dalam mengingat informasi jangka panjang. Kemampuannya dalam mengenali pola yang tersembunyi dalam data historis penjualan menjadikannya sangat efektif untuk prediksi jangka pendek maupun jangka panjang, terutama pada data yang memiliki fluktuasi tinggi[25][26].

Studi ini membahas kesenjangan dalam evaluasi komparatif model prediktif untuk penjualan perhiasan emas menggunakan data transaksi dunia nyata. Pemilihan ketiga metode ini didasarkan pada perbedaan pendekatan algoritmik dan kemampuannya dalam memproses data penjualan emas yang bersifat dinamis dan fluktuatif. Masing-masing metode memiliki karakteristik, kelebihan, dan kelemahan yang berbeda, sehingga penting untuk dilakukan studi komparatif dalam hal akurasi prediksi. Komparasi ini bertujuan untuk mengidentifikasi metode yang paling optimal dan sesuai dengan karakteristik data penjualan emas, sehingga dapat memberikan hasil prediksi yang andal dan aplikatif.

Dengan demikian, penelitian ini diharapkan tidak hanya memberikan kontribusi dalam pengembangan model prediksi yang lebih akurat, tetapi juga menjadi acuan bagi pelaku industri dan peneliti dalam memilih pendekatan analisis yang tepat. Evaluasi performa dari ketiga metode tersebut akan memberikan gambaran yang lebih jelas mengenai efektivitas masing-masing algoritma dalam konteks prediksi penjualan emas sebagai data time series yang kompleks.

2. METODE

Penelitian ini bertujuan untuk membandingkan tingkat akurasi dari tiga algoritma machine learning, yaitu Naive Bayes, Random Forest, dan Long Short-Term Memory (LSTM) dalam memprediksi penjualan perhiasan emas. Ketiga algoritma ini dipilih karena mewakili pendekatan berbeda dalam pembelajaran mesin: Naive Bayes sebagai model berbasis probabilitas yang sederhana namun efektif, Random Forest sebagai metode ensemble learning yang menggabungkan banyak decision tree untuk meningkatkan stabilitas dan akurasi, serta LSTM sebagai bagian dari deep learning yang unggul dalam memproses data berurutan atau time series, seperti pola penjualan dari waktu ke waktu. Dalam konteks machine learning, proses pengembangan model tidak hanya berfokus pada pemilihan algoritma, tetapi juga mencakup serangkaian tahapan penting yang memengaruhi performa akhir model[27]. Oleh karena itu, metode penelitian ini dirancang secara sistematis mulai dari identifikasi masalah, pengumpulan data, preprocessing data, hingga tahap evaluasi model, diilustrasikan secara detail pada Gambar 1.



Gambar 1. Metode Penelitian

2.1 Pengumpulan Dataset

Dalam penelitian ini, dataset yang digunakan diperoleh dari data historis transaksi penjualan perhiasan emas di Toko Emas Parimas. Penggunaan data historis transaksi ini memungkinkan analisis

terhadap pola penjualan, tren musiman, serta variabel-variabel yang memengaruhi permintaan konsumen terhadap perhiasan emas[28]. Dataset ini mencakup periode waktu selama tiga tahun, yaitu dari Januari 2022 hingga Desember 2024, dan terdiri dari 4.595 entri transaksi. Sumber data ini adalah sistem pencatatan internal toko, yang secara otomatis merekam setiap transaksi yang terjadi, sehingga keakuratan dan kelengkapan data relatif dapat diandalkan.

Proses pengumpulan data dilakukan dengan cara mengakses dan mengekstrak informasi transaksi dari sistem manajemen toko yang telah terdigitalisasi. Hal ini memungkinkan peneliti untuk memperoleh data dalam bentuk yang terstruktur dan siap diolah untuk keperluan analisis lebih lanjut. Karena data diambil langsung dari catatan resmi toko, maka validitas data dipastikan tinggi dan relevan dengan konteks penelitian. Setiap entri dalam dataset terdiri atas tiga variabel utama, yaitu:

a. **Tanggal:**

Tanggal terjadinya transaksi penjualan. Variabel ini bersifat temporal dan digunakan sebagai acuan untuk mengurutkan data secara kronologis serta mengidentifikasi pola musiman atau tren penjualan.

b. **Jumlah item:**

Jumlah unit perhiasan emas yang terjual dalam satu transaksi. Variabel ini bersifat numerik dan mencerminkan volume transaksi harian.

c. **Berat (gram):**

Total berat perhiasan emas yang terjual dalam satu transaksi, diukur dalam satuan gram. Variabel ini penting untuk menghitung total volume penjualan dalam satuan fisik, yang lebih representatif dibanding hanya jumlah item karena setiap item dapat memiliki berat yang berbeda-beda.

Data yang terkumpul kemudian dibersihkan (*data cleaning*) untuk menghilangkan duplikasi, mengisi nilai yang hilang (jika ada), dan memastikan format konsistensi antar entri. Selanjutnya, data diklasifikasikan ke dalam kategori penjualan tinggi, rendah menggunakan metode klasifikasi untuk memungkinkan penerapan model pembelajaran mesin, yaitu Naive Bayes, Random Forest, dan Long Short-Term Memory (LSTM).

2.2 Pra-Pemrosesan Data

Tahap pra-pemrosesan data merupakan bagian krusial dalam proses analisis data karena bertujuan untuk memastikan bahwa data yang digunakan dalam penelitian berada dalam kondisi yang bersih, konsisten, dan dalam format yang sesuai untuk dianalisis lebih lanjut. Data mentah sering kali mengandung ketidaksesuaian, nilai hilang, atau informasi yang tidak relevan yang dapat memengaruhi kinerja dan akurasi model[29]. Adapun tahapan-tahapan pra-pemrosesan yang dilakukan dalam penelitian ini meliputi:

a. **Pembersihan Data (*Data Cleaning*)**

Langkah awal adalah menghapus duplikasi data, yakni entri yang identik dalam hal tanggal, jumlah item, dan berat. Selanjutnya, dilakukan penanganan nilai kosong (*missing values*) dengan metode imputasi. Dalam proses ini, fungsi `errors='coerce'` digunakan untuk memaksa nilai tidak valid (seperti entri string pada kolom numerik atau format tanggal yang salah) menjadi NaN. Setelah itu, semua baris yang mengandung NaN dihapus secara permanen untuk menjaga integritas dataset.

b. **Transformasi Waktu (*Time-Series Formatting*)**

Variabel tanggal diubah dari format string menjadi objek bertipe `datetime` agar dapat dikenali sebagai data deret waktu (*time series*). Transformasi ini memungkinkan model, khususnya LSTM, untuk menangkap pola musiman (*seasonal patterns*) dan tren berdasarkan periode waktu bulanan, mingguan, atau harian. Selain itu, dari variabel tanggal juga diturunkan fitur baru seperti bulan, hari dalam seminggu, dan hari libur (jika tersedia) sebagai variabel tambahan yang dapat memengaruhi pola penjualan.

c. **Transformasi Kategorik ke Numerik**

Kolom item yang semula dalam format string (misalnya: “Cincin”, “Gelang”, “Kalung”) dikonversi menjadi representasi numerik menggunakan teknik One-Hot Encoding. Teknik ini menghasilkan sejumlah kolom biner yang merepresentasikan setiap kategori item secara eksplisit, sehingga dapat diproses oleh algoritma pembelajaran mesin. Contohnya, jika terdapat tiga kategori item, maka akan dihasilkan tiga kolom biner: `item_cincin`, `item_gelang`, dan `item_kalung`, dengan nilai 1 jika kategori tersebut muncul dalam transaksi dan 0 jika tidak.

d. **Normalisasi Fitur**

Seluruh variabel numerik seperti jumlah item dan berat (gram) dinormalisasi menggunakan metode Min-Max Scaler. Teknik ini mengubah rentang nilai setiap fitur menjadi skala antara 0 hingga 1 tanpa mengubah distribusinya, sehingga model tidak bias terhadap fitur yang memiliki skala lebih besar. Normalisasi sangat penting, khususnya untuk algoritma seperti LSTM dan Naive Bayes yang sensitif terhadap skala input.

2.3 Pemisahan Data

Dalam pengembangan model pembelajaran mesin, pemisahan data berperan penting untuk memastikan bahwa model tidak hanya bekerja baik pada data pelatihan, tetapi juga mampu memberikan prediksi yang akurat terhadap data baru yang belum dikenalnya. Dengan memisahkan data menjadi beberapa subset, peneliti dapat melakukan evaluasi yang lebih obyektif dan mencegah overfitting. Dalam penelitian ini, digunakan dua skema pemisahan data untuk membandingkan stabilitas dan konsistensi performa model:

- Skema 1: 80% data digunakan sebagai data latih (training set), dan 20% sisanya sebagai data uji (testing set).
- Skema 2: 70% data digunakan sebagai data latih, dan 30% sebagai data uji.

Kedua skema dirancang untuk mengevaluasi sensitivitas model terhadap proporsi data latih dan data uji yang berbeda. Pemisahan dilakukan dengan menggunakan pendekatan time-based split, yaitu berdasarkan urutan kronologis data dari waktu transaksi. Artinya, data latih diambil dari bagian awal periode (misalnya Januari 2022 hingga sebagian akhir tahun 2023), sedangkan data uji berasal dari periode setelahnya (misalnya akhir 2023 hingga Desember 2024). Pendekatan ini digunakan untuk menghindari data leakage, yaitu kondisi di mana informasi dari masa depan (data uji) tanpa sengaja digunakan dalam proses pelatihan model, yang dapat menyebabkan hasil evaluasi menjadi tidak valid atau terlalu optimis. Dengan mempertahankan urutan waktu dalam proses pemisahan, model juga diuji secara lebih realistis dalam memprediksi kejadian di masa depan berdasarkan pola-pola historis.

2.4 Penerapan dan Pelatihan Model

Dalam pembelajaran mesin, proses pelatihan (*training*) dan penerapan (*deployment*) model merupakan dua tahapan penting dalam membangun sistem prediksi yang andal[30]. Tahap pelatihan bertujuan untuk mengenalkan pola pada data kepada model melalui algoritma tertentu, sehingga model mampu memetakan hubungan antara fitur masukan dan target keluaran. Sementara itu, tahap penerapan menguji kemampuan model dalam melakukan prediksi terhadap data baru yang tidak dikenali sebelumnya. Keberhasilan dari keseluruhan proses ini sangat tergantung pada karakteristik algoritma yang digunakan serta kesesuaian antara jenis data dan pendekatan model.

Dalam penelitian ini, diterapkan tiga jenis algoritma pembelajaran mesin yang memiliki pendekatan berbeda, yaitu *Naïve Bayes*, *Random Forest*, dan *Long Short-Term Memory* (LSTM). Ketiganya dilatih menggunakan data historis penjualan emas untuk kemudian diuji kemampuannya dalam melakukan klasifikasi.). Setiap model diterapkan dan dilatih menggunakan data latih yang telah diproses, lalu

dievaluasi pada data uji untuk mengukur akurasi prediksi terhadap kategori penjualan (rendah, sedang, tinggi).

- a. Naïve Bayes adalah algoritma klasifikasi berbasis probabilistik yang mengasumsikan independensi antar fitur. Meskipun pendekatannya sederhana, Naïve Bayes memiliki keunggulan dalam kecepatan pelatihan dan efisiensi ketika diterapkan pada data berskala besar. Model ini cocok untuk kasus klasifikasi awal yang tidak terlalu kompleks, dan umumnya memberikan hasil yang cukup baik ketika asumsi dasar independensi terpenuhi. Model Naive Bayes digunakan sebagai baseline karena kesederhanaannya dan efisiensinya dalam menangani data dengan dimensi terbatas. Algoritma ini bekerja dengan mengasumsikan independensi antar fitur dan menghitung probabilitas posterior untuk setiap kelas berdasarkan Teorema Bayes. Model ini dilatih menggunakan fitur numerik dan kategorik (setelah *one-hot encoding*), dan sangat cocok untuk kasus klasifikasi multikelas dengan jumlah data yang cukup besar[22].
- b. Random Forest merupakan algoritma *ensemble learning* yang terdiri dari banyak pohon keputusan (*decision tree*) yang dilatih secara paralel. Setiap pohon memberikan prediksi, dan hasil akhir ditentukan melalui agregasi, seperti voting mayoritas. Keunggulan Random Forest terletak pada kemampuannya menangani data berdimensi tinggi, mencegah overfitting melalui teknik *bagging*, serta memberikan hasil prediksi yang stabil dan akurat. Random Forest adalah algoritma berbasis ensemble learning yang menggabungkan banyak pohon keputusan (*decision trees*) untuk meningkatkan akurasi dan mengurangi overfitting. Model ini dilatih dengan menggunakan fitur numerik yang telah dinormalisasi dan hasil one-hot encoding untuk fitur kategorik[23]. Selama proses pelatihan, digunakan parameter default sebagai titik awal, kemudian dilakukan tuning parameter seperti jumlah pohon (*n_estimators*) dan kedalaman maksimum pohon (*max_depth*) menggunakan pencarian grid (*grid search*) atau validasi silang (*cross-validation*) untuk mendapatkan performa optimal.
- c. Long Short-Term Memory (LSTM) adalah varian dari *Recurrent Neural Network (RNN)* yang dirancang khusus untuk memproses data sekuensial dan bersifat temporal, seperti data deret waktu (*time-series*). LSTM memiliki arsitektur khusus yang memungkinkan model mengingat informasi penting dari waktu sebelumnya dan melupakan informasi yang tidak relevan. Hal ini menjadikannya sangat efektif dalam menangkap dependensi jangka panjang dalam data, meskipun memerlukan waktu pelatihan yang lebih lama dan sumber daya komputasi yang lebih tinggi dibandingkan model-model lainnya. LSTM digunakan untuk menangani karakteristik data deret waktu (*time series*) dengan ketergantungan temporal. Model ini dilatih menggunakan urutan data historis penjualan yang diubah ke dalam format *sliding window*, di mana setiap input mencerminkan urutan data selama N hari sebelumnya untuk memprediksi kategori penjualan pada hari ke-N+1. Model LSTM diimplementasikan menggunakan kerangka kerja deep learning seperti TensorFlow atau PyTorch, dengan arsitektur dasar yang terdiri dari satu atau dua lapisan LSTM, diikuti oleh lapisan Dense untuk klasifikasi. Fungsi aktivasi softmax digunakan pada lapisan output untuk menghasilkan probabilitas dari masing-masing kelas[24].

Ketiga model tersebut dilatih dengan cara yang sesuai dengan karakteristik masing-masing, menggunakan data pelatihan yang telah disiapkan. Selanjutnya, performa model diuji menggunakan data uji yang belum pernah dilibatkan dalam proses pelatihan, sehingga dapat dievaluasi kemampuan generalisasi dari masing-masing algoritma terhadap data baru.

2.5 Evaluasi Kinerja Model

Dalam *machine learning*, evaluasi kinerja model merupakan tahapan penting untuk menilai sejauh mana model mampu menghasilkan prediksi yang akurat dan dapat diandalkan. Evaluasi ini dilakukan dengan menggunakan data uji (*testing data*) yang tidak pernah terlibat dalam proses pelatihan (*training*),

sehingga hasil evaluasi dapat mencerminkan kemampuan model dalam melakukan generalisasi terhadap data baru yang belum pernah dilihat sebelumnya. Evaluasi yang baik tidak hanya mempertimbangkan tingkat akurasi semata, tetapi juga berbagai metrik lain yang memberikan gambaran lebih menyeluruh mengenai performa model, terutama dalam konteks klasifikasi[31].

Untuk mengukur kinerja model klasifikasi yang digunakan dalam penelitian ini, dilakukan evaluasi menggunakan beberapa metrik performa yang umum digunakan dalam pembelajaran mesin, khususnya untuk masalah klasifikasi multikelas. Evaluasi dilakukan pada data uji yang sebelumnya tidak pernah dilibatkan dalam proses pelatihan model.

a. **Accuracy**

Akurasi digunakan sebagai metrik dasar untuk menilai seberapa banyak prediksi model yang sesuai dengan kelas sebenarnya. Akurasi dihitung dengan membandingkan jumlah prediksi yang benar terhadap total jumlah prediksi. Akurasi penting sebagai metrik awal atau dasar untuk mengevaluasi model, terutama dalam dataset yang seimbang

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN}$$

Akurasi memberikan gambaran langsung tentang seberapa baik model klasifikasi memprediksi data secara keseluruhan dengan menghitung persentase prediksi yang benar dari total prediksi.

b. **Precision**

Precision mengukur sejauh mana prediksi positif yang dilakukan oleh model benar-benar relevan. Dalam konteks klasifikasi multikelas, precision dihitung untuk setiap kelas dan dirata-ratakan.

$$Precision = \frac{TP}{TP + FP}$$

Precision penting ketika biaya kesalahan positif tinggi, misalnya jika salah mengklasifikasikan penjualan "tinggi" padahal sebenarnya "rendah".

c. **Recall**

Recall (atau sensitivitas) mengukur seberapa banyak kasus aktual positif yang berhasil diprediksi dengan benar oleh model.

$$Recall = \frac{TP}{TP + FN}$$

Recall menjadi penting ketika tujuan utama adalah menangkap sebanyak mungkin kasus positif, misalnya mengidentifikasi semua potensi hari penjualan tinggi.

d. **F1-Score**

F1-score merupakan rata-rata harmonik dari precision dan recall. Metode ini memberikan keseimbangan antara keduanya, terutama saat terdapat ketidakseimbangan kelas.

$$F1 - score = 2 \times \frac{Precision \times Recall}{Precision + Recall}$$

Metrik ini sangat relevan untuk model seperti Naive Bayes dan Random Forest yang digunakan dalam klasifikasi multikelas dan ketika kelas target tidak seimbang.

2.6 Hyperparameter yang digunakan

Pada penelitian ini, digunakan tiga algoritma pembelajaran mesin dan pembelajaran dalam: Naïve Bayes, Random Forest, dan Long Short-Term Memory (LSTM). Masing-masing algoritma memiliki sejumlah hyperparameter yang disesuaikan untuk mendukung performa model secara optimal. Berikut penjelasannya:

a. **Naïve Bayes**

Algoritma Naïve Bayes merupakan metode klasifikasi berbasis probabilitas yang sederhana namun efektif. Hyperparameter yang digunakan dalam penelitian ini meliputi:

- `priors=None`: Nilai ini mengindikasikan bahwa distribusi probabilitas awal untuk masing-masing kelas tidak ditentukan secara manual, melainkan diasumsikan seimbang secara default.
- `var_smoothing=1e-9`: Parameter ini menambahkan nilai kecil ke varians untuk menghindari pembagian dengan nol atau nilai yang sangat kecil, yang dapat menyebabkan ketidakstabilan numerik.

b. Random Forest

Random Forest adalah algoritma ensemble berbasis pohon keputusan yang menggunakan banyak pohon (decision trees) untuk meningkatkan akurasi dan mengurangi overfitting. Hyperparameter yang digunakan:

- `n_estimators=100`: Menentukan jumlah pohon keputusan dalam ensemble. Dalam hal ini, digunakan 100 pohon untuk meningkatkan stabilitas prediksi.
- `random_state=42`: Menetapkan seed acak untuk memastikan hasil eksperimen dapat direproduksi dengan hasil yang konsisten.

c. Long Short-Term Memory (LSTM)

LSTM merupakan jenis Recurrent Neural Network (RNN) yang efektif dalam menangani data berurutan. Arsitektur LSTM dalam penelitian ini terdiri dari beberapa lapisan dengan hyperparameter sebagai berikut:

- LSTM Layer:
 - `units=64`: Jumlah unit memori (neuron) dalam lapisan LSTM yang menentukan kapasitas model dalam mengingat informasi jangka panjang.
 - `return_sequences=False`: Hanya mengembalikan output dari langkah waktu terakhir, karena hanya satu time step yang digunakan dalam prediksi.
- Dropout Layer:
 - `rate=0.3`: Menonaktifkan 30% neuron secara acak selama pelatihan untuk mengurangi risiko overfitting dan meningkatkan generalisasi model.
- Dense Layer:
 - Layer pertama: Terdiri dari 32 neuron dengan fungsi aktivasi ReLU, yang bertugas mengekstraksi fitur non-linier.
 - Layer output: Terdiri dari 1 neuron dengan fungsi aktivasi sigmoid, digunakan untuk menghasilkan probabilitas klasifikasi biner.
- Pengaturan kompilasi model (compile):
 - `loss='binary_crossentropy'`: Fungsi loss yang umum digunakan untuk tugas klasifikasi biner.
 - `optimizer='adam'`: Optimizer berbasis gradien yang adaptif dan efisien untuk pelatihan neural network.
 - `metrics=['accuracy']`: Digunakan untuk mengevaluasi akurasi model selama proses pelatihan dan pengujian.

3. HASIL

3.1 Hasil Uji Training dan Testing Naïve Bayes

Evaluasi performa model Naïve Bayes dilakukan dengan menggunakan dua skenario pembagian data, yaitu 80:20 dan 70:30 untuk training dan testing. Tujuannya adalah untuk melihat pengaruh proporsi data pelatihan terhadap akurasi dan kemampuan generalisasi model pada data uji.

Tabel 1. Training dan Testing

No	Training	Testing
1	80	20
2	70	30

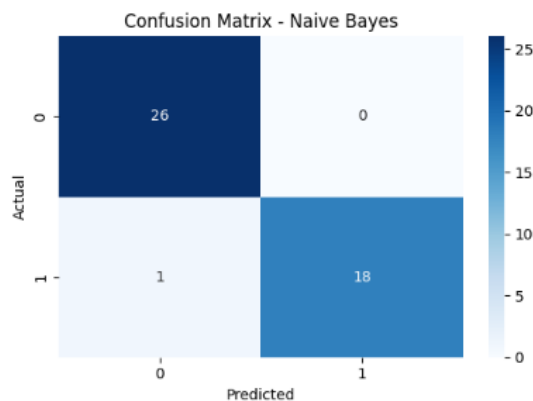
Untuk menguji konsistensi dan keandalan model prediksi, penelitian ini menggunakan dua skema pembagian data antara data latih (training) dan data uji (testing)[32]. Skema pertama menggunakan 80%

dari total data sebagai data pelatihan dan 20% sisanya sebagai data pengujian. Skema ini umum digunakan ketika dataset relatif besar, dengan tujuan memberikan model lebih banyak informasi untuk belajar. Sementara itu, skema kedua menggunakan 70% data pelatihan dan 30% data pengujian. Rasio ini bertujuan untuk mengevaluasi kemampuan generalisasi model dalam kondisi data pelatihan yang lebih terbatas. Dengan menerapkan dua rasio pembagian data tersebut, penelitian ini dapat membandingkan kinerja masing-masing model secara lebih menyeluruh dan memastikan bahwa performa model tidak hanya bergantung pada proporsi data yang digunakan. Evaluasi dilakukan pada kedua skema untuk melihat sejauh mana hasil prediksi model tetap stabil dan akurat terhadap variasi proporsi data latih dan uji.

```
Naive Bayes Metrics:  
Accuracy: 0.98  
Recall: 0.95  
Precision: 1.00  
F1 Score: 0.97
```

Gambar 2. Hasil Naïve Bayes Metrik

Model Naive Bayes merupakan salah satu algoritma klasifikasi yang berbasis pada Teorema Bayes, yang mengasumsikan independensi antar fitur. Meskipun asumsi ini sering kali tidak sepenuhnya terpenuhi dalam data dunia nyata, Naive Bayes dikenal mampu memberikan hasil klasifikasi yang kompetitif, terutama pada dataset berukuran besar dan bersifat spars. Model ini tetap efektif meskipun asumsi independensi tidak sepenuhnya akurat, karena pendekatannya yang probabilistik mampu menangani ketidakpastian dengan baik[33]. Gambar 2 Model Naive Bayes menunjukkan performa yang sangat baik dalam klasifikasi, dengan akurasi tinggi dan presisi sempurna. Hal ini menunjukkan bahwa model sangat andal dalam menghindari kesalahan klasifikasi positif (*false positives*), sekaligus cukup baik dalam mendeteksi semua kasus yang relevan (*true positives*). F1 Score yang mendekati sempurna (0.97) juga menunjukkan kestabilan kinerja model dalam situasi data yang mungkin tidak seimbang.



Gambar 3. Confusion Matrix Naïve Bayes

Gambar 3 memperlihatkan confusion matrix dari model Naive Bayes yang diterapkan dalam prediksi kelas penjualan emas. Model berhasil mengklasifikasikan 26 instance negatif dan 18 instance positif secara akurat. Hanya terdapat satu kesalahan klasifikasi, yaitu false negative, di mana model gagal mengenali satu instance kelas positif. Dengan akurasi sebesar 97,78%, precision mencapai 100%, recall 94,74%, dan F1-score sebesar 97,28%, model Naive Bayes menunjukkan performa yang sangat baik dan andal dalam konteks klasifikasi dataset ini. Ketidadaan false positive juga menegaskan bahwa model sangat presisi dan konservatif dalam mendeteksi kelas positif.

Pada skenario pertama, dengan 80% data untuk pelatihan dan 20% untuk pengujian, model menunjukkan performa yang sangat baik. Akurasi yang dicapai sebesar 0,98, menandakan bahwa sebagian

besar prediksi model sesuai dengan label aktual. Selain itu, nilai recall sebesar 0,95 menunjukkan kemampuan model dalam menangkap hampir seluruh instance positif. Nilai precision yang sempurna (1,00) menunjukkan bahwa seluruh prediksi positif yang dihasilkan oleh model adalah benar, tanpa kesalahan klasifikasi positif palsu (false positive). Hal ini menghasilkan nilai F1 Score sebesar 0,97, yang mencerminkan keseimbangan optimal antara precision dan recall.

Pada skenario kedua, dengan 70% data untuk pelatihan dan 30% untuk pengujian, performa model sedikit menurun namun tetap berada dalam kategori tinggi. Akurasi menurun menjadi 0,94, dengan nilai recall sebesar 0,93 dan precision sebesar 0,98. Nilai F1 Score sebesar 0,96 masih menunjukkan bahwa model tetap konsisten dalam melakukan klasifikasi dengan tingkat kesalahan yang rendah.

Secara keseluruhan, hasil ini menunjukkan bahwa model Naïve Bayes memiliki kinerja yang stabil dan andal, bahkan ketika jumlah data pelatihan dikurangi. Meskipun terdapat sedikit penurunan performa pada rasio 70:30, model tetap mampu mempertahankan presisi tinggi dan generalisasi yang baik. Temuan ini menegaskan bahwa Naïve Bayes merupakan algoritma klasifikasi yang efektif dan efisien untuk kasus prediksi seperti yang digunakan dalam penelitian ini, terutama ketika ketersediaan data pelatihan terbatas.

3.2 Hasil Uji Training dan Testing Random Forest

Dalam konteks pembelajaran mesin, proses pelatihan dan pengujian model menjadi kunci utama dalam mengevaluasi kemampuan model untuk melakukan generalisasi terhadap data yang belum pernah dilihat sebelumnya. Generalisasi yang baik tercermin dari konsistensi performa model, baik pada data pelatihan maupun data pengujian. Salah satu pendekatan yang lazim digunakan adalah membagi dataset menjadi dua bagian, yakni data pelatihan (training data) dan data pengujian (testing data). Semakin representatif data pelatihan, semakin besar kemungkinan model dapat menangkap pola yang relevan tanpa mengandalkan ingatan terhadap data semata.

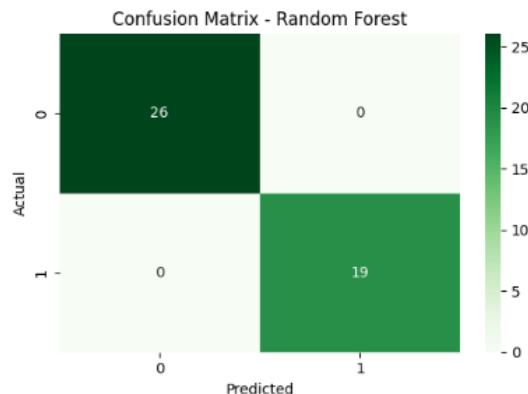
```
Random Forest Metrics:  
Accuracy: 1.00  
Recall: 1.00  
Precision: 1.00  
F1 Score: 1.00
```

Gambar 4. Random Forest Metrik

Pada studi ini, model Random Forest dilatih menggunakan 80% dari total dataset, sementara 20% sisanya dialokasikan untuk pengujian. Berdasarkan metrik evaluasi yang ditampilkan, model menunjukkan performa klasifikasi yang sangat tinggi, bahkan sempurna, pada skenario pembagian data ini. Hasil ini tercermin dalam nilai akurasi, precision, recall, dan F1-score yang mencapai angka maksimal, sebagaimana terlihat pada Gambar 4. Menariknya, ketika proporsi data pelatihan dikurangi menjadi 70% dan data pengujian ditingkatkan menjadi 30%, model Random Forest masih menunjukkan performa yang sangat tinggi. Jika hasil ini konsisten, maka hal ini menjadi indikator bahwa model memiliki kapasitas generalisasi yang sangat baik, yakni mampu mengenali pola secara akurat meskipun dengan data pelatihan yang lebih sedikit. Namun, apabila terjadi penurunan skor pada metrik evaluasi, hal tersebut dapat dimaklumi mengingat keterbatasan informasi yang tersedia selama proses pelatihan.

Model Random Forest merupakan salah satu algoritma ensemble learning yang dibangun dari sejumlah pohon keputusan (decision trees) dan menggabungkan hasilnya untuk meningkatkan akurasi prediksi serta mengurangi risiko overfitting. Pendekatan ini bekerja dengan membangun banyak pohon pada subset acak dari data pelatihan, kemudian menggabungkan hasilnya melalui voting (klasifikasi) atau rata-rata (regresi). Keunggulan utama Random Forest terletak pada kemampuannya untuk menangani variabel dalam jumlah besar, mendeteksi interaksi non-linear, dan secara umum memberikan performa

tinggi tanpa banyak penyetelan parameter[22]. Pada gambar 4 diketahui bahwa Model Random Forest dalam konteks ini menunjukkan performa klasifikasi yang sempurna, dengan nilai maksimal pada seluruh metrik evaluasi. Hasil ini bisa mengindikasikan bahwa model sangat cocok dengan data (fit dengan sangat baik) dan tidak terjadi kesalahan klasifikasi sama sekali. Namun, perlu diwaspadai kemungkinan overfitting, terutama jika performa ini tidak konsisten pada data uji lainnya atau data dunia nyata yang belum pernah dilihat sebelumnya. Jika digunakan dalam studi komparatif dengan model lain seperti Naive Bayes, maka Random Forest secara objektif menunjukkan performa yang lebih unggul dalam kasus ini. Namun, validitas hasil ini tetap perlu diuji pada data yang lebih besar dan beragam.



Gambar 5. *Confusion Matrix Random Forest*

Untuk memperkuat analisis ini, Gambar 5 menampilkan confusion matrix dari model Random Forest. Terlihat bahwa sebanyak 26 instance dari kelas negatif dan 19 instance dari kelas positif berhasil diklasifikasikan dengan benar, tanpa adanya kesalahan prediksi. Dengan demikian, seluruh metrik evaluasi mencapai nilai 1.00, mencerminkan akurasi dan konsistensi model dalam menangani klasifikasi tingkat penjualan emas. Hasil ini semakin memperkuat posisi Random Forest sebagai model yang sangat andal dalam konteks dataset ini, terutama dalam mengelola data yang kompleks dan multivariat.

Bila dibandingkan secara objektif dengan algoritma lain, seperti Naive Bayes, maka Random Forest menunjukkan keunggulan yang signifikan dalam studi ini. Meski demikian, validitas performa ini tetap perlu diuji lebih lanjut menggunakan dataset yang lebih luas dan heterogen, guna memastikan bahwa keunggulan tersebut bukan semata-mata hasil dari keterbatasan atau kekhasan data yang digunakan dalam eksperimen ini.

3.3 Hasil Uji Training dan Testing Long Short Term Memory

Model Long Short-Term Memory (LSTM) diuji dalam dua skenario pembagian data, yakni 80% data dialokasikan untuk proses pelatihan dan 20% sisanya untuk pengujian (80:20), serta skenario kedua dengan proporsi 70% untuk pelatihan dan 30% untuk pengujian (70:30). Pengujian ini bertujuan untuk mengevaluasi kemampuan model dalam mengenali pola data serta mengukur sejauh mana model mampu melakukan generalisasi terhadap data yang belum pernah dilihat sebelumnya.

Pada skenario 80:20, hasil evaluasi performa model LSTM menunjukkan capaian yang cukup baik. Berdasarkan metrik evaluasi yang digunakan, diperoleh sebagaimana Gambar 6.

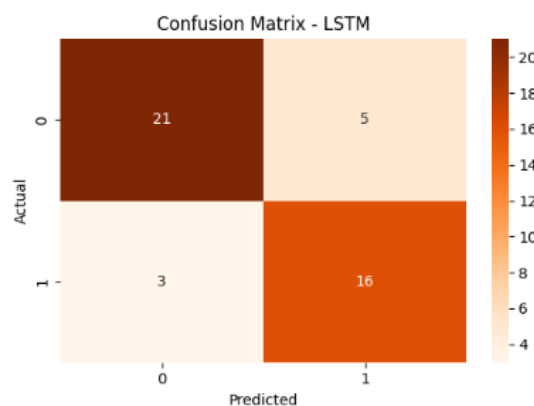
- *Accuracy* sebesar 0,82, yang menunjukkan bahwa 82% dari seluruh prediksi model sesuai dengan label aktual;
- *Recall* sebesar 0,84, mengindikasikan bahwa model cukup efektif dalam mengidentifikasi instance dari kelas positif secara benar;

- *Precision* sebesar 0,76, yang menggambarkan bahwa 76% dari seluruh prediksi positif yang dihasilkan oleh model merupakan prediksi yang tepat sasaran;
- *F1-Score* sebesar 0,80, yang mencerminkan keseimbangan antara *precision* dan *recall*, serta memberikan gambaran umum mengenai stabilitas performa model dalam konteks klasifikasi.

Nilai-nilai tersebut mengindikasikan bahwa model LSTM mampu mendeteksi dan mengklasifikasikan kedua kelas dengan tingkat ketepatan yang cukup seimbang, tanpa menunjukkan kecenderungan yang berat sebelah terhadap salah satu kelas. Hal ini menjadi indikator bahwa model memiliki potensi yang baik untuk diterapkan dalam kasus klasifikasi berbasis data sekuensial, seperti pada prediksi tingkat penjualan emas, meskipun akurasinya masih dapat ditingkatkan melalui optimasi lebih lanjut terhadap parameter pelatihan atau arsitektur model.

```
LSTM Metrics:
Accuracy: 0.82
Recall: 0.84
Precision: 0.76
F1 Score: 0.80
```

Gambar 6. Hasil LSTM Metrics



Gambar 7. Confusion Matrix Long Short-Term Memory

Long Short-Term Memory (LSTM) adalah varian dari Recurrent Neural Network (RNN) yang dirancang untuk mengatasi masalah vanishing gradient dan memungkinkan pembelajaran dependensi jangka panjang dalam data sekuensial[18]. Model LSTM menunjukkan kinerja yang cukup baik, namun masih lebih rendah dibandingkan dengan model Random Forest maupun Naive Bayes dalam konteks evaluasi ini. Nilai *Precision* yang relatif rendah mengindikasikan bahwa model cenderung memberikan lebih banyak prediksi positif palsu (false positive), meskipun performa *Recall*-nya cukup kuat. Dalam studi komparatif, LSTM mungkin memiliki keunggulan dalam memproses data sekuensial atau temporal (misalnya prediksi berdasarkan urutan waktu), tetapi untuk dataset ini, kinerjanya belum optimal.

Sedangkan gambar 7 menyajikan confusion matrix dari model Long Short-Term Memory (LSTM). Model ini berhasil mengklasifikasikan 21 instance negatif dan 16 instance positif secara benar, namun menghasilkan 5 kesalahan false positive dan 3 false negative. Hal ini menghasilkan nilai akurasi sebesar 82,22%, *precision* 76,19%, *recall* 84,21%, dan *F1-score* 79,96%. Dibandingkan dengan model Naive Bayes dan Random Forest, performa LSTM relatif lebih rendah. Hal ini kemungkinan disebabkan oleh karakteristik LSTM yang lebih optimal untuk prediksi berkelanjutan atau regresi time-series daripada tugas klasifikasi biner. Meski demikian, LSTM tetap menunjukkan performa yang kompeten dalam mengenali pola temporal dari data penjualan emas.

Tabel 2. Epoch 80% Pelatihan 20% Pengujian

Epoch	Accuracy	Loss
1	0.644	0.686
2	0.760	0.668
3	0.805	0.635
4	0.762	0.609
5	0.883	0.551
6	0.808	0.513
7	0.860	0.401
8	0.894	0.333
9	0.877	0.311
10	0.878	0.306

Tabel 3. Epoch 70% Pelatihan 30% Pengujian

Epoch	Accuracy	Loss
1	0.453	0.694
2	0.805	0.680
3	0.857	0.660
4	0.865	0.631
5	0.852	0.598
6	0.850	0.547
7	0.805	0.505
8	0.838	0.445
9	0.809	0.424
10	0.813	0.393

Penelitian ini membandingkan performa model dengan dua pembagian data yang berbeda, yaitu 80% data untuk pelatihan dan 20% untuk pengujian (Tabel 2), serta 70% untuk pelatihan dan 30% untuk pengujian (Tabel 3). Hasil pada Tabel 2 menunjukkan bahwa model mengalami peningkatan akurasi dari 64.4% menjadi 87.8% selama 10 epoch. Selain itu, nilai loss menurun dari 0.686 ke 0.306, menandakan bahwa model semakin baik dalam mengenali pola data. Peningkatan paling signifikan terjadi pada epoch ke-5, di mana akurasi melonjak ke 88.3%. Setelah itu, akurasi tetap cukup stabil dan tinggi, menunjukkan proses pelatihan yang baik.

Sementara itu, pada Tabel 3, model memulai pelatihan dengan akurasi yang lebih rendah, yaitu 45.3%. Namun, akurasi meningkat secara bertahap hingga mencapai 81.3% pada epoch ke-10. Loss juga menurun dari 0.694 menjadi 0.393. Meski nilainya tidak sebaik Tabel 2, model tetap menunjukkan peningkatan yang konsisten. Ada sedikit penurunan akurasi di beberapa epoch, seperti dari epoch ke-6 ke ke-7, yang mungkin mengindikasikan gejala overfitting ringan.

Secara keseluruhan, model dengan 80% data pelatihan menunjukkan performa yang lebih baik dalam hal akurasi dan loss. Namun, model dengan 70% data pelatihan memberikan gambaran yang lebih adil terhadap kemampuan generalisasi model karena diuji dengan lebih banyak data. Untuk aplikasi di dunia nyata, skenario ini mungkin lebih menggambarkan performa model saat diterapkan pada data baru.

Tabel 4. Komparasi Tingkat Akurasi Naïve Bayes, Random Forest dan LSTM

Naïve Bayes		Random Forest		LSTM	
80%	70%	80%	70%	80%	70%
0.98	0.97	1.00	1.00	0.82	0.88

Tabel 4 menunjukkan perbandingan akurasi antara tiga algoritma yang berbeda: Naïve Bayes, Random Forest, dan LSTM. Hasilnya menunjukkan bahwa Naïve Bayes dan Random Forest unggul secara signifikan, dengan akurasi masing-masing mencapai 0.98–0.97 dan 1.00 pada kedua skenario data (80% dan 70%). Sementara itu, akurasi LSTM lebih rendah, yaitu 0.82 dengan 80% data pelatihan dan justru meningkat menjadi 0.88 saat data pelatihan diubah menjadi 70%.

3.4 Metrik Evaluasi

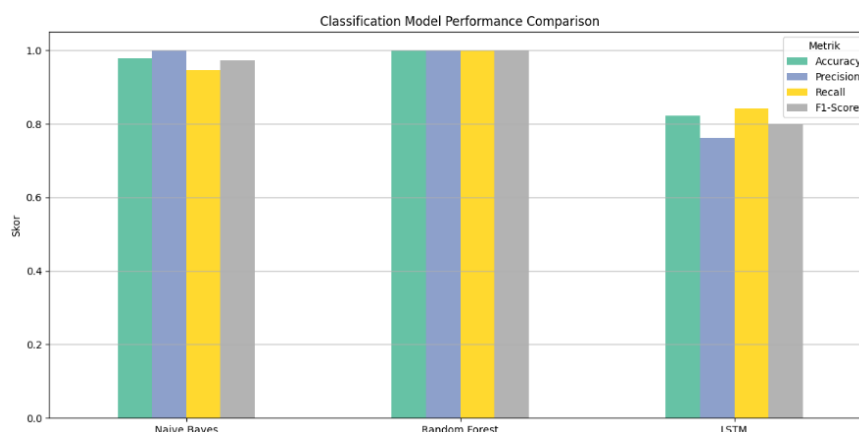
Penelitian ini melanjutkan dengan membandingkan kinerja tiga algoritma klasifikasi utama, yakni Naive Bayes, Random Forest, dan Long Short-Term Memory (LSTM). Perbandingan ini dilakukan menggunakan empat metrik evaluasi penting, yaitu akurasi (accuracy), presisi (precision), recall, dan F1-score, yang masing-masing merepresentasikan aspek berbeda dari kemampuan klasifikasi model (Figure 3.).

Akurasi mengukur proporsi keseluruhan prediksi yang benar terhadap total data, presisi menilai kemampuan model untuk mengidentifikasi kelas positif secara tepat tanpa banyak kesalahan, recall mencerminkan sensitivitas model dalam menangkap seluruh kejadian sebenarnya, sedangkan F1-score merupakan rata-rata harmonik dari presisi dan recall, memberikan gambaran keseimbangan performa.

Hasil evaluasi menunjukkan bahwa algoritma Random Forest memberikan performa terbaik secara konsisten di semua metrik, dengan skor sempurna sebesar 1.00 untuk akurasi, presisi, recall, dan F1-score. Ini menandakan bahwa Random Forest mampu melakukan klasifikasi dengan sangat tepat dan lengkap tanpa kehilangan informasi penting. Di posisi kedua, Naive Bayes juga menunjukkan hasil yang sangat baik dengan akurasi dan metrik lainnya berada di kisaran 0.97 hingga 0.99, mencerminkan kestabilan dan keandalan dalam prediksi.

Sebaliknya, LSTM meskipun merupakan model yang kuat untuk analisis data deret waktu, menunjukkan performa yang lebih rendah dalam konteks tugas klasifikasi ini. Skor akurasi LSTM berkisar sekitar 0.83, dengan presisi dan F1-score yang lebih rendah dibandingkan dua model sebelumnya. Hal ini mengindikasikan bahwa meskipun LSTM efektif dalam memproses urutan data, pada masalah klasifikasi yang dihadapi dalam penelitian ini, model ini kurang optimal dibanding Random Forest dan Naive Bayes.

Secara keseluruhan, perbandingan ini menegaskan bahwa Random Forest adalah model klasifikasi paling unggul dan konsisten pada dataset yang digunakan, sementara Naive Bayes juga menjadi alternatif yang efektif. Sementara itu, LSTM, dengan performa yang relatif lebih rendah, lebih cocok digunakan pada aplikasi yang memerlukan pemodelan deret waktu daripada klasifikasi langsung pada data yang bersifat non-sekuensial.



Gambar 8. Perbandingan Kinerja Model Klasifikasi 80% Pelatihan 20% Pengujian

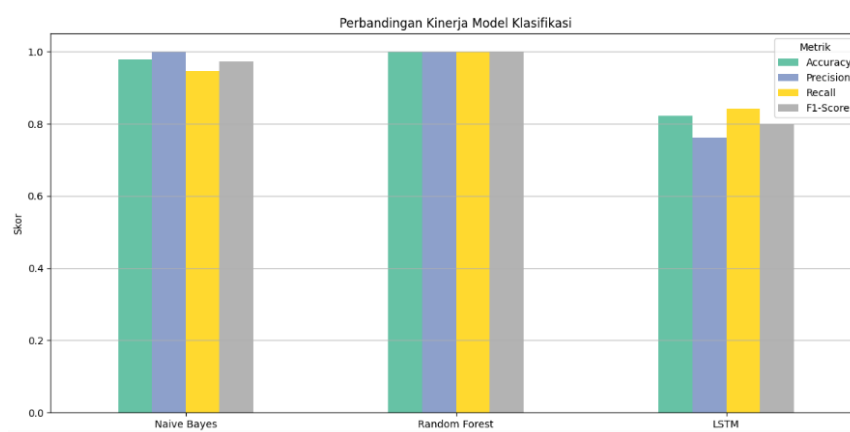
Penelitian ini juga melakukan perbandingan kinerja tiga algoritma klasifikasi, yaitu Naive Bayes, Random Forest, dan Long Short-Term Memory (LSTM). Evaluasi dilakukan menggunakan empat metrik utama: akurasi (accuracy), presisi (precision), recall, dan F1-score sebagaimana gambar 8. Setiap metrik memberikan gambaran spesifik tentang aspek performa model, di mana akurasi mengukur proporsi prediksi yang benar secara keseluruhan, presisi menilai ketepatan prediksi positif, recall menunjukkan kemampuan model mendeteksi seluruh kasus positif yang sebenarnya, dan F1-score merepresentasikan keseimbangan antara presisi dan recall.

Analisis hasil menunjukkan bahwa Naive Bayes memperoleh skor akurasi, presisi, dan F1-score yang sangat tinggi, berkisar antara 0.98 hingga 0.99, dengan recall sedikit lebih rendah sekitar 0.96. Hal ini mengindikasikan bahwa meskipun Naive Bayes sangat akurat, model ini memiliki sedikit keterbatasan dalam menjangkau seluruh kasus positif secara menyeluruh. Sebaliknya, Random Forest tampil sebagai model yang paling unggul dan stabil, dengan skor semua metrik mendekati atau mencapai nilai sempurna (1.00). Performanya yang seragam dan superior menjadikan Random Forest sebagai pilihan terbaik dalam klasifikasi pada dataset ini.

Sementara itu, LSTM menunjukkan performa yang lebih rendah dengan akurasi dan F1-score sekitar 0.82 hingga 0.85, presisi sekitar 0.78, dan recall sekitar 0.84. Hasil ini menunjukkan bahwa LSTM kurang seakurat dan seimbang dibandingkan dua model lainnya dalam tugas klasifikasi. Kondisi ini kemungkinan disebabkan oleh sifat LSTM yang lebih cocok untuk aplikasi regresi pada data deret waktu daripada klasifikasi data kategori.

Sebagai kelanjutan dari evaluasi performa model klasifikasi, Matriks gambar 9 memberikan gambaran rinci mengenai jumlah prediksi benar dan salah yang dilakukan model pada dua kelas utama, yaitu kelas 0 yang merepresentasikan kondisi negatif (misalnya, tidak terjadi penjualan tinggi) dan kelas 1 yang merepresentasikan kondisi positif (terjadi penjualan tinggi).

Confusion matrix menjadi alat penting untuk memahami sejauh mana model dapat membedakan antara dua kelas tersebut secara tepat. Nilai pada matriks menggambarkan jumlah instance yang diprediksi sebagai kelas 0 maupun kelas 1, dibandingkan dengan kelas sebenarnya. Dengan demikian, matriks ini membantu mengidentifikasi jenis kesalahan klasifikasi yang terjadi, seperti false positive (kelas negatif diprediksi positif) dan false negative (kelas positif diprediksi negatif).

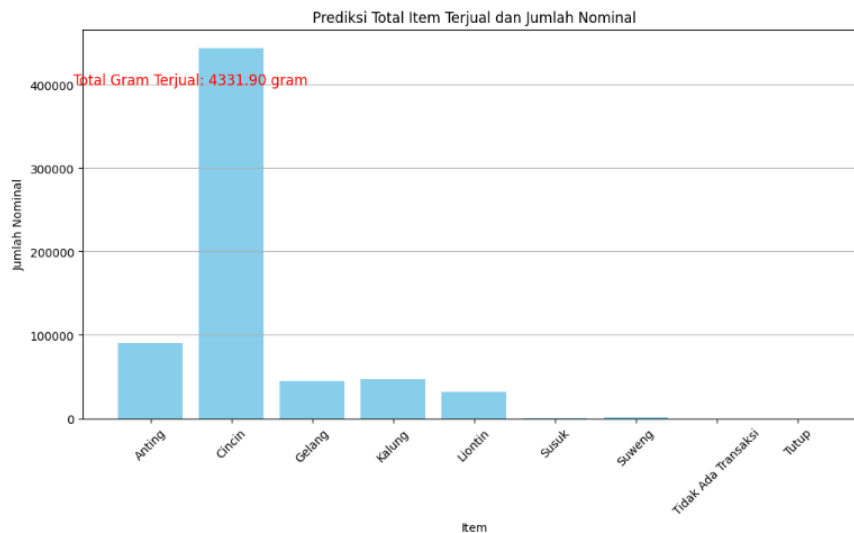


Gambar 9. Perbandingan Kinerja Model Klasifikasi 70% Pelatihan 30% Pengujian

Analisis matriks kebingungan ini melengkapi metrik evaluasi sebelumnya, memberikan gambaran yang lebih komprehensif terhadap keandalan model Naive Bayes dalam mengklasifikasikan data penjualan berdasarkan tingkat penjualan tinggi atau tidak. Dengan memahami pola kesalahan yang terjadi, dapat

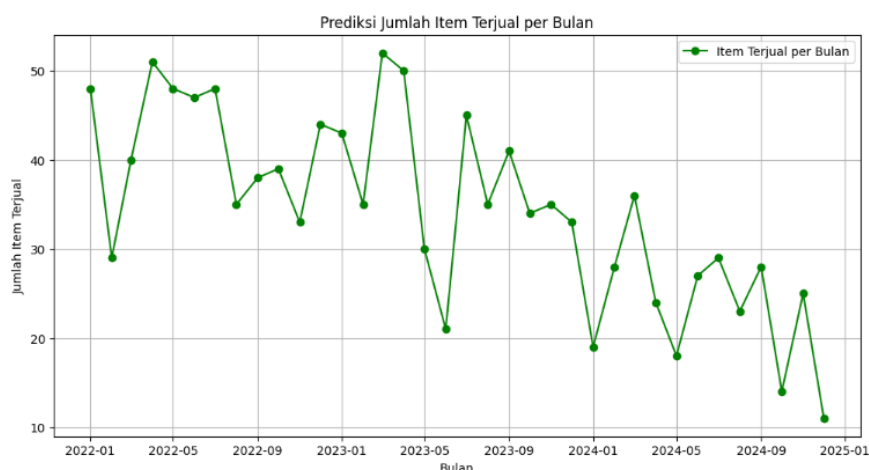
dilakukan perbaikan model atau penyesuaian strategi klasifikasi untuk meningkatkan performa di masa mendatang.

3.5 Prediksi Pola Penjualan Produk



Gambar 10. Prediksi Total Item Terjual dan Jumlah Nominal

Setelah mengevaluasi dan membandingkan performa berbagai model prediksi menggunakan dua skema pembagian data yang berbeda, penelitian ini melanjutkan analisis terhadap pola penjualan produk perhiasan berdasarkan data riil. Hasil data penjualan berdasarkan gambar 10 mengungkapkan bahwa cincin merupakan jenis perhiasan dengan volume penjualan tertinggi, dengan nilai penjualan yang secara signifikan melebihi jenis perhiasan lainnya. Urutan berikutnya ditempati oleh anting, gelang, kalung, dan liontin. Sebaliknya, jenis perhiasan seperti susuk, subeng, serta kategori “Tidak Ada Transaksi” dan “Tutup” hampir tidak menunjukkan aktivitas penjualan yang berarti. Untuk prediksi total jumlah gram terjual adalah 4331,90 gram.



Gambar 11. Prediksi Jumlah Item Terjual Perbulan

Setelah menguraikan pola penjualan berdasarkan jenis perhiasan, analisis selanjutnya difokuskan pada tren temporal penjualan item emas secara keseluruhan dari Januari 2022 hingga Januari 2025. Grafik tren penjualan ini menggunakan sumbu horizontal yang merepresentasikan waktu dalam format bulanan,

sementara sumbu vertikal menunjukkan jumlah item emas yang terjual setiap bulan. Garis berwarna hijau dengan titik-titik menandai prediksi perkembangan jumlah penjualan dari bulan ke bulan, yang pada legenda di kanan atas diberi label “Item Terjual per Bulan”.

Dari grafik gambar 11 bahwa jumlah item yang terjual mengalami fluktuasi yang cukup signifikan setiap bulannya, mencerminkan adanya variasi permintaan sepanjang waktu. Beberapa titik puncak penjualan tercatat pada pertengahan tahun 2022 dan awal tahun 2023, di mana jumlah item terjual mencapai lebih dari 50 unit per bulan. Namun, setelah periode tersebut, terdapat tren penurunan secara bertahap dalam volume penjualan yang berlangsung hingga akhir tahun 2024. Pada periode ini, jumlah item yang terjual per bulan cenderung menurun hingga di bawah 30 unit, bahkan mendekati angka 10 unit pada akhir grafik.

4. DISKUSI

Hasil penelitian ini menegaskan bahwa pemilihan algoritma memiliki dampak signifikan terhadap performa model prediksi dalam konteks klasifikasi tingkat penjualan emas. Dari tiga algoritma yang diuji—Naive Bayes, Random Forest, dan Long Short-Term Memory (LSTM)—algoritma Random Forest menunjukkan performa paling unggul dan konsisten. Model ini mencatat akurasi sempurna (100%) pada kedua skema pembagian data (80:20 dan 70:30), serta memperoleh nilai maksimal pada metrik evaluasi lainnya seperti presisi, recall, dan F1-score.

Temuan ini memperkuat literatur yang telah ada, seperti studi oleh Landge et al [34]. yang menerapkan Random Forest untuk memprediksi harga emas (GLD) dengan mempertimbangkan variabel ekonomi seperti indeks S&P 500 (SPX), harga minyak (USO), harga perak (SLV), dan nilai tukar EUR/USD. Dalam studi tersebut, Random Forest juga terbukti menjadi model terbaik dibandingkan pendekatan lain, menunjukkan kapasitas algoritma ini dalam memahami kompleksitas hubungan multivariat dalam data keuangan. Dengan demikian, efektivitas Random Forest dalam penelitian ini tidak hanya terbatas pada klasifikasi penjualan, tetapi juga relevan untuk berbagai aplikasi di ranah ekonomi dan pasar komoditas.

Selain itu, studi Langsetmo et al. [35], yang menunjukkan keunggulan Random Forest dalam menangani prediksi risiko patah tulang pinggul dan mortalitas pada lansia, meskipun pada konteks berbeda. Random Forest unggul karena sifat non-parametriknya yang memungkinkan pemodelan relasi kompleks antar variabel tanpa asumsi linearitas. Fleksibilitas ini menjadikannya sangat adaptif dalam klasifikasi multikelas dengan dimensi data tinggi dan interaksi variabel yang kompleks, seperti halnya dalam klasifikasi penjualan emas berdasarkan indikator harga dan tren historis.

Keunggulan Random Forest dalam penelitian ini juga memperkuat temuan dari studi oleh Syahputra et al. [36], yang juga menemukan bahwa algoritma ensemble seperti Random Forest memiliki performa prediksi yang kuat dalam konteks regresi penjualan (pada kasus kelapa sawit). Namun, menariknya, studi tersebut menemukan bahwa algoritma XGBoost lebih unggul dari Random Forest dalam hal efisiensi waktu dan akurasi (RMSE dan R^2) dalam prediksi penjualan sawit. Perbedaan hasil ini mengindikasikan bahwa keunggulan suatu algoritma sangat kontekstual terhadap jenis data, tujuan model (klasifikasi vs regresi), dan pola temporal yang terkandung dalam dataset.

Sementara itu, algoritma Naive Bayes meskipun tidak seunggul Random Forest, tetap menunjukkan performa yang baik dengan akurasi mencapai 0.97–0.98. Keunggulan model ini terletak pada kesederhanaannya, kecepatan pelatihan, dan presisi yang tinggi dalam mengklasifikasikan kelas positif. Meski recall sedikit lebih rendah, hal ini masih dalam batas wajar dan menjadikan Naive Bayes sebagai alternatif yang efisien, terutama dalam kondisi sumber daya terbatas.

Penggunaan LSTM (*Long Short-Term Memory*) dalam klasifikasi biner—seperti yang dijelaskan dalam artikel “Algoritma LSTM-CNN untuk Sentimen Klasifikasi dengan Word2vec pada Media Online”

merupakan pendekatan *deep learning* yang efektif untuk menganalisis dan mengklasifikasikan sentimen teks. Namun, efektivitas tersebut sangat dipengaruhi oleh karakteristik data yang digunakan[37].

Berbeda halnya dengan Long Short-Term Memory (LSTM), yang secara umum dikenal efektif dalam pemrosesan data deret waktu atau data yang memiliki ketergantungan antar waktu (*temporal dependencies*). Dalam penelitian ini, penerapan LSTM pada tugas klasifikasi statis, yaitu klasifikasi berdasarkan data yang tidak memiliki urutan waktu eksplisit, menunjukkan keterbatasan. Hal ini tercermin dari akurasi yang relatif lebih rendah (0.82–0.88), serta meningkatnya nilai *false positive* dan *false negative*.

Ketidakefektifan LSTM dalam konteks ini disebabkan oleh sifat dasar arsitekturnya yang dirancang untuk menangkap hubungan berurutan (*sequential patterns*) dalam data. Pada klasifikasi statis, seluruh fitur diasumsikan independen terhadap urutan atau waktu, sehingga keunggulan utama LSTM—yaitu memori jangka pendek dan panjang—tidak termanfaatkan secara optimal. Alih-alih memperkuat kinerja, kompleksitas arsitektur LSTM justru dapat menyebabkan *overfitting* atau kesulitan dalam konvergensi saat tidak ada informasi temporal yang signifikan.

Meskipun performa LSTM sedikit membaik pada rasio data latih yang lebih kecil (70%), hasilnya masih belum melampaui model konvensional seperti Random Forest atau Naive Bayes, yang lebih sederhana dan memang dirancang untuk klasifikasi statis. Dengan demikian, penggunaan LSTM dalam klasifikasi statis memerlukan modifikasi arsitektur, pengayaan fitur, atau bahkan transformasi data menjadi bentuk sekuensial agar model dapat mengidentifikasi pola dengan lebih efektif. Sebaliknya, potensi LSTM lebih sesuai diterapkan pada tugas regresi atau prediksi jangka panjang berbasis waktu, seperti forecasting tren penjualan atau harga, di mana pola temporal dan urutan historis memegang peran penting dalam pembentukan output prediksi.

Secara keseluruhan, hasil penelitian ini menggarisbawahi pentingnya pemilihan algoritma yang sesuai dengan karakteristik data dan tujuan analisis. Random Forest terbukti sebagai model yang paling andal untuk klasifikasi data penjualan dalam konteks ini, dengan kinerja yang tidak hanya konsisten terhadap variasi rasio data latih dan uji, tetapi juga superior dalam semua metrik evaluasi. Penelitian ini sekaligus memperkuat bukti empiris dari studi-studi sebelumnya tentang keunggulan Random Forest dalam aplikasi prediktif berbasis data keuangan dan ekonomi.

5. KESIMPULAN

Penelitian ini menyimpulkan bahwa algoritma Random Forest merupakan model paling unggul dan andal dalam memprediksi penjualan emas berdasarkan data historis. Dengan nilai akurasi, presisi, recall, dan F1-score yang mencapai angka sempurna (1.00), Random Forest menunjukkan performa klasifikasi yang konsisten, baik pada skema pembagian data 80:20 maupun 70:30. Keunggulan ini didukung oleh studi sebelumnya yang juga menempatkan Random Forest sebagai algoritma terbaik dalam konteks klasifikasi dan prediksi harga emas.

Model Naive Bayes turut menunjukkan performa yang tinggi dengan tingkat kesalahan yang sangat rendah, menjadikannya sebagai alternatif efisien untuk prediksi cepat dengan sumber daya komputasi terbatas. Sebaliknya, LSTM menunjukkan performa yang relatif lebih rendah dalam tugas klasifikasi statis, meskipun tetap relevan untuk aplikasi yang melibatkan data sekuensial atau prediksi deret waktu. Dengan demikian, Random Forest direkomendasikan sebagai model utama untuk sistem pendukung keputusan dalam konteks prediksi penjualan emas, khususnya pada data dengan karakteristik non-sekuensial dan klasifikasi biner. Studi ini berkontribusi pada bidang informatika terapan dengan menunjukkan keunggulan Random Forest dalam tugas klasifikasi pada data penjualan transaksional, serta memberikan arah yang jelas dalam pemilihan model prediktif berdasarkan konteks data dan kebutuhan sistem.

Berdasarkan hasil penelitian, disarankan agar algoritma Random Forest diimplementasikan dalam sistem pendukung keputusan untuk prediksi penjualan emas, mengingat performanya yang sangat tinggi

dan konsisten. Untuk meningkatkan generalisasi model, perlu dilakukan uji coba pada dataset yang lebih besar dan bervariasi, termasuk data dari periode dan wilayah berbeda.

Algoritma Naive Bayes, dengan efisiensi komputasinya, direkomendasikan sebagai alternatif pada sistem dengan keterbatasan sumber daya. Sementara itu, penggunaan LSTM lebih tepat diarahkan pada analisis deret waktu atau tren jangka panjang, bukan klasifikasi statis.

Penelitian di masa depan bisa lebih fokus pada *hybrid model* atau kumpulan data yang lebih besar yang melibatkan fitur deret waktu guna menghasilkan prediksi yang lebih komprehensif. Selain itu, penambahan variabel makroekonomi dan geopolitik sebagai fitur model juga disarankan agar prediksi lebih kontekstual dan akurat terhadap dinamika pasar.

REFERENCES

- [1] N. T. T. Binh, "Comparative Analysis of Gold, Art, and Wheat as Inflation Hedges," *J. Risk Financ. Manag.*, vol. 17, no. 7, 2024, doi: 10.3390/jrfm17070270.
- [2] R. Msi, "Gold Investment Decisions to Anticipate Various Economic Problems in the Future (Study of Generation Z Mothers in Makassar City)," *Mediterr. J. Soc. Sci.*, vol. 15, no. 5, p. 84, 2024, doi: 10.36941/mjss-2024-0047.
- [3] B. N. Vu and H. Y. Hoang, "Factors influencing gold demand: Evidence from developing countries," *J. Soc. Sci. Humanit.*, vol. 66, no. 3, 2024, doi: 10.31276/VMOSTJOSSH.2024.0011.
- [4] S. Pachiyappan, G. Chandrakala, A. Mishra, K. Keerthana, S. Kumari, and T. Verma, "Golden Insights: Analyzing the Influence of Economic Indicators on Sovereign Gold Bond Performance in India," *Int. Res. J. Multidiscip. Scope*, vol. 6, no. 2, pp. 564–576, 2025, doi: 10.47857/irjms.2025.v06i02.03155.
- [5] A. Yavuz and S. Eken, "Gold Returns Prediction: Assessment based on Major Events," *EAI Endorsed Trans. Scalable Inf. Syst.*, vol. 10, no. 5, pp. 1–10, 2023, doi: 10.4108/eetsis.3323.
- [6] S. S. Bannerjee, R. Pillai, M. I. Tabash, and M. S. M. Al-Absy, "Unveiling Inter-Market Reactions to Different Asset Classes/Commodities Pre- and Post-COVID-19: An Exploratory Qualitative Study," *Economies*, vol. 13, no. 3, pp. 1–24, 2025, doi: 10.3390/economies13030066.
- [7] R. Anggara, "Investing In Gold In Times Of Global Crisis," *Int. Asia Law Money Laund.*, vol. 1, no. 4, pp. 228–233, 2022, doi: 10.59712/iaml.v1i4.43.
- [8] M. F. Ghazali, H. H. Lean, and Z. Bahari, "Does gold investment offer protection against stock market losses? evidence from five countries," *Singapore Econ. Rev.*, vol. 65, no. 2, pp. 275–301, 2020, doi: 10.1142/S021759081950036X.
- [9] T. Bunnag, "The Importance of Gold's Effect on Investment and Predicting the World Gold Price Using the ARIMA and ARIMA-GARCH Model," *Ekon. J. Econ.*, vol. 2, no. 1, pp. 38–52, 2024, doi: 10.60084/eje.v2i1.155.
- [10] V. Stasytytė, N. Maknickienė, and R. Martinkutė-Kaulienė, "Hedging basic materials equity portfolios using gold futures," *J. Int. Stud.*, vol. 17, no. 2, pp. 132–145, 2024, doi: 10.14254/2071-8330.2024/17-2/7.
- [11] O. Peretz, M. Koren, and O. Koren, "Naive Bayes classifier – An ensemble procedure for recall and precision enrichment," *Eng. Appl. Artif. Intell.*, vol. 136, no. March, 2024, doi: 10.1016/j.engappai.2024.108972.
- [12] H. S. Putra, Vihi Atina, and Dwi Hartanti, "Application of Naive Bayes Algorithm for Sentiment Analysis of Service and Facility Satisfaction (Case Study: PKU Muhammadiyah Sukoharjo General Hospital)," *J. Adv. Inf. Ind. Technol.*, vol. 6, no. 1, pp. 61–72, 2024, doi: 10.52435/jaiit.v6i1.566.
- [13] H. Chen, S. Hu, R. Hua, and X. Zhao, "Improved naive Bayes classification algorithm for traffic risk management," *EURASIP J. Adv. Signal Process.*, vol. 2021, no. 1, 2021, doi: 10.1186/s13634-021-00742-6.
- [14] Rais, A. M. Soleh, and B. Susetyo, "Rotation Double Random Forest Algorithm to Predict The Food Insecurity Status of Households," *J. RESTI*, vol. 8, no. 1, pp. 33–41, 2024, doi: 10.29207/resti.v8i1.5540.
- [15] T. Zhu, "Analysis on the applicability of the random forest," *J. Phys. Conf. Ser.*, vol. 1607, no. 1,

- 2020, doi: 10.1088/1742-6596/1607/1/012123.
- [16] M. E. Saleh and N. Abd-alsabour, "Improved Decision Tree, Random Forest, and XGBoost Algorithms for Predicting Client Churn in the Telecommunications Industry," *Int. J. Adv. Comput. Sci. Appl.*, vol. 15, no. 12, pp. 674–682, 2024, doi: 10.14569/IJACSA.2024.0151268.
- [17] H. Chen, G. Zhang, X. Pan, and R. Jia, "Using dual evolutionary search to construct decision tree based ensemble classifier," *Complex Intell. Syst.*, vol. 9, no. 2, pp. 1327–1345, 2023, doi: 10.1007/s40747-022-00855-x.
- [18] I. Malashin, V. Tynchenko, A. Gantimurov, V. Nelyub, and A. Borodulin, "Applications of Long Short-Term Memory (LSTM) Networks in Polymeric Sciences: A Review," *Polymers (Basel)*, vol. 16, no. 18, pp. 1–44, 2024, doi: 10.3390/polym16182607.
- [19] B. S. S. Ulama, R. A. Putra, F. Hibatullah, M. R. Habibi, and M. A. Nafis, "Comparison of Artificial Neural Network and Long Short-Term Memory for Modelling Crude Palm Oil Production in Indonesia," *Int. J. Adv. Comput. Sci. Appl.*, vol. 16, no. 1, pp. 741–747, 2025, doi: 10.14569/IJACSA.2025.0160171.
- [20] H. Jdi and N. Falih, "Comparison of time series temperature prediction with auto-regressive integrated moving average and recurrent neural network," *Int. J. Electr. Comput. Eng.*, vol. 14, no. 2, pp. 1770–1778, 2024, doi: 10.11591/ijece.v14i2.pp1770-1778.
- [21] I. Abu-Doush, B. Ahmed, M. A. Awadallah, M. A. Al-Betar, and A. R. Rababaah, "Enhancing multilayer perceptron neural network using archive-based harris hawks optimizer to predict gold prices," *J. King Saud Univ. - Comput. Inf. Sci.*, vol. 35, no. 5, p. 101557, 2023, doi: 10.1016/j.jksuci.2023.101557.
- [22] X. Zhao and P. Keikhosrokiani, "Sales prediction and product recommendation model through user behavior analytics," *Comput. Mater. Contin.*, vol. 70, no. 2, pp. 3855–3874, 2022, doi: 10.32604/cmc.2022.019750.
- [23] A. C. Dewi, A. Hermawan, and D. Avianto, "Classification of Customers' Repeat Order Probability Using Decision Tree, Naïve Bayes and Random Forest," *J. Pilar Nusa Mandiri*, vol. 20, no. 1, pp. 52–59, 2024, doi: 10.33480/pilar.v20i1.5243.
- [24] K. Nurfibia, "Sentiment Analysis of Skincare Products Using the Naive Bayes Method," *J. Inf. Syst. Informatics Vol.*, vol. 6, no. 3, pp. 1663–1676, 2024, doi: 10.51519/journalisi.v6i3.817.
- [25] A. Ecevit, İ. Öztürk, M. Dağ, and T. Özcan, "Short-Term Sales Forecasting Using LSTM and Prophet Based Models in E-Commerce," *Acta Infologica*, vol. 0, no. 0, pp. 0–0, 2023, doi: 10.26650/acin.1259067.
- [26] T. de Castro Moraes, X. M. Yuan, and E. P. Chew, "Hybrid convolutional long short-term memory models for sales forecasting in retail," *J. Forecast.*, vol. 43, no. 5, pp. 1278–1293, 2024, doi: 10.1002/for.3073.
- [27] R. Pugliese, S. Regondi, and R. Marini, "Machine learning-based approach: Global trends, research directions, and regulatory standpoints," *Data Sci. Manag.*, vol. 4, no. November, pp. 19–29, 2021, doi: 10.1016/j.dsm.2021.12.002.
- [28] J. Jabbar, "Sistem Informasi Stok Barang Menggunakan Metode Clustering Kmeans (Studi Kasus Rmd Store)," *INFOTECH J.*, vol. 8, no. 1, pp. 70–75, 2022, doi: 10.31949/infotech.v8i1.2280.
- [29] A. Agung, A. Daniswara, I. Kadek, and D. Nuryana, "Data Preprocessing Pola Pada Penilaian Mahasiswa Program Profesi Guru," *J. Informatics Comput. Sci.*, vol. 05, pp. 97–100, 2023.
- [30] S. Wiguna, S. P. A. Alkadri, and Istikomah, "Prediksi Harga Smartphone Berdasarkan Fitur Smartphone Dengan Random Forest Regression," *Kohesi J. Multidisiplin Saintek*, vol. 6, no. 9, pp. 1–14, 2025, [Online]. Available: <https://ejournal.warunayama.org/kohesi>.
- [31] Slamet Kusworo, Nugroho Adhi Santoso, and Rifki Dwi Kurniawan, "Prediksi Nilai Akhir Semester Siswa Menggunakan Algoritma Random Forest," *J. Sains dan Ilmu Terap.*, vol. 7, no. 2, pp. 246–256, 2024, doi: 10.59061/jsit.v7i2.918.
- [32] D. Arianyah and I. V. Paputungan, "Jurnal Sains, Nalar, dan Aplikasi Teknologi Informasi," *J. Sains, Nalar, dan Apl. Teknol. Inf.*, vol. 3, no. 2, pp. 50–57, 2024, doi: 10.20885/snati.v4.i2.38599.
- [33] F. S. Pamungkas, B. D. Prasetya, and I. Kharisudin, "Perbandingan Metode Klasifikasi Supervised Learning pada Data Bank Customers Menggunakan Python," *Prism. Pros. Semin. Nas. Mat.*, vol. 3, pp. 692–697, 2020, [Online]. Available:

-
- <https://journal.unnes.ac.id/sju/index.php/prisma/article/view/37875>.
- [34] U. L. O. P. N. B. P. Shelke and I. Information Technology, Vishwakarma Institute of Information Technology, Pune, "Gold Price Prediction using Random Forest Algorithm," *Int. Conf. Appl. Artif. Intell. Comput.*, vol. 3, 2024, [Online]. Available: doi: 10.1109/ICAAIC60222.2024.10575316.
- [35] L. Langsetmo *et al.*, "Advantages and Disadvantages of Random Forest Models for Prediction of Hip Fracture Risk Versus Mortality Risk in the Oldest Old," *JBMR Plus*, vol. 7, no. 8, pp. 1–11, 2023, doi: 10.1002/jbm4.10757.
- [36] S. W. Irwan Syahputra, Muhammad Iqbal, "ANALYSIS OF DATA MINING METHODS IN PALM OIL SALES PREDICTION USING RANDOM FOREST AND XGBOOST (CASE STUDY : PT GEKA ADHI RAKSA)," *Jatilima J. Multimed. Dan Teknol. Inf.*, vol. 07, no. 02, pp. 68–79, 2025, [Online]. Available: doi.org/10.54209/jatilima.v7i02.1346.
- [37] D. T. Hermanto, A. Setyanto, and E. T. Luthfi, "Algoritma LSTM-CNN untuk Binary Klasifikasi dengan Word2vec pada Media Online," *Creat. Inf. Technol. J.*, vol. 8, no. 1, p. 64, 2021, doi: 10.24076/citec.2021v8i1.264.