

Evaluasi Komparatif Random Forest, XGBoost, LightGBM, dan K-Nearest Neighbors untuk Prediksi Cuaca di Kota Semarang

Maulana Wahyu Ibrahim*, Wahyu Aji Eko Prabowo

Fakultas Ilmu Komputer, Teknik Informatika, Universitas Dian Nuswantoro, Semarang, Indonesia

Email: ¹*111202214086@mhs.dinus.ac.id, ²prabowo@dsn.dinus.ac.id

Email Penulis Korespondensi: 111202214086@mhs.dinus.ac.id

Submitted 09-01-2026; Accepted 24-02-2026; Published 28-02-2026

Abstrak

Prediksi cuaca yang akurat memiliki peranan penting dalam membantu keputusan strategis di berbagai bidang, dari pertanian hingga penanganan bencana. Namun, ada tantangan mendasar dalam menciptakan model prediksi otomatis yaitu sifat *dataset* meteorologi yang sering kali tidak seimbang dalam distribusi kelas. Fenomena ini membuat algoritma pembelajaran mesin konvensional cenderung condong pada kelas yang dominan dan kurang mampu mendeteksi kelas yang jarang muncul (hujan), yang terlihat dari rendahnya nilai sensitivitas. Penelitian ini bertujuan untuk mengatasi masalah bias ini dan meningkatkan presisi klasifikasi curah hujan harian menggunakan pendekatan komparatif dengan empat algoritma: *Random Forest*, *K-Nearest Neighbor* (KNN), *LightGBM*, dan *XGBoost*. Sebagai metode utama untuk mengatasi ketidakseimbangan data, teknik *Synthetic Minority Over-sampling Technique* (SMOTE) diterapkan untuk menghasilkan sampel baru pada kelas yang kurang terwakili. Kinerja model dievaluasi secara menyeluruh dengan menggunakan *confusion matrix*, strategi *One-vs-Rest* (OvR), dan metrik evaluasi konvensional. Hasil eksperimen pada model *baseline* menunjukkan kegagalan deteksi kelas minoritas dengan nilai *Recall* dan *F1-Score* yang sangat rendah ($< 0,30$). Penerapan SMOTE terbukti berhasil meningkatkan *Recall* dan *F1-Score* dengan signifikan jika dibandingkan dengan kondisi tanpa SMOTE. *LightGBM* yang menggunakan SMOTE tercatat sebagai model paling unggul yang berhasil menyeimbangkan semua matriks evaluasi, dengan *F1-Score* mencapai 0,870 dan AUC ROC 0,976. Pencapaian ini menunjukkan bahwa penggabungan metode *ensemble* dengan teknik *resampling* adalah solusi yang efektif untuk menciptakan sistem prediksi cuaca yang kuat dan akurat.

Kata Kunci: Prediksi Cuaca; Machine Learning; Imbalanced data; SMOTE; Evaluasi Kinerja Model

Abstract

Accurate weather predictions play an important role in assisting strategic decisions in various fields, from agriculture to disaster management. However, there is a fundamental challenge in creating automatic prediction models, namely the nature of meteorological datasets, which are often imbalanced in class distribution. This phenomenon causes conventional machine learning algorithms to favor the dominant class and be less capable of detecting the rare class (rain), as seen in the low sensitivity values. This study aims to overcome this bias problem and improve the accuracy of daily rainfall classification using a comparative approach with four algorithms: *Random Forest*, *K-Nearest Neighbor* (KNN), *LightGBM*, and *XGBoost*. As the main method to overcome data imbalance, the *Synthetic Minority Over-sampling Technique* (SMOTE) was applied to generate new samples in the underrepresented class. Model performance was evaluated comprehensively using a confusion matrix, *One-vs-Rest* (OvR) strategy, and conventional evaluation metrics. The results of the experiments on the baseline model showed a failure to detect the minority class with very low Recall and F1-Score values (< 0.30). The application of SMOTE was proven to significantly improve Recall and F1-Score compared to the SMOTE. *LightGBM* using SMOTE was recorded as the most superior model that successfully balanced all evaluation metrics.

Keywords: Weather Predictions; Machine Learning; Imbalanced data; SMOTE; Model Performance Evaluation

1. PENDAHULUAN

Prediksi cuaca merupakan bidang penelitian yang krusial karena peranannya yang signifikan dalam mendukung produktivitas sektor pertanian serta upaya mitigasi dan pengurangan risiko bencana. Karena pola cuaca memiliki sifat yang tidak pasti dan kompleks, maka diperlukan metode yang mampu mengolah data dalam jumlah besar serta menghadapi variasi yang tinggi agar hasil prediksi bisa akurat dan dapat diandalkan [1], [2]. Model *numerical weather prediction* (NWP) masih menjadi fondasi dalam prakiraan cuaca operasional, tetapi membutuhkan kapasitas komputasi yang besar dan sering kali belum cukup akurat pada skala lokal tanpa adanya koreksi statistik. Karena alasan tersebut, teknik *data-driven* seperti algoritma pembelajaran mesin berbasis *ensemble* telah berkembang pesat sebagai pendukung atau bentuk *hybrid* terhadap NWP untuk meningkatkan tingkat akurasi prediksi pada skala lokal dan menengah [1], [3]. Namun, dalam beberapa dekade terakhir, kemajuan teknologi *machine learning* (ML) telah memberikan peluang besar untuk meningkatkan keakuratan prediksi cuaca. Metode ini memanfaatkan data historis dan data secara *real time* dari berbagai sumber, seperti sensor IoT (*Internet of Things*), gambar satelit, serta data meteorologi konvensional [4]. Sebuah penelitian menyimpulkan bahwa “Model berbasis data yang didukung oleh *machine learning* menunjukkan potensi yang sangat besar dalam peramalan cuaca” [5]. Lebih lanjut, survei terkini menegaskan bahwa model *machine learning* untuk prakiraan cuaca “Menunjukkan kemampuan luar biasa dalam mengelola *dataset* kompleks dan berdimensi tinggi serta memanfaatkan volume besar data historis dan *real time*” [6].

Model pembelajaran mesin seperti *Random Forest* (RF), *Extreme Gradient Boosting* (XGBoost), dan *Light Gradient Boosting Machine* (*LightGBM*) sering dipilih karena kemampuan mereka dalam mengolah data yang memiliki banyak dimensi dan menganalisis pola-pola non-linear pada data cuaca. *Random Forest* dikenal baik dalam menangani data yang tidak bersih dan memberikan hasil yang mudah diinterpretasikan, sedangkan *XGBoost* dan *LightGBM* menawarkan kecepatan pemrosesan yang lebih cepat serta tingkat akurasi yang kompetitif berkat metode *boosting* yang membantu memperbaiki kesalahan prediksi secara bertahap [1], [7], [8]. Selain memilih model yang tepat, penggunaan

teknik evaluasi yang baik juga sangat penting untuk memastikan kualitas model prediksi. Salah satu metrik yang sering digunakan adalah *Area Under the Receiver Operating Characteristic Curve* (AUC ROC), yang digunakan untuk mengukur kemampuan model dalam membedakan antara kelas positif dan negatif secara efektif. AUC ROC memberikan gambaran yang objektif mengenai kinerja model klasifikasi, terutama dalam menangani data cuaca yang tidak seimbang akibat fenomena langka seperti hujan lebat atau badai [9].

Kendati algoritma pembelajaran mesin modern menawarkan arsitektur yang canggih, efektivitas prediksinya kerap kali dibatasi oleh tantangan fundamental yang melekat pada struktur data meteorologi. Ketidakseimbangan data menjadi tantangan besar dalam memprediksi cuaca, karena peristiwa ekstrem cenderung lebih sedikit dibandingkan kondisi cuaca normal. Untuk mengatasinya, teknik *Synthetic Minority Over sampling Technique* (SMOTE) digunakan untuk menghasilkan data minoritas secara sintesis. Hal ini membantu model belajar dengan lebih baik dan tidak terbiasa hanya pada kelas mayoritas. Penerapan SMOTE telah terbukti meningkatkan kemampuan model seperti *Random Forest*, *XGBoost*, dan *LightGBM* dalam memprediksi cuaca jarang terjadi dengan tingkat presisi yang lebih baik [10], [11].

Pada penelitian yang dilakukan oleh Syahreza dkk. yang mengkaji performa *Random Forest*, *XGBoost*, dan *LightGBM* menunjukkan hasil yang memuaskan. Dengan menggunakan data meteorologi yang diperoleh dari sensor IoT, model *XGBoost* menunjukkan peningkatan signifikan dalam akurasi prediksi suhu dan curah hujan dibandingkan metode tradisional [1]. Pada penelitian yang dilakukan oleh Dwiyantri dan Prianto, melaporkan kemampuan *Random Forest* mencapai nilai ROC AUC lebih dari 0,9 untuk memprediksi cuaca di kota Jakarta, menunjukkan bahwa model ini cukup andal dalam konteks kota yang kompleks [2]. Selain itu, Pringandana dan Kusnawi, membandingkan beberapa algoritma, termasuk *XGBoost* dan *LightGBM*, menunjukkan bahwa kombinasi metode *boosting* dengan pengelolaan data tidak seimbang secara efektif meningkatkan kualitas prediksi curah hujan dalam upaya mengurangi risiko bencana kebakaran hutan dan lahan [12].

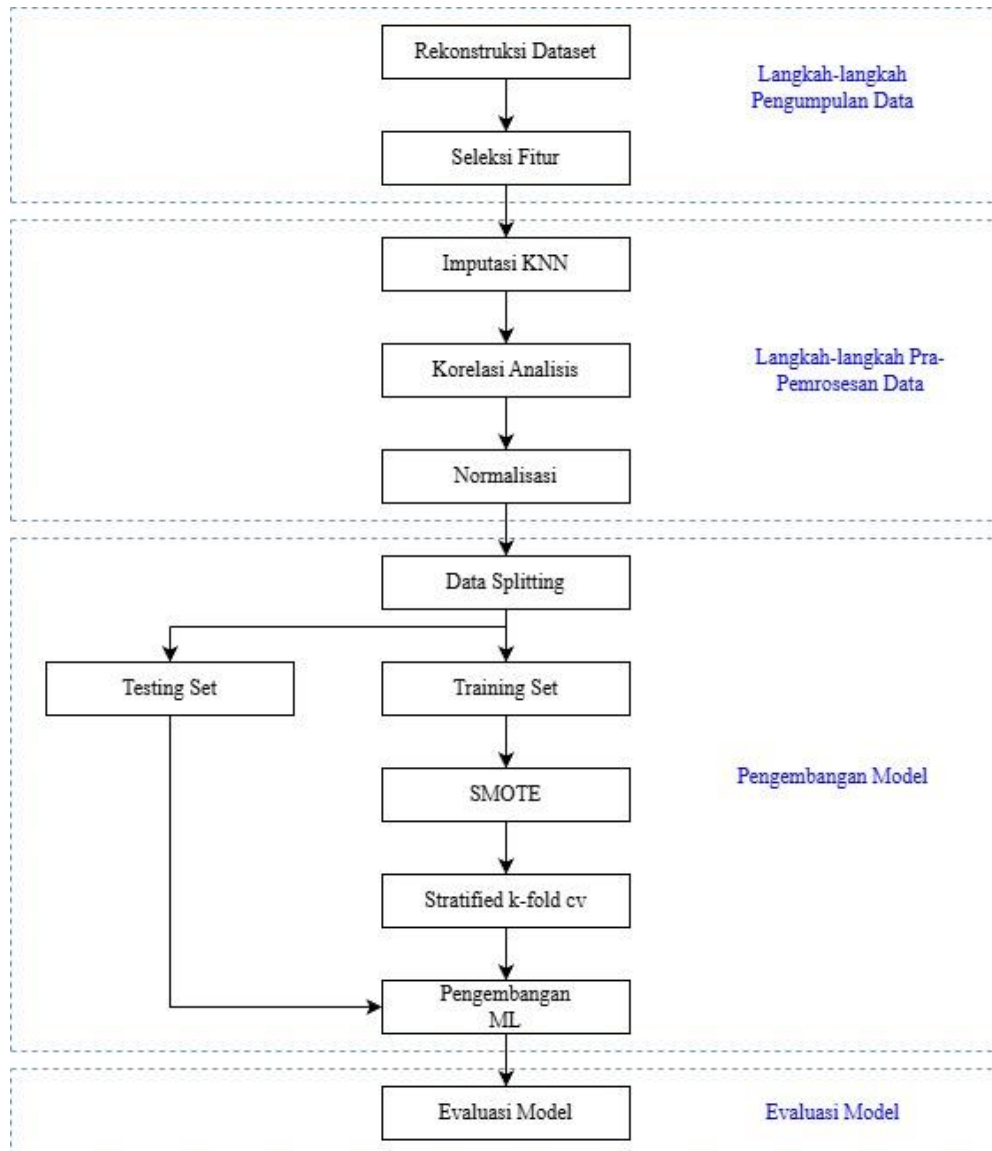
Meskipun demikian, terdapat kesenjangan yang signifikan dalam literatur yang ada. Mayoritas penelitian tersebut cenderung berfokus pada metrik akurasi umum dan belum secara komprehensif menangani masalah ketidakseimbangan kelas (*imbalanced data*) yang ekstrem pada data cuaca lokal. Penggunaan metrik akurasi pada data yang tidak seimbang sering kali menghasilkan performa semu (bias pada kelas mayoritas/cuaca normal) dan gagal mendeteksi fenomena langka namun krusial, seperti cuaca ekstrem. Selain itu, belum banyak studi yang membandingkan secara simultan efektivitas metode *ensemble* berbasis *tree* (RF, *XGBoost*, *LightGBM*) dengan metode KNN yang dikombinasikan secara spesifik dengan teknik SMOTE untuk merekonstruksi keseimbangan data.

Oleh karena itu, pendekatan yang lebih komprehensif diterapkan dalam studi ini dengan mengintegrasikan empat algoritma, yaitu *Random Forest*, *XGBoost*, *LightGBM*, dan KNN. Keempat metode tersebut dipilih berdasarkan keandalannya dalam memproses *dataset* bervolume besar serta kemampuannya mengidentifikasi pola non-linear yang sering kali luput dari jangkauan metode statistik konvensional. Meskipun banyak studi terdahulu telah mengeksplorasi prediksi cuaca, kesenjangan penelitian yang krusial terletak pada minimnya penanganan efektif terhadap ketidakseimbangan distribusi kelas (*imbalanced data*). Mayoritas model yang ada cenderung bias pada kondisi cuaca dominan (cerah/berawan) sehingga menghasilkan akurasi semu yang gagal mendeteksi fenomena cuaca ekstrem. Guna mengisi celah tersebut, teknik *Synthetic Minority Oversampling Technique* (SMOTE) diimplementasikan secara spesifik untuk merekonstruksi keseimbangan data, memastikan model memiliki sensitivitas yang setara terhadap seluruh kelas. Validasi kualitas model pun tidak lagi bertumpu pada akurasi semata, melainkan memprioritaskan metrik AUC ROC untuk mengukur efektivitas pemisahan antar kelas pada situasi data yang kompleks. Melalui strategi ini, hasil yang diharapkan tidak hanya berupa model yang adaptif terhadap variabilitas iklim Indonesia, tetapi juga kontribusi nyata bagi penguatan sistem peringatan dini, mitigasi bencana, serta optimalisasi perencanaan di berbagai sektor strategis.

2. METODOLOGI PENELITIAN

2.1 Tahapan Penelitian

Alur penelitian ini dirancang secara sistematis untuk memastikan setiap tahapan pemrosesan data dan pengembangan model berjalan secara terstruktur. Kerangka kerja penelitian ini terdiri dari empat tahapan utama yaitu Langkah-langkah Pengumpulan Data, Langkah-langkah Pra-Pemrosesan Data, Pengembangan Model, dan Evaluasi Model. Secara rinci, tahapan-tahapan yang dilakukan dalam penelitian ini diilustrasikan pada Gambar 1. Pada gambar 1 menunjukkan bahwa proses diawali dengan tahap pengumpulan data yang mencakup rekonstruksi *dataset* dan seleksi fitur untuk menyaring variabel yang paling relevan. Setelah data terkumpul, tahapan berlanjut ke pra-pemrosesan data untuk meningkatkan kualitas input model. Pada tahap ini, dilakukan penanganan data yang hilang (*missing values*) menggunakan metode Imputasi KNN, dilanjutkan dengan analisis korelasi untuk melihat hubungan antar fitur, serta proses normalisasi agar data memiliki skala yang seragam. Setelah data siap, proses masuk ke tahap pengembangan model. Seperti terlihat pada diagram, dilakukan *data splitting* untuk memisahkan data menjadi *training set* dan *testing set*. Hal yang krusial pada tahap ini adalah penerapan SMOTE yang dilakukan secara khusus hanya pada *training set* untuk menyeimbangkan kelas data tanpa menyebabkan kebocoran data (*data leakage*) ke data uji. Selanjutnya, diterapkan *Stratified k-fold cross-validation* untuk memastikan model yang dikembangkan (*Random Forest*, *XGBoost*, *LightGBM*, dan KNN) memiliki performa yang stabil. Rangkaian proses ini bermuara pada tahap akhir, yaitu evaluasi model, di mana performa prediksi diukur menggunakan metrik evaluasi yang komprehensif.



Gambar 1. Alur Penelitian

2.2 Rekonstruksi Dataset

Data penelitian bersumber dari Badan Meteorologi, Klimatologi, dan Geofisika (BMKG) yang difokuskan secara spesifik pada wilayah Kota Semarang [13]. Data tersebut secara spesifik mewakili kondisi cuaca di Kota Semarang selama lima tahun terakhir. Dalam struktur *dataset*, variabel-variabel disimpan dalam dua jenis tipe data numerik utama, yaitu Integer (Int) dan *Float*. Tipe data Integer (Int) digunakan untuk merepresentasikan bilangan bulat yang tidak memiliki bagian desimal. Sementara itu, tipe data *Float* digunakan untuk menyimpan bilangan riil dengan presisi desimal, sehingga bisa mencatat parameter pengukuran cuaca secara lebih detail. Untuk memberikan gambaran spesifik mengenai karakteristik data yang digunakan, rincian parameter beserta tipe datanya dirangkum dalam Tabel 1. Seperti yang ditunjukkan pada tabel tersebut, struktur data dikelompokkan menjadi dua tipe numerik, yaitu *Float* dan *Integer*. Tipe data *Float* diterapkan pada parameter yang memerlukan presisi desimal tinggi untuk menangkap fluktuasi iklim mikro, meliputi variabel Suhu (Maksimum Tx, Minimum Tn, Rata-rata Tavg), Kelembapan (RH_avg), Lama Penyinaran (ss), dan Curah Hujan (RR). Sementara itu, tipe data Integer digunakan untuk merepresentasikan nilai bulat, khususnya pada parameter Kecepatan Angin (Sf_x dan ff_avg) dan Arah Angin (DDD_x). Berikut adalah tabel informasi *dataset*.

Tabel 1. Informasi *Dataset*

Parameter	Tipe Data
Suhu Maksimum (Tx)	Float
Suhu Minimum (Tn)	Float
Suhu Rata-rata (Tavg)	Float
Kelembapan Rata-rata (RH_avg)	Float

Lama Penyiraman Matahari (ss)	Float
Kecepatan Angin Maksimum (ff_x)	Int
Kecepatan Angin Rata-rata (ff_avg)	Int
Arah Kecepatan Angin Maksimum (DDD_x)	Int
Curah Hujan (RR)	Float

2.3 Pra-Pemrosesan Data

Tahap pra-pemrosesan data dilakukan secara sistematis untuk menyiapkan *dataset* secara optimal sebelum pemodelan. Proses ini diawali dengan transformasi data untuk mengonversi seluruh variabel menjadi format numerik, diikuti oleh penanganan *missing values* menggunakan metode imputasi *K-Nearest Neighbor* (KNN) yang memanfaatkan kesamaan karakteristik antar data tetangga. Langkah krusial selanjutnya adalah pengelompokan variabel target curah hujan ke dalam enam kategori kelas sesuai standar operasional BMKG. Rincian batasan nilai dan label kategori tersebut disajikan dalam Tabel 2.

Tabel 2. Klasifikasi curah Hujan

Tingkat Curah Hujan	Keterangan
0 mm/hari	Berawan
0,5 – 20 mm/hari	Hujan Ringan
20 – 50 mm/hari	Hujan Sedang
50 – 100 mm/hari	Hujan Lebat
100 – 150 mm/hari	Hujan Sangat Lebat
>150 mm/hari	Hujan Ekstrem

Berdasarkan informasi pada Tabel 2, variabel target diklasifikasikan berdasarkan intensitas curah hujan harian. Kategori dimulai dari kondisi "Berawan" (0 mm/hari), diikuti dengan "Hujan Ringan" (0,5 – 20 mm/hari), dan "Hujan Sedang" (20 – 50 mm/hari). Intensitas yang lebih tinggi dikelompokkan menjadi "Hujan Lebat" (50 – 100 mm/hari) dan "Hujan Sangat Lebat" (100 – 150 mm/hari). Kategori terakhir adalah "Hujan Ekstrem", yang mencakup kejadian hujan dengan intensitas melebihi 150 mm/hari. Setelah klasifikasi target terbentuk, tahap selanjutnya berfokus pada fitur *input*. Mengingat *dataset* meteorologi memiliki variasi satuan dan rentang nilai yang sangat tidak seimbang seperti variabel arah angin (DDD_x) yang mencapai ratusan, berbeda dengan variabel curah hujan (RR) atau penyinaran matahari (ss) yang bernilai kecil—maka proses normalisasi menjadi sangat penting. Untuk menangani ketimpangan ini, metode Min-Max Scaling digunakan untuk menyamakan skala semua fitur dalam rentang seragam (0 hingga 1). Transformasi ini bertujuan mengeliminasi dominasi fitur bernilai besar, sehingga setiap variabel memberikan kontribusi bobot yang setara dalam mempercepat konvergensi algoritma.

2.3.1 Seleksi Fitur

Dataset awal mencakup serangkaian variabel meteorologi harian yang terdiri dari suhu minimum (Tn), suhu maksimum (Tx), suhu rata-rata (Tavg), kelembapan rata-rata (RH_avg), lama penyiraman matahari (ss), kecepatan angin maksimum (ff_x), kecepatan angin rata-rata (ff_avg), arah angin kecepatan maksimum (DDD_x), dan curah hujan (RR). Seleksi fitur dilakukan dengan meninjau kelengkapan data pada setiap kolom dan memprioritaskan variabel yang memiliki dampak meteorologis langsung terhadap pembentukan hujan berdasarkan studi literatur terkait. Fitur yang tidak relevan, seperti atribut tanggal atau ID stasiun, serta fitur dengan jumlah *missing values* yang melebihi ambang batas toleransi, dihapus dari *dataset* untuk menjaga kualitas model.

2.3.2 Imputasi KNN

Mengingat data cuaca sering kali memuat nilai yang hilang (*missing values*) akibat kendala teknis pada sensor pencatatan, metode imputasi KNN diterapkan untuk melengkapi kekosongan tersebut. Algoritma ini bekerja dengan mengestimasi nilai yang hilang berdasarkan kemiripan fitur dengan data tetangga terdekat, sehingga sangat efektif dalam menangani pola hubungan antar variabel cuaca yang bersifat non-linear dan dinamis. Berbeda dengan metode rata-rata sederhana, pendekatan KNN terbukti mampu mempertahankan variabilitas data asli dan menjaga kestabilan distribusi statistik dalam *dataset* klimatologi [14].

2.3.3 Korelasi Analisis dan Normalisasi

Analisis korelasi digunakan sebagai langkah awal untuk mendeteksi adanya ketergantungan antar variabel independen dalam data cuaca. Tujuan utamanya adalah memilih variabel-variabel yang memiliki hubungan linier kuat, sehingga hanya variabel-variabel yang memberikan informasi baru yang dipertahankan sebagai masukan dalam model. Menghilangkan variabel-variabel yang sudah terlalu banyak dan tidak perlu sangat penting agar model tetap stabil saat dilatih, serta mencegah terjadinya kompleksitas berlebihan yang bisa mengurangi kemampuan model dalam memprediksi data baru [15], [16].

2.4 Data Splitting

Dataset dipartisi menjadi dua *subset* independen dengan rasio 80:20, di mana 80% dialokasikan sebagai data latih (*training set*) dan 20% sisanya sebagai data uji (*testing set*). Tujuan pembagian ini adalah agar model dapat dievaluasi dengan data yang belum pernah dilihat sebelumnya, sehingga hasil pengecekan lebih adil dan bisa menunjukkan sejauh mana kemampuan model dalam menangani data nyata [17], [18]. Hal yang krusial dalam pembagian ini adalah memastikan bahwa data uji dipisahkan secara ketat dan sama sekali tidak dilibatkan ataupun 'dilihat' oleh model selama proses pelatihan berlangsung. Isolasi total ini bertujuan untuk menjamin objektivitas evaluasi, sehingga metrik performa yang dihasilkan benar-benar merefleksikan kemampuan generalisasi model terhadap data baru (*unseen data*) yang belum pernah dikenali sebelumnya.

2.5 Pengembangan Model *Machine Learning*

2.5.1 SMOTE

Dalam *dataset* curah hujan yang digunakan dalam penelitian meteorologi, distribusi kelas intensitas hujan umumnya tidak merata. Jumlah data pada kelas hujan ringan jauh lebih banyak dibandingkan kelas hujan sedang atau lebat. Hal ini menyebabkan model prediksi lebih cenderung mengenali hujan ringan, sehingga kemampuan model dalam mendeteksi hujan sedang atau lebat menjadi kurang baik [19], [20]. Untuk mengatasi masalah tersebut, metode *Synthetic Minority Oversampling Technique* (SMOTE) diterapkan pada *dataset* pelatihan. SMOTE bekerja dengan membuat data baru dari kelas minoritas melalui interpolasi antar data yang sudah ada, sehingga distribusi data antar kelas menjadi lebih seimbang. Dengan cara ini, model mendapatkan gambaran yang lebih baik untuk setiap kelas dan dapat membuat prediksi yang lebih adil serta akurat terhadap berbagai jenis hujan. Tabel 3 akan menampilkan perbandingan distribusi data sebelum dan sesudah penerapan SMOTE untuk menilai efek oversampling terhadap ketidakseimbangan kelas curah hujan.

Tabel 3. Perbandingan Distrubusi Data Sebelum dan Sesudah SMOTE

Label Hujan	Normal	SMOTE
Berawan	960	768
Hujan Ringan	670	768
Hujan Sedang	149	768
Hujan Lebat	44	768
Hujan Sangat Lebat	4	768

Berdasarkan data yang tersaji pada Tabel 3, terlihat ketimpangan yang sangat ekstrem pada kondisi awal (kolom "Normal"). Kelas "Berawan" mendominasi dataset dengan jumlah 960 sampel, sangat kontras dibandingkan dengan kelas "Hujan Sangat Lebat" yang hanya memiliki 4 sampel. Ketidakseimbangan ini berpotensi membuat model gagal mempelajari pola hujan ekstrem. Namun, setelah proses *resampling* SMOTE, distribusi data berhasil disetarakan. Terlihat bahwa seluruh kelas, mulai dari "Berawan" hingga "Hujan Sangat Lebat", kini memiliki jumlah sampel yang seragam, yaitu masing-masing 768 data. Pemerataan ini memastikan bahwa setiap kategori hujan memiliki proporsi representasi yang seimbang dalam proses pelatihan, sehingga model dapat mempelajari karakteristik setiap kelas cuaca secara adil dan objektif tanpa bias terhadap kelas mayoritas.

2.5.2 Pengembangan Model dan *Hyperparameter Tuning*

Model yang dikembangkan dalam penelitian ini menggunakan empat algoritma berbasis *tree ensemble* dan satu *supervised learning*, yaitu *Random Forest*, XGBoost, LightGBM dan KNN. Keempat algoritma ini terbukti efektif dalam berbagai studi prediksi cuaca karena mampu menangani data yang tidak linear dan bisa mengeksplorasi hubungan antar fitur yang kompleks secara efisien. *Random Forest* menggunakan teknik *bagging* dengan menggabungkan banyak pohon keputusan untuk meningkatkan stabilitas dan akurasi model. Sementara itu, XGBoost dan LightGBM menggunakan teknik *boosting* yang sangat adaptif, terutama untuk data yang memiliki pola rumit atau kasus ketidakseimbangan kelas masalah yang sering ditemukan pada data cuaca harian [12], [21]. Sementara itu, dalam algoritma KNN dilakukan pemilihan nilai *k* yang terbaik serta penentuan jenis metrik jarak, seperti jarak Euclidean, agar bisa mendapatkan hasil klasifikasi yang akurat dalam memprediksi intensitas hujan berkelas banyak [22]. Pelatihan model dilakukan dengan metode *k-fold cross validation* (*k* = 5), di mana seluruh data dibagi menjadi lima bagian yang berbeda. Pada setiap tahap, data latihan dan data validasi dipertukarkan. Pendekatan ini dilakukan agar model dapat dicek kinerjanya pada berbagai bagian data secara mandiri, sehingga mengurangi bias dalam pelatihan dan meningkatkan kemampuan model dalam memprediksi data nyata. Langkah ini juga membantu mencegah *overfitting* dan memberikan gambaran yang lebih benar mengenai metrik performa seperti akurasi, *recall*, dan AUC-ROC dalam prediksi cuaca multi kelas [12], [23].

2.6 Evaluasi Model

Evaluasi kinerja model dilakukan secara menyeluruh dan tidak hanya berdasarkan metrik akurasi, karena akurasi rata-rata bisa bias terhadap kelas yang dominan. Fokus utama evaluasi adalah mengukur efektivitas penggunaan SMOTE, di mana keberhasilannya ditandai dengan peningkatan signifikan pada nilai *recall* dan *F1-Score* untuk kelas minoritas (hujan) dibandingkan model dasar (tanpa SMOTE). Selain itu, metrik AUC-ROC multikelas serta *Confusion Matrix* digunakan untuk memastikan kemampuan model dalam memisahkan kelas secara menyeluruh. Di akhir proses, model yang dipilih adalah model yang memiliki kinerja seimbang antar kelas serta hasil pengujian yang stabil di semua *fold*

pada validasi silang. Setelah model klasifikasi dibuat, langkah berikutnya adalah mengevaluasi hasil kerjanya. Proses evaluasi ini dilakukan dengan menghitung beberapa metrik seperti presisi, *recall*, F1-score, dan akurasi. Metrik tersebut dihitung berdasarkan nilai-nilai *True Positive* (TP), *True Negative* (TN), *False Positive* (FP), dan *False Negative* (FN) yang diperoleh dari *confusion matrix*. Rumus yang digunakan adalah sebagai berikut:

$$\text{Presisi} = \frac{TP}{TP+FP} \quad (1)$$

$$\text{Recall} = \frac{TP}{TP+FN} \quad (2)$$

$$\text{F1 - Score} = 2 \times \frac{(\text{Presisi} \times \text{Recall})}{(\text{Presisi} + \text{Recall})} \quad (3)$$

$$\text{Akurasi} = \frac{TP+TN}{TP+TN+FP+FN} \quad (4)$$

Selanjutnya, untuk mengetahui seberapa baik model membedakan setiap kelas secara khusus, digunakan *Area Under Curve* (AUC). Nilai AUC untuk kelas ke-k (AUC_k) dihitung dengan mengintegrasikan fungsi *True Positive Rate* (TPR) terhadap *False Positive Rate* (FPR) seperti berikut:

$$AUC_k = \int_0^1 TP R_k(FP R_k) d(FP R_k) \quad (5)$$

Nilai akhir performa model secara keseluruhan ditentukan dengan menghitung rata-rata makro (*Macro Average*) dari nilai AUC seluruh kelas menggunakan persamaan:

$$AUC_{macro} = \frac{1}{K} \sum_{k=1}^K AUC_k \quad (6)$$

Keterangan variabel pada persamaan 1 hingga 6:

TP (*True Positive*) : Jumlah data kelas positif yang diprediksi secara benar.

TN (*True Negative*) : Jumlah data kelas negatif yang diprediksi secara benar.

FP (*False positive*) : Jumlah data negatif yang salah diprediksi sebagai positif (Kesalahan Tipe I).

FN (*False Negative*) : Jumlah data positif yang salah diprediksi sebagai negatif (Kesalahan Tipe II).

TP R_k : *True Positive Rate* untuk kelas k (Sensitivitas).

FP R_k : *False Positive Rate* untuk kelas k (1 - Spesifisitas).

K : Total jumlah kelas kategori dalam dataset

AUC_k : Nilai *Area Under Curve* untuk kelas spesifik ke-k

AUC_{macro} : Rata-rata nilai AUC dari seluruh kelas.

3. HASIL DAN PEMBAHASAN

Bagian ini menyajikan hasil pengujian dari hasil empat algoritma *machine learning* yang digunakan untuk memprediksi hujan. Fokus utama penelitian adalah membandingkan kinerja keempat algoritma sebelum dan sesudah penerapan teknik SMOTE. Tujuan dari teknik ini adalah mengatasi ketidakseimbangan distribusi data antar kelas sentimen minoritas. Selain itu, juga ditampilkan perbandingan distribusi label sebelum dan sesudah diterapkan SMOTE untuk mengilustrasikan perubahan proporsi kelas akibat penerapan teknik SMOTE. Penelitian ini juga mencakup analisis terhadap permasalahan ketidakseimbangan kelas, pengaruh SMOTE, serta alasan mengapa model tertentu memberikan performa lebih tinggi pada kasus prediksi intensitas hujan multikelas.

3.1 Dataset

Dataset yang digunakan berasal dari Badan Meteorologi, Klimatologi, dan Geofisika (BMKG), yaitu data sekunder yang difokuskan secara spesifik untuk merepresentasikan kondisi cuaca di wilayah Kota Semarang selama lima tahun terakhir. *Dataset* ini berisi parameter meteorologi harian seperti suhu (maksimum, minimum, rata-rata), kelembapan, penyinaran matahari, serta kecepatan dan arah angin, dengan setiap entri memiliki target variabel berupa curah hujan. Pada penelitian ini, data dipartisi menggunakan rasio 80:20 yang menghasilkan *subset* independen untuk proses pelatihan dan pengujian guna menjamin objektivitas evaluasi model. Seleksi fitur dilakukan dengan memprioritaskan variabel yang memiliki dampak meteorologis langsung dan menghapus atribut non-informatif untuk menjaga kualitas input. Tabel 4 menyajikan informasi mengenai tipe data dan parameter yang digunakan untuk memperlihatkan konteks awal *dataset* BMKG Kota Semarang ini.

Tabel 4. Dataset Penelitian

Tanggal	TN	TX	TAVG	RH AVG	SS	FF X	DDD X	FF AVG	RR
01/06/2020	24,6	30,4	28	85	-	4	2	100	0
02/06/2020	24,6	31,8	27,7	87	-	4	1	120	0
03/06/2020	25,8	32,2	29	80	1,8	6	3	110	0

3.2 Pra-Pemrosesan

Mengingat adanya data yang hilang akibat kendala teknis pada sensor pencatatan, dilakukan prosedur Imputasi KNN untuk mengisi kekosongan informasi berdasarkan kemiripan pola data tetangga terdekat. Selanjutnya, dilakukan transformasi data menggunakan *Min-Max Scaling* untuk menyeragamkan seluruh fitur ke dalam rentang nilai 0 hingga 1. Tahapan normalisasi ini sangat krusial agar variabel dengan skala besar tidak mendominasi proses perhitungan matematis algoritma, sehingga mempercepat proses konvergensi saat pelatihan model. Setiap entri akan mendapatkan variabel yaitu Label curah hujan ("Berawan" (0 mm/hari), "Hujan Ringan" (0,5 – 20 mm/hari), "Hujan Sedang" (20 – 50 mm/hari), "Hujan Lebat" (50 – 100 mm/hari), dan "Hujan Sangat Lebat" (100 – 150 mm/hari). "Hujan Ekstrem" (>150 mm/hari)). Pada tabel 5 menampilkan dataset setelah dilakukan pra-pemrosesan.

Tabel 5. Hasil Pra-Pemrosesan

Tanggal	TN	TX	TAVG	RH AVG	SS	FF X	DDD X	FF AVG	RR	Kategori
01/06/2020	0,622	0,333	0,443	0,75	0,824	0,181	0,277	0,333	0	Berawan
02/06/2020	0,622	0,444	0,405	0,795	0,771	0,181	0,333	0,166	0	Berawan
03/06/2020	0,755	0,476	0,569	0,636	0,157	0,363	0,305	0,5	0	Berawan

3.3 Implementasi Algoritma Komparatif

Implementasi algoritma dalam penelitian ini mengintegrasikan pendekatan berbasis *tree ensemble* dan *instance-based learning* untuk memetakan pola non-linear pada data meteorologi Kota Semarang. Seluruh tahapan pelatihan dilakukan menggunakan metode *Stratified 5-Fold Cross-Validation* guna memastikan stabilitas performa dan mencegah terjadinya *overfitting* pada setiap lipatan data. Pada algoritma *Random Forest*, model dikonversi menjadi kumpulan pohon keputusan yang dibangun secara independen melalui teknik *bagging* untuk meningkatkan akurasi melalui mekanisme suara terbanyak. Sementara itu, pada pendekatan *XGBoost* dan *LightGBM*, proses optimasi dilakukan secara iteratif menggunakan mekanisme *gradient boosting* yang secara bertahap memperbaiki kesalahan prediksi dari pohon sebelumnya. *XGBoost* difokuskan pada penanganan pola data yang rumit, sedangkan *LightGBM* menerapkan pendekatan *leaf-wise tree growth* yang menawarkan efisiensi komputasi lebih tinggi dalam memproses *dataset* bervolume besar. Di sisi lain, algoritma KNN diterapkan dengan melakukan klasifikasi berdasarkan metrik jarak *Euclidean* untuk menentukan intensitas hujan melalui kedekatan karakteristik antar observasi di dalam ruang fitur. Seluruh konfigurasi algoritma ini dirancang untuk bekerja secara optimal setelah distribusi kelas diseimbangkan melalui teknik SMOTE, sehingga setiap model memiliki sensitivitas yang setara dalam mengenali fenomena cuaca ekstrem.

3.4 Evaluasi Model Tanpa SMOTE

Bagian ini menampilkan hasil analisis performa model dalam skenario pengujian pertama menggunakan data pelatihan asli tanpa teknik penyeimbangan kelas. Penilaian ini krusial untuk mendemonstrasikan masalah bias yang dihadapi algoritma *machine learning* saat berhadapan dengan dominasi kelas mayoritas. Rangkuman hasil penilaian kinerja untuk keempat algoritma disajikan secara rinci dalam Tabel 6, yang mencakup metrik akurasi, *Recall*, F1-Score, presisi, dan AUC-ROC.

Tabel 6. Evaluasi Model pada Data Uji Sebelum Penerapan SMOTE

Model	RF	XGBoost	LGBM	KNN
Akurasi	0,631	0,609	0,609	0,584
Presisi	0,245	0,295	0,276	0,265
Recall	0,273	0,285	0,279	0,267
F1-Score	0,258	0,283	0,274	0,264
AUC-ROC	0,787	0,749	0,762	0,581

Berdasarkan data yang disajikan pada Tabel 6, terlihat bahwa meskipun nilai akurasi berada pada rentang 0,584 hingga 0,631, metrik evaluasi lainnya menunjukkan performa yang sangat rendah. Fenomena ini mengonfirmasi adanya masalah ketidakseimbangan kelas (*imbalanced data*), di mana model memiliki nilai *Recall* dan F1-Score yang rendah, rata-rata di bawah 0,30 untuk semua algoritma. Rendahnya nilai *Recall* pada Tabel 6 misalnya 0,273 pada model *Random Forest* dan 0,267 pada KNN menunjukkan bahwa model gagal mendeteksi sebagian besar kejadian hujan (kelas minoritas) karena algoritma lebih cenderung memprediksi kelas mayoritas yang dominan. Selain itu, nilai AUC-ROC pada Tabel 6 yang bervariasi menunjukkan bahwa tanpa intervensi teknik penyeimbangan data, model belum memiliki kemampuan pemisahan kelas yang optimal, terutama pada algoritma KNN yang hanya mencatat nilai 0,581.

3.5 Evaluasi Model Dengan SMOTE

Setelah pengujian pada model *baseline* dilakukan, hasil evaluasi menunjukkan bahwa keempat algoritma mengalami kesulitan besar dalam mendeteksi kelas cuaca minoritas akibat ketidakseimbangan data. Pada bagian ini, dibahas performa keempat model klasifikasi setelah mengintegrasikan teknik SMOTE untuk menyeimbangkan distribusi kelas. Rangkuman detail mengenai peningkatan metrik evaluasi untuk seluruh algoritma disajikan secara komprehensif pada Tabel 7.

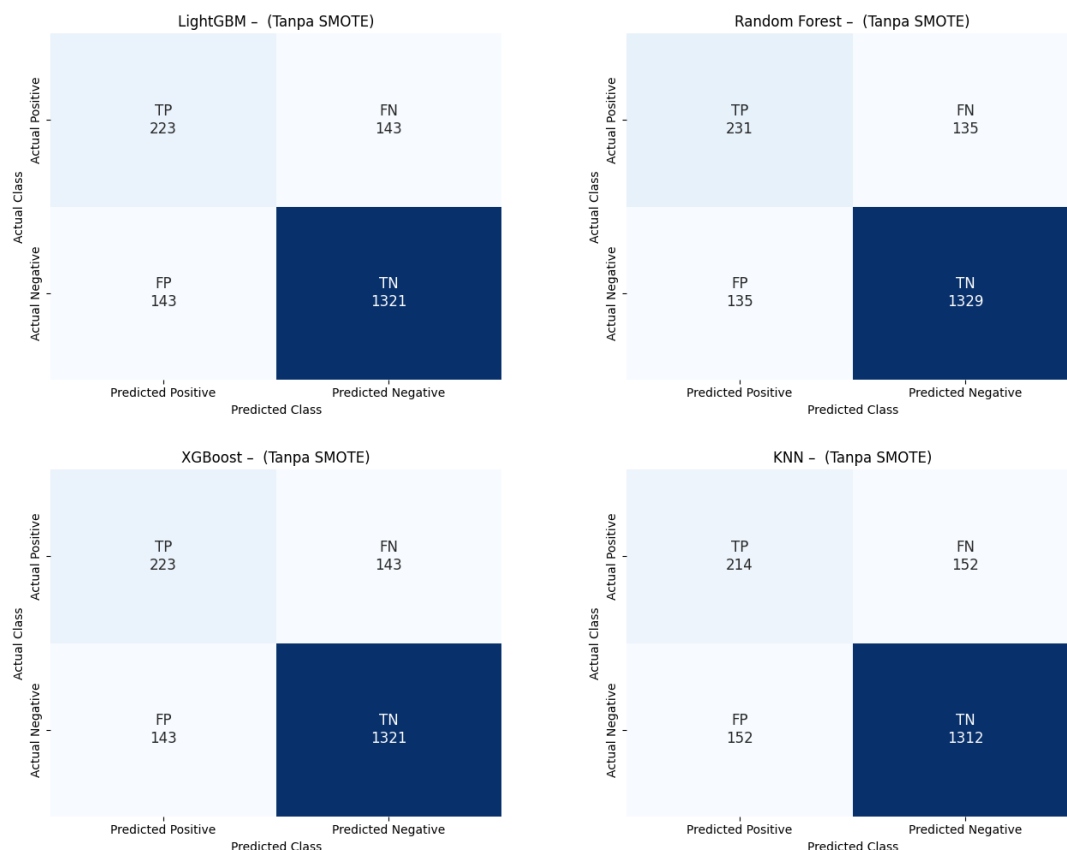
Tabel 7. Evaluasi Model pada Data Uji Setelah Penerapan SMOTE

Model	RF	XGBoost	LGBM	KNN
Akurasi	0,631	0,609	0,609	0,584
Presisi	0,245	0,295	0,276	0,265
Recall	0,273	0,285	0,279	0,267
F1-Score	0,258	0,283	0,274	0,264
AUC-ROC	0,787	0,749	0,762	0,581

Berdasarkan hasil pengujian yang tertera pada Tabel 7, terlihat adanya transformasi performa yang sangat signifikan dibandingkan dengan kondisi sebelum penyeimbangan data. Sebagai contoh, pada algoritma *LightGBM* yang tercatat di Tabel 7, nilai F1-Score melonjak drastis hingga mencapai 0,870 dengan AUC-ROC sebesar 0,976. Peningkatan serupa juga terlihat pada model *XGBoost* dan KNN yang menunjukkan nilai *Recall* pada kisaran 0,87, yang berarti model kini mampu mengenali hampir 87% kejadian hujan yang sebelumnya gagal terdeteksi. Keberhasilan yang ditunjukkan pada Tabel 7 ini membuktikan bahwa penggunaan SMOTE secara efektif mampu menghilangkan bias pada kelas mayoritas, sehingga model memiliki sensitivitas yang seimbang dalam memprediksi berbagai intensitas curah hujan, mulai dari berawan hingga hujan sangat lebat.

3.6 Analisis Confusion Matrix Sebelum SMOTE

Pada bagian ini, analisis terhadap hasil prediksi dilakukan untuk memberikan perspektif yang lebih terperinci tentang performa model dalam setiap kategori intensitas hujan. Untuk mencapai hal tersebut, evaluasi dengan *Confusion Matrix* dilakukan dengan menggunakan strategi *One-vs-Rest* (OvR). Skenario pertama menguji performa model pada *dataset* asli (*baseline*) yang memiliki distribusi kelas tidak seimbang. Representasi visual dari hasil evaluasi baseline ini disajikan pada Gambar 2, yang mencakup hasil dari empat algoritma yaitu *Random Forest*, *XGBoost*, *LightGBM*, dan KNN.



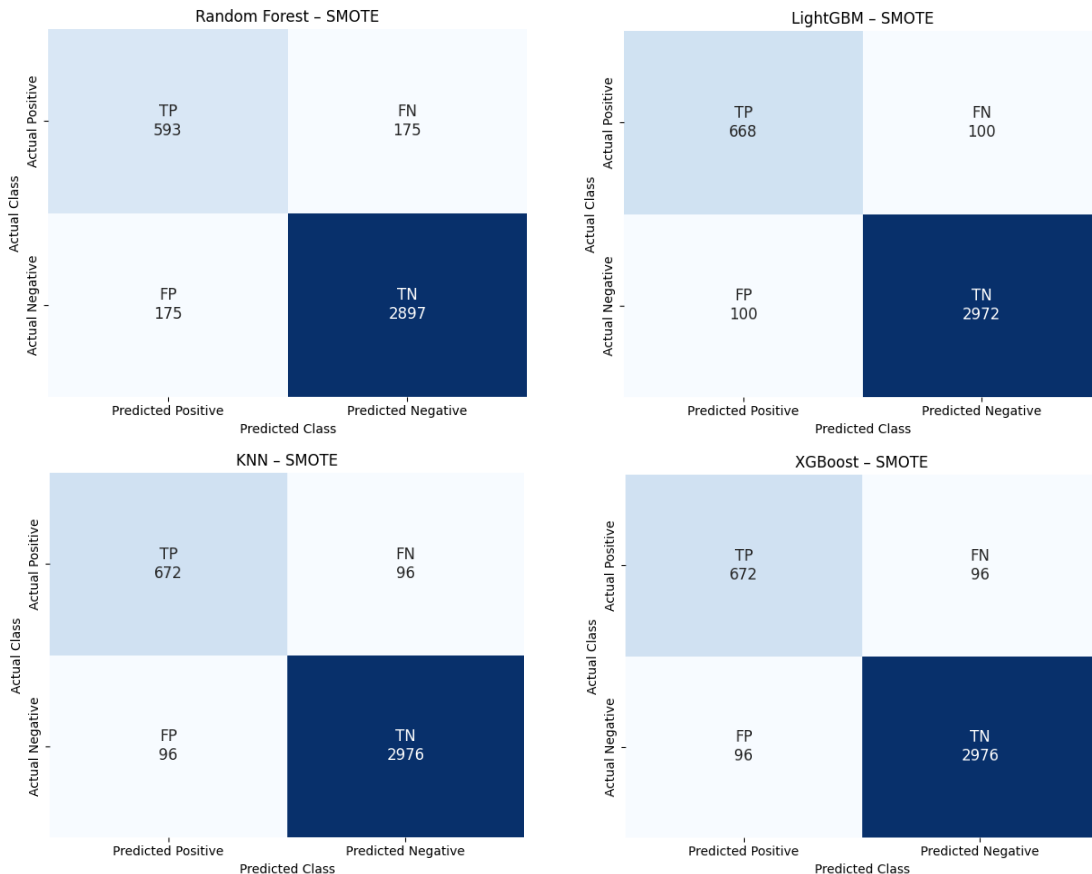
Gambar 2. Confusion Matrix Sebelum Penerapan SMOTE

Visualisasi *Confusion Matrix* pada Gambar 2 menunjukkan adanya bias mayoritas yang sangat jelas pada seluruh model sebelum diterapkannya teknik penyeimbangan data. Berdasarkan data yang disajikan pada Gambar 2, model *Random Forest* mencatat nilai *True Negative* (TN) yang tinggi sebesar 1329 kasus, namun masih gagal mendeteksi 135 kasus kelas minoritas atau *False Negative* (FN). Kondisi serupa terlihat pada model *LightGBM* dalam Gambar 2, di mana model cenderung memprediksi kelas mayoritas secara dominan dengan TN sebanyak 1321 kasus, sementara angka FN mencapai 143 kasus yang menunjukkan lemahnya sensitivitas deteksi pada kelas hujan. Analisis pada Gambar 2 untuk model *XGBoost* juga menunjukkan hasil yang identik dengan *LightGBM*, yaitu menghasilkan 1321 TN namun tetap memiliki tingkat ketidakakuratan pada kelas minoritas dengan FN sebesar 143 kasus. Sementara itu, algoritma KNN pada

Gambar 2 mencatatkan performa yang paling rendah dalam mengenali kelas minoritas dibandingkan model lainnya, dengan angka FN tertinggi mencapai 152 kasus dan nilai *True Positive* (TP) hanya sebesar 214 kasus. Secara keseluruhan, data pada Gambar 2 mengonfirmasi bahwa tanpa adanya perbaikan pada distribusi kelas, seluruh model akan terus mengabaikan pola dari kelas minoritas sehingga mengurangi keandalan sistem dalam mendeteksi fenomena cuaca ekstrem. Oleh karena itu, penerapan teknik SMOTE menjadi langkah krusial untuk merekonstruksi keseimbangan data dan meningkatkan sensitivitas model terhadap seluruh kelas secara setara.

3.7 Analisis Confusion Matrix Setelah Penerapan SMOTE

Sebaliknya, Gambar 3 menunjukkan adanya transformasi kinerja yang mencolok pada seluruh algoritma setelah penerapan metode penyeimbangan data SMOTE. Masalah bias terhadap kelas mayoritas yang sebelumnya terdeteksi pada model *baseline* berhasil diselesaikan secara signifikan



Gambar 3. Confusion Matrix Setelah Penerapan SMOTE

Berdasarkan visualisasi pada Gambar 3, model *LightGBM* mencatatkan lonjakan nilai *True Positive* (TP) menjadi 668 kasus dengan rasio kesalahan prediksi (*False Negative* dan *False Positive*) yang seimbang pada angka 100 kasus. Peningkatan serupa juga terlihat pada algoritma KNN dan *XGBoost* dalam Gambar 3, di mana keduanya berhasil mencapai nilai TP tertinggi sebesar 672 kasus dengan tingkat kesalahan prediksi yang sangat minim, yaitu hanya 96 kasus baik untuk FN maupun FP. Sementara itu, model *Random Forest* pada Gambar 3 menunjukkan peningkatan TP menjadi 593 kasus, meskipun masih memiliki rasio kesalahan sebesar 175 kasus. Keseimbangan dalam distribusi kesalahan yang tertera pada Gambar 3 membuktikan bahwa garis keputusan seluruh model telah bergeser menjadi lebih efektif dan tidak lagi condong kepada kelas mayoritas. Dengan peningkatan tajam pada *Recall*, model-model yang disajikan dalam Gambar 3 kini memiliki tingkat sensitivitas yang jauh lebih tinggi, terutama pada model *LightGBM* dan *XGBoost* yang mencapai sekitar 87% dalam mendeteksi curah hujan. Hasil ini menjadikan sistem prediksi tersebut jauh lebih dapat diandalkan untuk kebutuhan mitigasi risiko dan peringatan dini dibandingkan dengan kondisi model pada tahap *baseline*.

3.8 Perbandingan Hasil

Bagian ini menyajikan analisis perbandingan yang mendalam tentang kinerja model klasifikasi dalam dua konteks pengujian yang berbeda, yaitu sebelum penerapan metode penyeimbangan data (*baseline*) dan setelah penerapan teknik SMOTE. Rangkuman hasil evaluasi kuantitatif untuk semua algoritma yang diuji disajikan secara terperinci dalam Tabel 8. Tabel tersebut menunjukkan perbandingan kinerja berdasarkan lima metrik evaluasi utama, yaitu Akurasi, Presisi, Recall, F1-Score, dan AUC ROC, untuk memberikan gambaran menyeluruh tentang efektivitas SMOTE dalam meningkatkan kualitas prediksi model, terutama pada kelas minoritas.

Tabel 8. Perbandingan Evaluasi Model

Model	Akurasi	Presisi	Recall	F1-Score	AUC ROC
Random Forest	0,631	0,245	0,273	0,258	0,787
RF + SMOTE	0,772	0,766	0,772	0,762	0,946
XGBoost	0,609	0,295	0,285	0,283	0,749
XGB + SMOTE	0,875	0,875	0,875	0,875	0,976
LightGBM	0,609	0,276	0,279	0,274	0,762
LGBM + SMOTE	0,869	0,871	0,870	0,870	0,976
KNN	0,584	0,265	0,267	0,264	0,581
KNN + SMOTE	0,875	0,871	0,875	0,869	0,957

Berdasarkan data yang disajikan pada Tabel 8, penilaian kinerja difokuskan pada metrik *Recall* dan F1-Score sebagai indikator utama efektivitas model dalam menangani ketidakseimbangan data. Dalam skenario dasar tanpa SMOTE pada Tabel 8, terlihat bahwa seluruh model mengalami kesulitan besar dalam mengenali kelas minoritas dengan nilai *Recall* yang sangat rendah, berkisar antara 0,267 hingga 0,285. Angka tersebut menunjukkan bahwa tanpa pengolahan data yang seimbang, model cenderung memihak kelas mayoritas dan gagal mendeteksi hampir 70% kejadian hujan yang sebenarnya terjadi (*False Negative*). Namun, penerapan teknik SMOTE memberikan dampak yang sangat signifikan sebagaimana terekam dalam Tabel 8. Terjadi lonjakan besar pada nilai *Recall* untuk semua model, dengan peningkatan rata-rata hingga tiga kali lipat mencapai rentang 0,772 hingga 0,875. Sebagai contoh spesifik pada Tabel 8, model XGBoost mencatatkan kenaikan *Recall* yang sangat tajam dari 0,285 menjadi 0,875. Peningkatan ini selaras dengan kenaikan nilai F1-Score dari 0,283 menjadi 0,875, yang menandakan bahwa model kini memiliki keseimbangan optimal antara kemampuan presisi dan sensitivitas. Hasil evaluasi pada Tabel 8 juga menunjukkan bahwa model LightGBM dan KNN mengalami peningkatan metrik secara masif, dengan nilai AUC ROC yang mencapai 0,976 dan 0,957. Temuan ini menjawab permasalahan utama penelitian bahwa penggunaan akurasi semata dalam analisis cuaca sering kali menyesatkan. Kenaikan drastis pada seluruh metrik di Tabel 8 membuktikan bahwa integrasi SMOTE berhasil memperkuat kemampuan model dalam mengenali kondisi cuaca ekstrem secara efektif dibandingkan metode konvensional, sehingga sistem peringatan dini yang dihasilkan menjadi jauh lebih akurat.

4. KESIMPULAN

Penelitian ini telah berhasil menunjukkan bahwa penggunaan teknik *Synthetic Minority Over-sampling Technique* (SMOTE) memiliki peran penting dalam menyelesaikan masalah ketidakseimbangan data pada klasifikasi curah hujan. Temuan utama dari kajian ini menunjukkan adanya peningkatan signifikan pada nilai *Recall* dan *F1-Score* setelah penerapan SMOTE, yang sebelumnya sangat rendah pada model dasar karena adanya kecenderungan terhadap kelas mayoritas. Perubahan ini menunjukkan bahwa SMOTE berhasil meningkatkan sensitivitas model, sehingga kemampuan dalam mendeteksi kelas minoritas jauh lebih baik. Dari evaluasi perbandingan, algoritma *LightGBM* yang digabungkan dengan SMOTE terbukti menjadi metode yang paling efektif dengan performa yang paling konsisten, mencapai akurasi 87%, F1-Score 0,870, dan AUC ROC sebesar 0,976. Hasil ini tidak hanya mencerminkan keseimbangan optimal antara Presisi dan *Recall*, tetapi juga mengungguli kinerja penelitian sebelumnya yang menggunakan metode KNN (akurasi 86,09%). Implikasi dari temuan ini menekankan bahwa penanganan data yang tidak seimbang adalah elemen kunci untuk membangun sistem peringatan dini yang efisien serta mengurangi kesalahan prediksi (*false negative*). Untuk langkah pengembangan studi berikutnya, disarankan untuk mengkaji teknik *resampling* adaptif yang lain seperti ADASYN atau kombinasi hibrida *SMOTE-Tomek Links* untuk mengurangi pengaruh *noise* di perbatasan kelas. Selain itu, penambahan data deret waktu (*time-series*) yang lebih panjang serta penerapan arsitektur *deep learning* seperti *Long Short-Term Memory* (LSTM) sangat dianjurkan untuk menangkap kompleksitas dinamika pola cuaca dengan lebih akurat.

REFERENCES

- [1] A. Syahreza, N. K. Ningrum, and M. A. Syahrazy, "Perbandingan Kinerja Model Prediksi Cuaca: Random Forest, Support Vector Regression, dan XGBoost," *Edumatic: Jurnal Pendidikan Informatika*, vol. 8, no. 2, pp. 526–534, Dec. 2024, doi: 10.29408/edumatic.v8i2.27640.
- [2] Z. A. Dwiyantri and C. Prianto, "Prediksi Cuaca Kota Jakarta Menggunakan Metode Random Forest," *Jurnal Tekno Insentif*, vol. 17, no. 2, pp. 127–137, Oct. 2023, doi: 10.36787/jti.v17i2.1136.
- [3] A. F. D. Putra, M. N. Azmi, H. Wijayanto, S. Utama, and I. G. P. W. Wedashwara Wirawan, "Optimizing Rain Prediction Model Using Random Forest and Grid Search Cross-Validation for Agriculture Sector," *MATRIK: Jurnal Manajemen, Teknik Informatika dan Rekayasa Komputer*, vol. 23, no. 3, pp. 519–530, Jul. 2024, doi: 10.30812/matrik.v23i3.3891.
- [4] A. Elsa Damayanti, E. Andriani, and M. Giyanfranco Zola, "Implementasi Machine Learning Untuk Prediksi Curah Hujan di Daerah Rawan Banjir," 2025.

- [5] Z. Ben Bouallègue et al., “The Rise of Data-Driven Weather Forecasting A First Statistical Assessment of Machine Learning–Based Weather Forecasts in an Operational-Like Context,” *Bull Am Meteorol Soc*, vol. 105, no. 6, pp. E864–E883, Jun. 2024, doi: 10.1175/BAMS-D-23-0162.1.
- [6] H. Zhang, Y. Liu, C. Zhang, and N. Li, “Machine Learning Methods for Weather Forecasting: A Survey,” Jan. 01, 2025, Multidisciplinary Digital Publishing Institute (MDPI). doi: 10.3390/atmos16010082.
- [7] I. Hapsari and S. Pandya Wisesa, “Evaluasi Model Prediksi Curah Hujan Berbasis Machine Learning di Kota Bandung,” *Jurnal Nasional Teknologi dan Sistem Informasi*, vol. 11, no. 2, pp. 136–143, Sep. 2025, doi: 10.25077/teknosi.v11i2.2025.136-143.
- [8] A. Elsa Damayanti, E. Andriani, and M. Giyanfranco Zola, “Implementasi Machine Learning Untuk Prediksi Curah Hujan di Daerah Rawan Banjir,” 2025.
- [9] Z. A. Dwiyanti and C. Prianto, “Prediksi Cuaca Kota Jakarta Menggunakan Metode Random Forest,” *Jurnal Tekno Insentif*, vol. 17, no. 2, pp. 127–137, Oct. 2023, doi: 10.36787/jti.v17i2.1136.
- [10] N. Cahyani, W. A. Putri, and R. Irsyada, “Improving Multiclass Rainfall Prediction with Multilayer Perceptron and SMOTE: Addressing Class Imbalance Challenges,” *Brilliance: Research of Artificial Intelligence*, vol. 4, no. 2, pp. 901–908, Jan. 2025, doi: 10.47709/brilliance.v4i2.5203.
- [11] C. Nur Azzahra et al., “Sistemasi: Jurnal Sistem Informasi Klasifikasi Cuaca Jawa Barat menggunakan Ensemble Learning pada Data Meteorologi Weather Classification in West Java using Ensemble Learning on Meteorological Data,” 2025. [Online]. Available: <http://sistemasi.ftik.unisi.ac.id>
- [12] C. G. L. Pringandana and K. Kusnawi, “A Comparative Analysis of Hyperparameter-Tuned XGBoost and LightGBM for Multiclass Rainfall Classification in Jakarta,” *Jurnal Teknik Informatika (Jutif)*, vol. 6, no. 4, pp. 2467–2483, Aug. 2025, doi: 10.52436/1.jutif.2025.6.4.4965.
- [13] “Data Online - Direktorat Data dan Komputasi BMKG.” Accessed: Dec. 19, 2025. [Online]. Available: <https://dataonline.bmkg.go.id/dataonline-home>
- [14] untuk Pemodelan Prediktif Jumlah Peserta Ajar Mata Kuliah Anggi Perwitasari, R. Septiriana, J. Hadari Nawawi, and K. Barat, “JEPIN (Jurnal Edukasi dan Penelitian Informatika),” 2023.
- [15] R. Arisandi, D. Ruhiat, and D. E. Marlina, “Implementasi Ridge Regression untuk Mengatasi Gejala Multikolinearitas pada Pemodelan Curah Hujan Berbasis Data Time Series Klimatologi,” *Jurnal Riset Matematika dan Sains Terapan*, vol. 1, no. 1, pp. 1–11, 2021.
- [16] R. C. Chen, C. Dewi, S. W. Huang, and R. E. Caraka, “Selecting critical features for data classification based on machine learning methods,” *J Big Data*, vol. 7, no. 1, Dec. 2020, doi: 10.1186/s40537-020-00327-4.
- [17] M. Yusuf et al., “PENERAPAN ALGORITMA K-NEAREST NEIGHBOR (KNN) DALAM MEMREDIKSI DAN MENGHITUNG TINGKAT AKURASI DATA CUACA DI INDONESIA,” vol. 2, no. 2, 2021.
- [18] J. Pebralia, “JIFP (Jurnal Ilmu Fisika dan Pembelajarannya) Analisis Curah Hujan Menggunakan Machine Learning Metode Regresi Linier Berganda Berbasis Python dan Jupyter Notebook Rainfall Analysis using Machine Learning-Multiple Linear Regression Method Based on Python and Jupyter Notebook,” vol. 6, no. 2, pp. 23–30, 2022, [Online]. Available: <http://jurnal.radenfatah.ac.id/index.php/jifp/>
- [19] Y. Hasnataeni, A. Saefuddin, and A. M. Soleh, “Comparison of Ensemble Forest-Based Methods Performance for Imbalanced Data Classification,” *Scientific Journal of Informatics*, vol. 12, no. 2, pp. 183–198, Jun. 2025, doi: 10.15294/sji.v12i2.24269.
- [20] C. Nur Azzahra et al., “Sistemasi: Jurnal Sistem Informasi Klasifikasi Cuaca Jawa Barat menggunakan Ensemble Learning pada Data Meteorologi Weather Classification in West Java using Ensemble Learning on Meteorological Data,” 2025. [Online]. Available: <http://sistemasi.ftik.unisi.ac.id>
- [21] Y. M. Indah, R. Aristawidya, A. Fitrianto, E. Erfiani, and L. M. R. D. Jumansyah, “Comparison of Random Forest, XGBoost, and LightGBM Methods for the Human Development Index Classification,” *Jambura Journal of Mathematics*, vol. 7, no. 1, pp. 14–18, Jan. 2025, doi: 10.37905/jjom.v7i1.28290.
- [22] V. R. Danestiara, “Algoritma k-Nearest Neighbor Classifier untuk Prediksi Curah Hujan di Kabupaten Bandung,” 2023.
- [23] S. Tırnık, “Machine learning-based forecasting of air quality index under long-term environmental patterns: A comparative approach with XGBoost, LightGBM, and SVM,” *PLoS One*, vol. 20, no. 10 October, Oct. 2025, doi: 10.1371/journal.pone.0334252.