

Perbandingan Algoritma Machine Learning untuk Klasifikasi Kopi Menggunakan Data Sensor Electronic Nose dan Tongue

Dwi Issadari Hastuti^{1,*}, Mula Agung Barata², Ifnu Wisma Dwi Prastya²

¹ Fakultas Sains dan Teknologi, Sistem Komputer, Universitas Nahdlatul Ulama Sunan Giri, Bojonegoro, Indonesia

² Fakultas Sains dan Teknologi, Teknik Informatika, Universitas Nahdlatul Ulama Sunan Giri, Bojonegoro, Indonesia

Email: ^{1,*}dwiastuti@unugiri.ac.id, ²mula.ab26@gmail.com, ³ifnudwi867@gmail.com

Email Penulis Korespondensi: dwiastuti@unugiri.ac.id

Submitted 15-11-2025; Accepted 10-01-2026; Published 28-02-2026

Abstrak

Kopi merupakan komoditas unggulan Indonesia yang memiliki keragaman aroma dan cita rasa yang dipengaruhi oleh varietas serta daerah asal. Namun, proses identifikasi dan klasifikasi jenis kopi masih banyak dilakukan secara konvensional melalui uji sensorik, yang bersifat subjektif, membutuhkan waktu lama, dan bergantung pada keahlian panelis. Kondisi ini mendorong perlunya pendekatan otomatis yang lebih objektif dan konsisten berbasis teknologi sensor dan machine learning. Penelitian ini bertujuan untuk membandingkan kinerja beberapa algoritma machine learning, yaitu Logistic Regression, Support Vector Classifier (SVC), dan Random Forest, dalam mengklasifikasikan jenis kopi Indonesia menggunakan data multisensor Electronic Nose dan Electronic Tongue. Data yang digunakan berasal dari sensor gas, suhu, dan pH dengan jumlah 1.503 sampel yang merepresentasikan sepuluh kelas kopi. Tahapan pra-proses meliputi pembersihan data menggunakan metode Interquartile Range (IQR) untuk menghilangkan outlier serta reduksi noise menggunakan metode Moving Average. Hasil penelitian menunjukkan bahwa penerapan pembersihan dan peredaman noise data secara signifikan meningkatkan kinerja seluruh model klasifikasi. Di antara algoritma yang diuji, Random Forest menunjukkan performa paling stabil dan unggul dalam mengklasifikasikan jenis kopi. Temuan ini menegaskan bahwa kombinasi pra-proses data yang tepat dan pemilihan algoritma yang sesuai berperan penting dalam meningkatkan akurasi sistem klasifikasi kopi berbasis machine learning.

Kata Kunci: Electronic Nose; Electronic Tongue; IQR; Moving Average; Klasifikasi Kopi

Abstract

Coffee is a leading Indonesian commodity with a diversity of aromas and flavors influenced by variety and region of origin. However, the process of identifying and classifying coffee types is still often carried out conventionally through sensory testing, which is subjective, time-consuming, and dependent on panelist expertise. This situation encourages the need for a more objective and consistent automated approach based on sensor technology and machine learning. This study aims to compare the performance of several machine learning algorithms, namely Logistic Regression, Support Vector Classifier (SVC), and Random Forest, in classifying Indonesian coffee types using multisensor Electronic Nose and Electronic Tongue data. The data used comes from gas, temperature, and pH sensors with a total of 1,503 samples representing ten coffee classes. The preprocessing stage includes data cleaning using the Interquartile Range (IQR) method to remove outliers and noise reduction using the Moving Average method. The results show that the application of data cleaning and noise reduction significantly improves the performance of all classification models. Among the algorithms tested, Random Forest showed the most stable and superior performance in classifying coffee types. These findings confirm that the combination of appropriate data preprocessing and appropriate algorithm selection plays a crucial role in improving the accuracy of machine learning-based coffee classification systems.

Keywords: Electronic Nose; Electronic Tongue; IQR; Moving Average; Coffee Classification

1. PENDAHULUAN

Kopi merupakan salah satu komoditas pertanian unggulan dan minuman paling populer di dunia yang memiliki nilai ekonomi tinggi, baik dari sisi konsumsi maupun perdagangan internasional. Indonesia sebagai salah satu produsen kopi terbesar dunia [1] memiliki keragaman varietas dan daerah asal kopi yang luas, terutama jenis Arabika dan Robusta yang tersebar di berbagai wilayah nusantara. Setiap daerah memiliki karakteristik cita rasa dan aroma yang berbeda, dipengaruhi oleh variasi geografis, iklim, ketinggian tempat, jenis tanah, hingga teknik pascapanen yang diterapkan petani [2]. Keragaman ini menjadikan kopi Indonesia sangat kaya secara organoleptik, namun sekaligus menimbulkan tantangan dalam proses penilaian kualitas secara konsisten. Oleh karena itu, klasifikasi kopi berdasarkan varietas dan daerah asal menjadi langkah penting dalam menjaga mutu, meningkatkan daya saing produk di pasar global, serta mencegah pemalsuan atau ketidaksesuaian label produk. Namun demikian, metode klasifikasi yang banyak digunakan hingga kini masih bergantung pada uji sensori manusia yang bersifat subjektif, memerlukan pelatihan khusus, dipengaruhi kondisi fisiologis penilai, serta sulit distandardisasi antar penilai [3]. Kondisi ini menegaskan perlunya pendekatan yang objektif, cepat, reproducible, dan akurat dalam proses klasifikasi kopi.

Seiring perkembangan teknologi kecerdasan buatan, *Electronic Nose* (E-Nose) dan *Electronic Tongue* (E-tongue) menjadi solusi alternatif yang menjanjikan dalam pengujian kualitas pangan [4], [5] dan minuman [6], [7]. E-Nose bekerja menyerupai fungsi hidung manusia, menggunakan kumpulan sensor gas untuk mendeteksi senyawa volatil yang dilepaskan oleh sampel, sementara E-tongue berfungsi menyerupai lidah manusia dengan memanfaatkan sensor pH atau ion untuk menangkap profil rasa. Kedua teknologi ini mampu menghasilkan pola data khas yang dapat dijadikan “sidik jari” aroma dan rasa suatu produk. Kombinasi E-Nose dan E-tongue memberikan representasi yang lebih lengkap terhadap karakteristik kimia kopi, sehingga semakin meningkatkan peluang memperoleh hasil klasifikasi yang akurat. Teknologi

ini telah diaplikasikan pada berbagai bidang seperti industri makanan, minuman, farmasi [8], dan kesehatan [9], serta terbukti efektif dalam meningkatkan objektivitas proses penilaian mutu dan deteksi kontaminan.

Berbagai penelitian sebelumnya menunjukkan potensi besar penggunaan E-Nose dan *E-tongue* dalam klasifikasi kopi. Pada penelitian klasifikasi kopi menggunakan kombinasi parameter statistik dan algoritma *Decision Tree* untuk membedakan kopi luwak dan non-luwak, hasil akurasi yang diperoleh mencapai 97% [10]. Studi lain yang menerapkan *K-Nearest Neighbour* (KNN) pada campuran kopi Arabika luwak dan Robusta non-luwak menunjukkan akurasi 97,7% [11], menegaskan potensi machine learning dalam mengolah data sensor kopi. Penelitian serupa yang mengimplementasikan *Support Vector Machine* (SVM) dan Perceptron untuk membedakan Arabika dan Robusta masing-masing memperoleh akurasi 71% dan 57% [12], menunjukkan perlunya pemilihan algoritma dan teknik pra-pemrosesan yang tepat. Lebih lanjut, penggunaan Artificial Neural Network (ANN) *backpropagation* dengan tiga sensor gas juga mampu mengenali kopi Arabika, Robusta, dan Liberika dengan akurasi 73,3% [13], meskipun performanya masih terbatas oleh jumlah sensor dan kompleksitas jaringan yang digunakan.

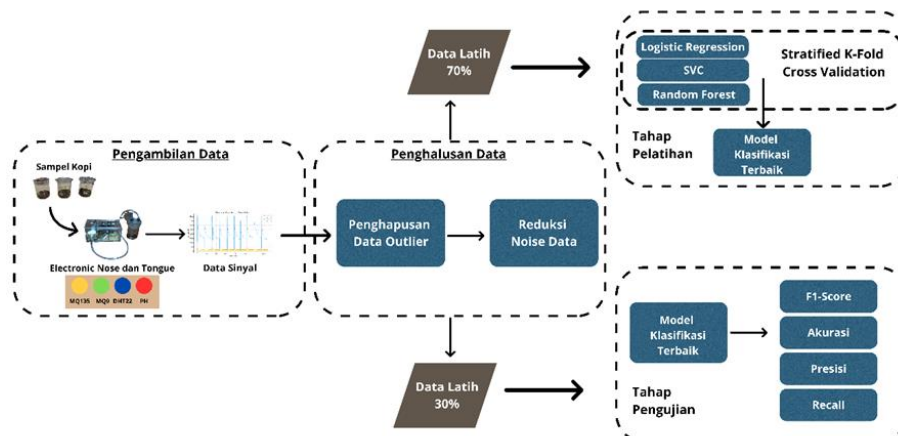
Selain pemilihan algoritma, tahap pra-pemrosesan data menjadi faktor signifikan yang menentukan kualitas model klasifikasi. Sinyal sensor pada E-Nose dan *E-tongue* sering kali mengandung *noise*, fluktuasi nilai yang tidak stabil, serta outlier yang muncul akibat gangguan lingkungan, variasi suhu, kelembapan, atau kondisi teknis alat. Sejumlah penelitian menunjukkan bahwa proses seperti deteksi dan penghapusan *outlier* serta reduksi derau (*noise reduction*) dapat memperbaiki distribusi data, menstabilkan pola sinyal, serta meningkatkan performa model secara keseluruhan [14]. Teknik-teknik ini menjadi semakin penting terutama ketika bekerja dengan data multi-sensor yang bersifat nonlinier dan rentan terhadap variabilitas tinggi.

Berbagai studi juga menunjukkan bahwa algoritma machine learning klasik seperti *Logistic Regression* (LR), *Support Vector Classifier* (SVC), dan *Random Forest* (RF) memiliki performa tinggi dalam berbagai kasus klasifikasi. Pada penelitian yang dilakukan untuk mengklasifikasikan kopi murni dan kopi dengan campuran susu menggunakan *Logistic Regression*, model menghasilkan akurasi 98% yang menunjukkan bahwa algoritma ini dapat bekerja efektif pada pola data yang terstruktur dan linear [15]. Sementara itu, penelitian lain yang melakukan klasifikasi data sensor menggunakan *Support Vector Classifier* memperoleh akurasi sebesar 96,7%, menegaskan kemampuan SVC dalam mengklasifikasikan data sensor volatil yang dihasilkan oleh E-Nose [16]. Pada penelitian lain yang menggunakan algoritma *Random Forest* dalam klasifikasi sinyal hasil E-Nose juga menunjukkan performa akurasi tinggi karena kemampuannya mengatasi data berdimensi besar dan memiliki struktur pohon keputusan yang kuat terhadap *noise*. Dengan demikian, perbandingan ketiga algoritma tersebut menjadi penting untuk mengetahui algoritma yang paling optimal dalam klasifikasi kopi berbasis E-Nose dan *E-tongue*.

Meskipun penelitian terdahulu telah memberikan kontribusi berarti, sebagian besar masih menggunakan pendekatan algoritma tunggal, bergantung pada teknik ekstraksi fitur manual, atau belum memanfaatkan pra-pemrosesan data secara terpadu. Di sisi lain, data sensor kopi mengandung karakteristik yang kompleks sehingga membutuhkan pipeline analisis yang lebih menyeluruh untuk mencapai performa terbaik. Dengan mempertimbangkan hal tersebut, penelitian ini bertujuan untuk membandingkan performa tiga algoritma *machine learning*, yaitu *Logistic Regression* (LR), *Support Vector Classifier* (SVC), dan *Random Forest* (RF) dalam mengklasifikasikan kopi Indonesia berdasarkan data sensor *Electronic Nose* dan *Electronic Tongue*. Selain itu, penelitian ini juga menerapkan tahapan *outlier detection* dan *noise reduction* untuk meningkatkan kestabilan serta akurasi hasil klasifikasi.

2. METODOLOGI PENELITIAN

Penelitian ini bertujuan untuk membangun sistem *Electronic Nose* (E-Nose) dan *Electronic Tongue* (E-tongue) yang digunakan untuk mengklasifikasikan kopi berdasarkan varietas (Arabika dan Robusta) dan daerah asal. Metode yang diusulkan ditunjukkan pada Gambar 1.



Gambar 1. Metodologi Penelitian

Metodologi penelitian meliputi beberapa tahapan pertama pengambilan data, penghalusan data dengan menggunakan penghapusan data *outlier* dan reduksi *noise* data. Setelah itu data dibagi menjadi data uji dan data latih dengan proporsi 80:20. Data latih dimodelkan dengan menggunakan tiga algoritma *machine learning* meliputi *Logistic Regression*, *Support Vector Classifier* dan *Random Forest*. Pemodelan data latih dilakukan dengan *Stratified K-Fold Cross Validation*. Model terbaik pada data latih diuji menggunakan data uji untuk mengukur akurasi, presisi, *recall* dan *f1-score*.

2.1 Bahan dan Sampel Penelitian

Penelitian ini menggunakan 10 kelas kopi yang terdiri dari dua varietas utama, yaitu Arabika (Bengkulu, Jawa Tengah, Lampung, Sumatera, dan Jawa Timur) serta Robusta (Bengkulu, Jawa Tengah, Lampung, Sumatera, dan Jawa Timur). Setiap sampel kopi ditimbang sebanyak 15gram dalam bentuk bubuk. Untuk menjaga konsistensi, seluruh sampel diperoleh dari sumber yang terkontrol dan disimpan pada kondisi lingkungan yang seragam sebelum dilakukan pengujian.

2.2 Pengambilan Data

Sistem deteksi terdiri atas *Electronic Nose* (E-Nose) dan *Electronic Tongue* (E-tongue) yang dikembangkan secara mandiri. E-Nose menggunakan dua sensor gas, yaitu MQ135 dan MQ9, yang memiliki sensitivitas tinggi terhadap senyawa volatil. E-tongue menggunakan sensor pH untuk mendeteksi tingkat keasaman, serta sensor suhu untuk merekam perubahan temperatur. Seluruh sensor dihubungkan dengan mikrokontroler yang berfungsi untuk mengakuisisi data dalam bentuk sinyal analog (A0).

2.3 Penghapusan Data Outlier

Sebelum data sensor digunakan dalam proses klasifikasi, dilakukan tahap deteksi outlier untuk memastikan bahwa data yang digunakan benar-benar merepresentasikan karakteristik aroma dan rasa kopi secara akurat. Outlier merupakan data yang memiliki nilai ekstrem atau menyimpang jauh dari pola umum data lainnya. Keberadaan outlier dapat menyebabkan model *machine learning* menghasilkan hasil klasifikasi yang bias dan kurang akurat karena algoritma akan berusaha menyesuaikan diri terhadap data yang tidak representatif tersebut.

Pada penelitian ini, deteksi outlier dilakukan menggunakan metode *Interquartile Range* (IQR), salah satu pendekatan statistik yang umum digunakan untuk mengidentifikasi nilai-nilai ekstrem dalam distribusi data tanpa mengasumsikan bentuk distribusi tertentu (non-parametrik)[17]. Metode ini dipilih karena mampu bekerja secara efisien pada data sensor yang bersifat fluktuatif dan tidak selalu berdistribusi normal, seperti sinyal hasil pembacaan *Electronic Nose* (E-Nose) dan *Electronic Tongue* (E-tongue).

2.4 Reduksi Noise Data

Setelah proses deteksi dan penghapusan *outlier*, langkah berikutnya adalah reduksi derau (*noise reduction*) untuk memperoleh sinyal yang lebih halus, stabil, dan representatif terhadap karakteristik aroma serta rasa kopi. Reduksi derau menjadi tahap penting dalam pengolahan data sensor *Electronic Nose* (E-Nose) dan *Electronic Tongue* (E-tongue) karena kedua perangkat ini sangat sensitif terhadap gangguan lingkungan, perubahan suhu, kelembapan, dan fluktuasi tekanan udara selama proses akuisisi. Kehadiran *noise* dapat menurunkan kualitas sinyal, mempersulit proses ekstraksi pola, serta mengurangi performa algoritma klasifikasi. Dalam penelitian ini menerapkan metode *moving average* untuk reduksi *noise* data sinyal. Metode *Moving Average* merupakan pendekatan sederhana namun efektif untuk mereduksi fluktuasi acak pada sinyal sensor. Prinsip dasarnya adalah menggantikan setiap titik data dengan rata-rata dari sejumlah nilai tetangga di sekitarnya [18]. Dengan demikian, komponen derau berfrekuensi tinggi dapat diredam, sementara tren utama sinyal tetap dipertahankan. Rumus umum dari *Moving Average* dituliskan pada persamaan (1).

$$\bar{x}_t = \frac{1}{n} \sum_{i=0}^{n-1} x_{t-i} \quad (1)$$

dengan $\bar{x}_t$ adalah nilai hasil perataan pada waktu ke- t , x_{t-i} adalah nilai asli data pada titik ke- $t - i$, dan n adalah ukuran jendela (*window size*).

2.5 Pembagian Data

Setelah data sensor melalui tahap pra-pemrosesan seperti deteksi *outlier* dan reduksi *noise*, langkah berikutnya adalah pembagian data menjadi dua subset utama, yaitu data latih dan data uji. Pembagian data dilakukan agar proses pembangunan dan evaluasi model *machine learning* dapat berjalan secara objektif dan terukur. Dalam penelitian ini, proporsi pembagian data ditetapkan sebesar 70% untuk data latih dan 30% untuk data uji. Proporsi ini dipilih berdasarkan praktik umum dalam penelitian *machine learning*, di mana sebagian besar data digunakan untuk pelatihan model guna mempelajari pola hubungan antara fitur (nilai sensor) dengan label kelas kopi, sementara sebagian lainnya disisihkan untuk menguji kemampuan generalisasi model terhadap data baru yang belum pernah dilihat sebelumnya.

Data latih digunakan untuk membangun dan menyesuaikan parameter model klasifikasi. Pada tahap ini, algoritma *machine learning* seperti *Logistic Regression* (LR), *Support Vector Classifier* (SVC), dan *Random Forest* (RF) menerima input berupa fitur hasil pengolahan sinyal dari sensor E-Nose dan E-tongue. Model kemudian mempelajari hubungan antara pola sinyal dan kategori kopi setiap kelas. Pemodelan data training diukur akurasi untuk menentukan model dengan akurasi terbaik.

Data uji berfungsi untuk mengevaluasi performa model terbaik dari hasil komparasi ketiga algoritma yang telah dibangun menggunakan data latih. Seluruh data uji tidak dilibatkan dalam proses pelatihan, sehingga hasil evaluasi mencerminkan kemampuan model dalam mengklasifikasikan data baru secara independen. Melalui pengujian ini diperoleh metrik performa seperti akurasi, presisi, recall, dan f1-score, yang digunakan untuk membandingkan efektivitas algoritma klasifikasi yang diuji.

2.6 Klasifikasi

Klasifikasi Tahap klasifikasi merupakan inti dari penelitian ini, di mana data sensor yang telah melalui proses preprocessing seperti deteksi *outlier*, reduksi *noise*, dan pembagian data kemudian digunakan untuk melatih model machine learning agar mampu mengidentifikasi jenis kopi berdasarkan pola aroma dan rasa yang terekam oleh sensor *Electronic Nose* (E-Nose) dan *Electronic Tongue* (E-tongue). Pemilihan algoritma klasifikasi dalam penelitian ini didasarkan pada karakteristik data sensor yang bersifat non-linear, multivariat, dan memiliki potensi korelasi antar-fitur yang kompleks. Tiga algoritma yang digunakan adalah *Logistic Regression* (LR), *Support Vector Classifier* (SVC), dan *Random Forest* (RF). Ketiganya dipilih karena telah terbukti memberikan performa yang baik pada berbagai kasus klasifikasi berbasis sinyal sensor dan data eksperimental serupa.

2.6.1 Logistic Regression

Logistic Regression merupakan algoritma klasifikasi linier yang digunakan untuk memprediksi probabilitas suatu data termasuk ke dalam kelas tertentu. Meskipun awalnya dikembangkan untuk kasus biner, LR mampu menangani klasifikasi multikelas melalui pendekatan *one-vs-rest* [19]. Persamaan utama LR didasarkan pada fungsi sigmoid dituliskan pada persamaan (2).

$$P(y = 1|x) = \frac{1}{1 + e^{-(\beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_n x_n)}} \quad (2)$$

di mana $P(y = 1 | x)$ menunjukkan probabilitas bahwa data x termasuk ke kelas positif, dan β_i merupakan koefisien yang dipelajari selama proses pelatihan. LR memiliki keunggulan dalam interpretabilitas dan efisiensi komputasi. Dalam konteks penelitian ini, LR digunakan sebagai model pembandingan dasar untuk mengukur sejauh mana hubungan linier antara fitur sensor dan kategori kopi dapat dimanfaatkan untuk klasifikasi.

2.6.2 Support Vector Classifier

Support Vector Classifier atau *Support Vector Machine* (SVM) merupakan algoritma klasifikasi berbasis teori margin maksimum. SVC bekerja dengan mencari *hyperplane* terbaik yang memisahkan data antar kelas dengan margin terlebar [20]. Untuk menangani data yang tidak terpisahkan secara linier, digunakan fungsi kernel seperti *Radial Basis Function* (RBF) yang memproyeksikan data ke ruang berdimensi lebih tinggi agar dapat dipisahkan secara efektif. Fungsi optimasi SVC didefinisikan seperti pada persamaan (3).

$$\min_{w,b} \frac{1}{2} \|w\|^2 \text{ dengan syarat } y_i(w \cdot x_i + b) \geq 1 \quad (3)$$

Di mana w adalah vektor bobot dan b merupakan bias. SVC dikenal unggul dalam mengatasi data berdimensi tinggi serta robust terhadap *outlier* yang masih tersisa setelah proses pembersihan. Pada penelitian ini, SVC digunakan untuk menguji kemampuan model dalam membedakan pola aroma dan rasa kopi yang kompleks dan tumpang tindih antar-kelas, terutama pada data hasil fusi multi-sensor (E-Nose dan E-tongue).

2.6.3 Random Forest

Random Forest merupakan algoritma berbasis ensemble learning yang menggabungkan sejumlah decision tree untuk meningkatkan akurasi dan stabilitas prediksi. Setiap tree dilatih pada subset acak dari data dan fitur, sehingga menghasilkan variasi dalam proses pelatihan. Prediksi akhir diperoleh melalui mekanisme voting atau rata-rata dari seluruh tree [21]. Konsep dasar RF dapat dituliskan pada persamaan (4).

$$y^{\wedge} = \text{mode}(h_1(x), h_2(x), \dots, h_k(x)) \quad (4)$$

Dengan $h_i(x)$ adalah hasil prediksi dari setiap decision tree. Kelebihan utama *Random Forest* terletak pada kemampuannya mengatasi data non-linear, menangani fitur yang saling berkorelasi, serta mengurangi risiko *overfitting*. Dalam penelitian ini, RF digunakan untuk mengevaluasi efektivitas pendekatan ensemble terhadap performa klasifikasi kopi, serta dibandingkan dengan hasil model LR dan SVC.

2.7 Evaluasi

Tahap evaluasi klasifikasi bertujuan untuk menilai sejauh mana model machine learning yang telah dilatih mampu mengklasifikasikan jenis kopi dengan benar berdasarkan data sensor *Electronic Nose* (E-Nose) dan *Electronic Tongue* (E-tongue). Evaluasi ini dilakukan untuk membandingkan performa tiga algoritma yang digunakan, yaitu *Logistic Regression* (LR), *Support Vector Classifier* (SVC), dan *Random Forest* (RF), dalam mengenali pola aroma dan rasa kopi. Penilaian kinerja model dilakukan menggunakan data uji sebesar 30% dari total dataset, yang tidak dilibatkan dalam proses pelatihan. Untuk mengukur performa klasifikasi, digunakan beberapa metrik evaluasi yang umum dalam analisis

machine learning, yaitu akurasi (*accuracy*), presisi (*precision*), recall, dan *f1-score*. Keempat metrik ini dihitung berdasarkan nilai *True Positive* (TP), *True Negative* (TN), *False Positive* (FP), dan *False Negative* (FN) yang diperoleh dari *confusion matrix*. Dengan demikian, hasil evaluasi menggambarkan kemampuan generalisasi model terhadap data baru yang belum pernah dilihat sebelumnya.

3. HASIL DAN PEMBAHASAN

Bab ini menyajikan hasil dan pembahasan dari rangkaian tahapan penelitian yang telah dilakukan, mulai dari pra-pemrosesan data, pembagian data, hingga penerapan algoritma *machine learning* untuk klasifikasi jenis kopi Indonesia berdasarkan data sensor *Electronic Nose* (E-Nose) dan *Electronic Tongue* (E-tongue). Seluruh proses dilakukan untuk mengevaluasi performa tiga algoritma utama *Logistic Regression* (LR), *Support Vector Classifier* (SVC), dan *Random Forest* (RF) dalam mengenali pola aroma dan rasa kopi dari berbagai daerah penghasil di Indonesia.

3.1 Akuisisi Data Sensor

Berisi hasil pembahasan dan bisa perbandingan dari hasil penelitian sebelumnya. Tahapan awal penelitian ini dimulai dengan proses akuisisi data sensor yang dilakukan menggunakan perangkat *Electronic Nose* (E-Nose) dan *Electronic Tongue* (E-tongue). Proses pengambilan data dilakukan terhadap 10 kelas kopi yang mewakili beberapa daerah penghasil utama kopi di Indonesia. Setiap kelas kopi diuji menggunakan empat jenis sensor, yaitu MQ135, MQ9, PH dan DHT22, dengan waktu pengambilan data selama 10 menit per sampel, di mana pompa aroma diaktifkan selama 1 menit untuk menyedot udara dari wadah sampel kopi ke dalam ruang sensor. Hasil pengukuran berupa tegangan analog dari setiap sensor yang dikonversi menjadi satuan digital menggunakan *Analog-to-Digital Converter* (ADC) dan disimpan dalam format numerik. Tabel 1 berikut menunjukkan distribusi jumlah data yang diperoleh dari setiap kelas kopi hasil akuisisi sistem sensor.

Tabel 1. Data Sinyal

Data Ke	MQ9	MQ135	DHT22	PH	Label
1	149	1073	35,3	14	Robusta Jawa Timur
2	319	4095	35,3	14	Robusta Jawa Timur
...	...	...	...	...	...
4	178	2200	33	14	Arabica Jawa Tengah
5	180	2179	33	14	Arabica Jawa Tengah
...	...	...	...	...	...
1502	162	1330	35	14	Jawa Timur Arabica
1503	161	1329	35	14	Jawa Timur Arabica

Secara keseluruhan, diperoleh sebanyak 1.503 data sensor yang terdiri dari 10 kategori kopi berbeda. Setiap data mencerminkan satu siklus pengukuran yang terdiri dari kombinasi nilai keluaran empat sensor pada satu waktu pengambilan tertentu. Hasil akuisisi ini menunjukkan bahwa sistem sensor yang dikembangkan mampu bekerja secara stabil dan konsisten dalam menangkap pola aroma dan rasa kopi dari berbagai daerah di Indonesia dan dataset ini menjadi dasar kuat dalam tahap analisis selanjutnya, meliputi deteksi *outlier*, *noise reduction*, pembagian data, dan evaluasi performa algoritma klasifikasi.

3.2 Penghapusan Data Outlier

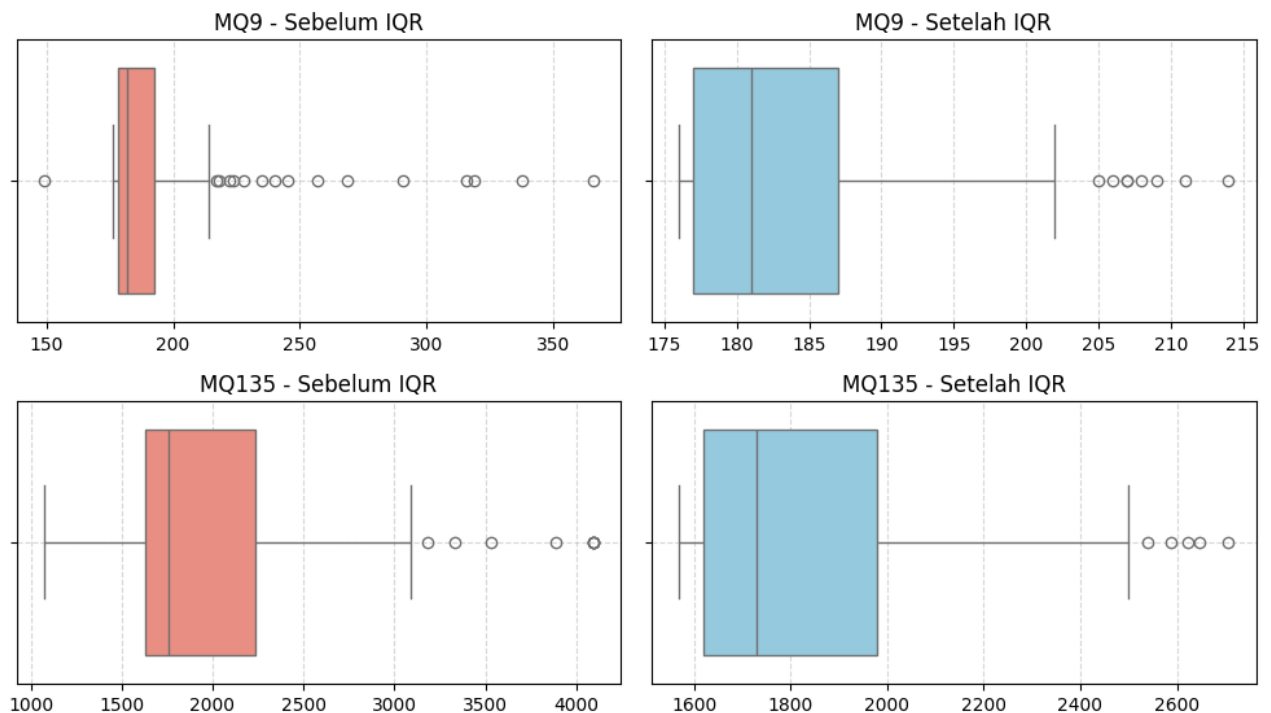
Berisi Tahapan berikutnya setelah akuisisi data sensor adalah penghapusan data outlier guna memastikan bahwa dataset yang digunakan dalam proses klasifikasi benar-benar merepresentasikan pola aroma dan rasa kopi secara konsisten. Proses ini dilakukan menggunakan metode *Interquartile Range* (IQR) yang diterapkan secara terpisah pada setiap label kopi. Pendekatan per-label ini penting karena setiap jenis kopi memiliki karakteristik sensorik yang berbeda, baik dari nilai ambang gas (MQ9 dan MQ135) maupun parameter fisik seperti suhu (Temp C) dan pH. Hasil pembersihan menunjukkan bahwa dari total 1.503 data awal, tersisa 1.321 data setelah outlier dihapus, sehingga terdapat 182 data yang dieliminasi. Persentase penghapusan bervariasi antar kelas, dengan rentang antara 9 % hingga 15 %, tergantung pada tingkat dispersi nilai sensor pada tiap jenis kopi. Ringkasan hasil penghapusan disajikan pada Tabel 2.

Tabel 2. Penghapusan Outlier

No	Label Kopi	Sebelum	Dihapus	Sisa	% Dihapus
1	Arabika Bengkulu	133	18	115	13.53%
2	Arabika Jawa Tengah	131	13	118	9.92%
3	Arabika Lampung	159	16	143	10.06%
4	Arabika Sumatera Selatan	155	20	135	12.90%
5	Arabika Jawa Timur	147	14	133	9.79%
6	Robusta Jawa Tengah	157	18	139	11.46%

7	Robusta Jawa Timur	151	16	135	10.60%
8	Robusta Lampung	166	26	140	15.66%
9	Robusta Bengkulu	162	24	138	14.81%
10	Robusta Sumatera Selatan	146	17	129	11.64%

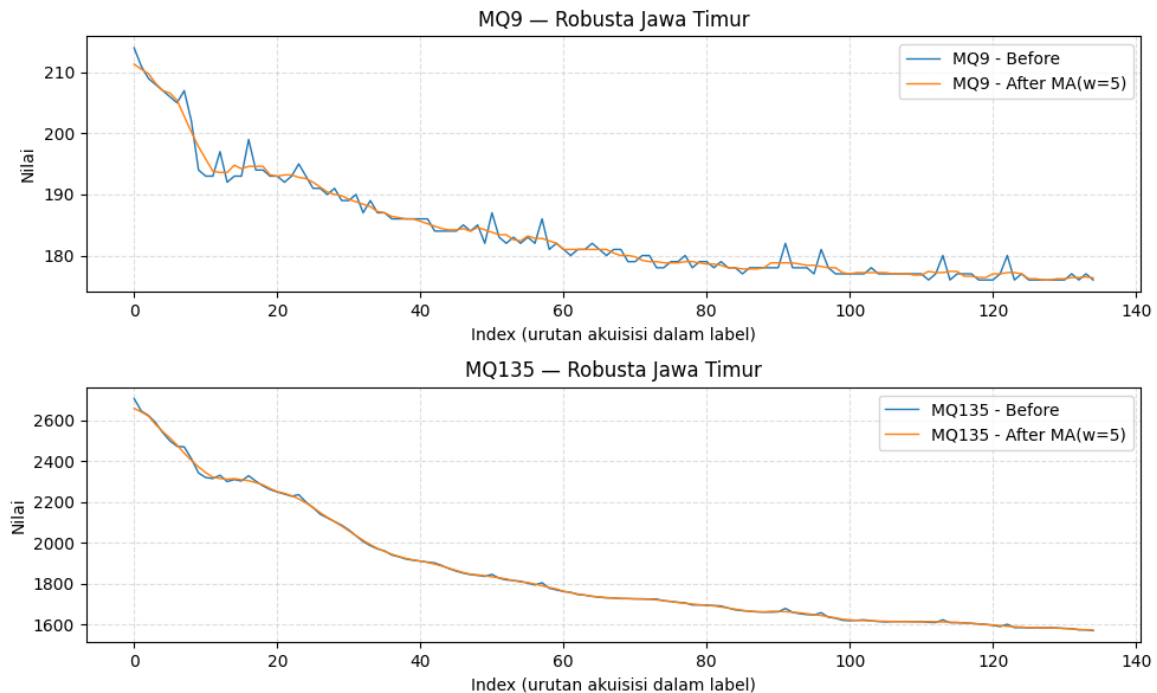
Hasil tersebut memperlihatkan bahwa kelas dengan variasi sensor yang lebih tinggi, seperti Robusta Lampung dan Sumatera Bengkulu Robusta, mengalami tingkat penghapusan terbesar (sekitar 15 %), sementara kelas dengan nilai sensor yang lebih stabil, seperti Arabika Jawa Tengah dan Jawa Timur Arabika, hanya mengalami penghapusan di bawah 10 %. Proses penghapusan *outlier* ini berdampak positif terhadap kualitas data, karena mengurangi noise ekstrem yang dapat memengaruhi performa model klasifikasi. Dataset hasil pembersihan digunakan sebagai dasar untuk tahap reduksi noise dan pelatihan model machine learning pada langkah selanjutnya. Visualisasi penghapusan data outlier seperti pada Gambar 2.



Gambar 2. Penghapusan Data Outlier

3.3 Reduksi Noise Data

Setelah proses penghapusan outlier menggunakan metode *Interquartile Range* (IQR), langkah berikutnya adalah reduksi noise atau smoothing terhadap data sensor untuk meningkatkan kestabilan sinyal dan mengurangi fluktuasi acak yang dapat mengganggu proses klasifikasi. Reduksi *noise* dilakukan menggunakan metode *Moving Average* (MA) yang diterapkan pada setiap fitur sensor dalam masing-masing kelas kopi secara terpisah. Metode *Moving Average* bekerja dengan menggantikan setiap nilai data dengan rata-rata dari sejumlah titik data di sekitarnya, berdasarkan ukuran jendela tertentu. Pada penelitian ini digunakan window size sebesar 5, yang berarti setiap titik data direpresentasikan sebagai rata-rata dari dua nilai sebelumnya, nilai saat ini, dan dua nilai sesudahnya. Dengan pendekatan ini, komponen noise berfrekuensi tinggi dapat dihilangkan, sementara tren utama sinyal tetap terjaga. Proses *smoothing* diterapkan pada empat variabel utama hasil pengukuran sensor, yaitu MQ9, MQ135, *Temperature* (°C), dan pH, untuk setiap label kopi. Berdasarkan hasil perbandingan grafik antara data sebelum dan sesudah penerapan *Moving Average*, terlihat bahwa pola sinyal dari masing-masing sensor menjadi lebih halus dan stabil. Nilai-nilai ekstrem yang sebelumnya masih muncul setelah tahap penghapusan *outlier* berhasil direduksi, terutama pada sensor MQ135 dan MQ9 yang memiliki variasi pembacaan lebih tinggi akibat sensitivitas terhadap gas volatil. Sementara itu, parameter suhu (Temp C) dan pH menunjukkan fluktuasi yang lebih kecil sehingga efek smoothing relatif tidak mengubah bentuk sinyal secara signifikan. Hasil reduksi *noise* data dapat dilihat pada Gambar 3.



Gambar 3. Reduksi Noise Data

3.4 Klasifikasi

3.4.1 Klasifikasi Setelah Penghapusan *Outlier*

Tahap klasifikasi dilakukan untuk mengetahui pengaruh penghapusan *outlier* terhadap peningkatan performa model *machine learning* dalam mengidentifikasi jenis kopi berdasarkan data sensor *Electronic Nose* dan *Electronic Tongue*. Sebelum dilakukan penghapusan *outlier*, data yang digunakan masih mengandung nilai ekstrem yang berpotensi mengganggu proses pembelajaran model. Setelah pembersihan data menggunakan metode *Interquartile Range* (IQR), dilakukan kembali proses klasifikasi dengan dataset hasil filtrasi. Hasil perbandingan akurasi antara model sebelum dan sesudah penghapusan *outlier* disajikan pada Tabel 3.

Tabel 3. Peforma Penghapusan *Outlier*

Metode	Akurasi Sebelum IQR	Akurasi Sesudah IQR
Random Forest	90.68%	93.63%
Support Vector Classifier (SVC)	69.01%	80.99%
Logistic Regression	56.65%	73.58%

Berdasarkan tabel tersebut, terlihat bahwa penghapusan *outlier* menggunakan metode IQR memberikan dampak positif terhadap peningkatan kinerja seluruh algoritma. Model Random Forest menunjukkan performa terbaik dengan akurasi meningkat dari 90.68% menjadi 93.63%, menandakan bahwa model ini paling robust dalam menangani variasi data sensor. Model SVC juga mengalami peningkatan signifikan dari 69.01% menjadi 80.99%, yang menunjukkan bahwa data bersih tanpa nilai ekstrem membantu margin pemisahan kelas menjadi lebih optimal. Sementara itu, *Logistic Regression* yang semula hanya mencapai 56.65%, meningkat menjadi 73.58% setelah penghapusan *outlier*. Hal ini menunjukkan bahwa model linear seperti LR sangat sensitif terhadap *outlier*, sehingga proses pembersihan data dapat meningkatkan generalisasi model secara substansial. Secara keseluruhan, hasil ini membuktikan bahwa deteksi dan penghapusan *outlier* merupakan langkah penting dalam *data preprocessing* untuk meningkatkan kualitas prediksi model *machine learning*. Peningkatan performa di semua algoritma mengindikasikan bahwa data sensor setelah tahap IQR menjadi lebih representatif dan stabil, sehingga model dapat mengenali pola karakteristik aroma dan rasa kopi dengan lebih baik.

3.4.2 Klasifikasi Setelah Reduksi Noise Data

Tahap selanjutnya dalam proses klasifikasi adalah penerapan model *machine learning* pada data yang telah melalui dua tahap pra-pemrosesan, yaitu penghapusan *outlier* menggunakan metode IQR dan reduksi *noise* menggunakan metode *Moving Average* (MA). Tiga algoritma klasifikasi yang sama digunakan pada tahap ini, yaitu *Random Forest* (RF), *Support Vector Classifier* (SVC), dan *Logistic Regression* (LR). Hasil pengujian menunjukkan adanya peningkatan performa yang signifikan dibandingkan tahap sebelumnya. Hasil perbandingan akurasi antara model sebelum dan sesudah reduksi *noise* data disajikan pada Tabel 4.

Tabel 4. Akurasi Model Setelah Reduksi Noise

Metode	Akurasi
Random Forest	94.36%
Support Vector Classifier (SVC)	81.70%
Logistic Regression	94.36%

Berdasarkan tabel di atas, terlihat bahwa model *Random Forest* dan *Logistic Regression* mencapai tingkat akurasi tertinggi, masing-masing sebesar 94.36%, sedangkan SVC memperoleh akurasi sebesar 81.70%. Hasil ini menunjukkan bahwa proses reduksi *noise* menggunakan *Moving Average* berhasil memperbaiki kestabilan sinyal dan membantu model mengenali pola sensorik dengan lebih baik. Peningkatan akurasi pada model *Logistic Regression* yang semula 73.58% menjadi 94.36% menunjukkan bahwa *smoothing* berperan penting bagi algoritma linear yang sensitif terhadap fluktuasi data. Sementara itu, *Random Forest* yang sebelumnya sudah memiliki performa tinggi tetap menunjukkan peningkatan, menandakan bahwa metode ini sangat efektif dalam menangkap variasi pola antar kelas setelah data distabilkan dan menunjukan metode ini menjadi *best model* pada klasifikasi kopi. Secara keseluruhan, hasil ini menegaskan bahwa kombinasi penghapusan *outlier* dan reduksi *noise* memberikan dampak positif yang signifikan terhadap performa model klasifikasi.

3.4.3 Klasifikasi Data Uji Menggunakan Model Terbaik

Setelah melalui tahapan pra-pemrosesan yang mencakup penghapusan *outlier* dengan metode *Interquartile Range* dan reduksi *noise* menggunakan *Moving Average*, dilakukan pengujian akhir terhadap model terbaik untuk menilai kemampuan generalisasinya pada data uji. Berdasarkan hasil evaluasi sebelumnya, *Random Forest* terpilih sebagai algoritma dengan performa tertinggi. Model *Random Forest* kemudian diuji menggunakan data uji yang telah diproses dengan metode yang sama. Evaluasi dilakukan menggunakan empat metrik utama, yaitu akurasi, presisi, recall, dan F1-score, guna memberikan gambaran menyeluruh terhadap kinerja model.

Tabel 5. Hasil Klasifikasi Data Uji Menggunakan Model Terbaik

Metode	Akurasi	Presisi	Recall	F1-score
Random Forest	96.03 %	96.15 %	96.07 %	96.05 %

Hasil pada Tabel 5 menunjukkan bahwa *Random Forest* memiliki performa yang sangat baik dan seimbang pada seluruh metrik evaluasi, dengan akurasi 96.03 %, presisi 96.15 %, recall 96.07 %, dan F1-score 96.05 %. Nilai-nilai tersebut mencerminkan konsistensi tinggi antara kemampuan model dalam mengenali kelas dengan benar (*recall*) dan menjaga ketepatan prediksi (*precision*). Performa yang hampir identik antara *precision*, *recall*, dan F1-score mengindikasikan bahwa model tidak hanya mampu mengidentifikasi sampel dengan benar, tetapi juga tidak cenderung memberikan prediksi yang keliru pada kelas tertentu. Dengan kata lain, tingkat kesalahan model relatif rendah dan distribusi prediksi antar kelas cukup stabil. Kinerja unggul *Random Forest* dalam penelitian ini dapat dijelaskan dari beberapa aspek. Secara teoretis, *Random Forest* merupakan algoritma ensemble berbasis bagging yang menggabungkan banyak pohon keputusan untuk meningkatkan ketahanan terhadap noise dan variabilitas data. Setiap pohon keputusan dibangun dari subset fitur dan subset sampel secara acak, sehingga menghasilkan keragaman (*diversity*) antar pohon. Keragaman ini penting karena mengurangi risiko model “belajar terlalu spesifik” terhadap karakteristik tertentu pada data latih, sebuah fenomena yang dikenal sebagai *overfitting*. Ketika hasil prediksi seluruh pohon digabungkan melalui mekanisme voting, model secara keseluruhan menjadi lebih stabil dan lebih mampu menangkap pola-pola non-linear yang kompleks pada data sensor. Kemampuan *Random Forest* untuk menangani data non-linear menjadi sangat relevan dalam konteks penelitian ini, mengingat data sensor aroma dan rasa dari *E-Nose* dan *E-tongue* memiliki struktur sinyal yang tidak sederhana. Respons sensor cenderung dipengaruhi oleh volatil aroma, dinamika waktu, serta fluktuasi nilai akibat kondisi lingkungan seperti suhu dan kelembapan. Oleh karena itu, model yang mampu mengakomodasi hubungan antar fitur yang tidak linier dan tidak tersusun secara seragam sangat diperlukan. *Random Forest* memberikan keunggulan ini melalui mekanisme pembentukan split decision yang adaptif pada setiap node pohon. Selain itu, keberhasilan *Random Forest* juga tidak terlepas dari tahapan pra-pemrosesan yang diterapkan sebelum pelatihan model. Penghapusan nilai ekstrem menggunakan metode IQR membantu menstabilkan distribusi data dengan mengurangi dampak pengukuran yang menyimpang akibat gangguan sensor atau ketidakkonsistenan selama akuisisi data. Tahap *smoothing* menggunakan *Moving Average* turut memperhalus fluktuasi sinyal, sehingga pola utama dalam karakteristik aroma dan rasa menjadi lebih terlihat. Ketika data yang masuk ke model sudah lebih terstruktur dan minim *noise*, proses pembelajaran pada *Random Forest* dapat berjalan lebih efektif. Dengan tingkat akurasi di atas 96 %, model *Random Forest* terbukti paling efektif dalam mengklasifikasikan jenis kopi Indonesia berdasarkan data gabungan sensor *Electronic Nose* dan *Electronic Tongue*. Temuan ini memperkuat kesimpulan bahwa kombinasi IQR + *Moving Average* + *Random Forest* merupakan pendekatan optimal untuk sistem klasifikasi berbasis sensor. Selain memberikan akurasi tinggi, pendekatan ini menawarkan *robustness* tinggi terhadap noise, interpretasi yang relatif mudah, serta kemampuan generalisasi yang kuat. Dengan demikian, metode ini berpotensi diterapkan pada proses autentikasi varietas kopi, deteksi campuran, dan pengawasan kualitas secara real-time dalam lingkungan industri pangan dan sektor pertanian yang membutuhkan ketepatan dan kecepatan analisis.

4. KESIMPULAN

Penelitian ini menunjukkan bahwa integrasi *Electronic Nose* dan *Electronic Tongue* dengan pendekatan *machine learning* mampu menjadi solusi yang efektif dan objektif dalam proses klasifikasi jenis kopi Indonesia. Sistem multi-sensor ini mampu menangkap karakteristik aroma dan rasa secara komprehensif sehingga menghasilkan representasi data yang lebih kaya dibandingkan metode pengujian konvensional. Tahap preprocessing memainkan peran krusial dalam keseluruhan proses analisis. Penghapusan outlier menggunakan metode Interquartile Range (IQR) dan perataan sinyal melalui Moving Average secara signifikan meningkatkan kualitas data dengan mengurangi anomali serta fluktuasi yang tidak relevan. Hasil peningkatan performa pada seluruh algoritma setelah tahap pembersihan data menguatkan temuan bahwa kualitas input menjadi faktor fundamental dalam keberhasilan proses klasifikasi berbasis sensor. Di antara tiga algoritma yang diuji, yaitu *Logistic Regression*, *Support Vector Classifier* (SVC), dan *Random Forest*, model *Random Forest* menunjukkan performa paling unggul dan paling stabil dengan akurasi tertinggi mencapai 96,03%. Keberhasilan ini tidak hanya menunjukkan kemampuan *Random Forest* dalam menangani data non-linear dan kompleks, tetapi juga memperlihatkan kecocokannya ketika dipadukan dengan teknik pra-pemrosesan berbasis IQR dan *Moving Average*. Kombinasi tersebut terbukti menghasilkan pola data yang lebih mudah dipelajari oleh model sehingga klasifikasi dapat dilakukan dengan tingkat kesalahan yang minimal. Secara praktis, hasil penelitian ini berpotensi menjadi dasar pengembangan sistem klasifikasi kopi otomatis yang lebih efisien, murah, dan dapat diimplementasikan dalam rantai pasok industri kopi di Indonesia, mulai dari proses pascapanen, verifikasi kualitas, hingga deteksi pemalsuan produk. Temuan ini membuka peluang penelitian lanjutan mengenai peningkatan akurasi, perluasan dataset, serta implementasi real-time pada perangkat portabel berbasis sensor.

UCAPAN TERIMA KASIH

Penulis menyampaikan ucapan terima kasih kepada Kementerian Pendidikan Tinggi Sains dan Teknologi melalui Program Penelitian Dosen Pemula (PDP) yang telah memberikan dukungan pendanaan dan kepercayaan sehingga penelitian ini dapat terlaksana dengan baik dan ucapan terima kasih juga disampaikan kepada Lembaga Penelitian dan Pengabdian kepada Masyarakat (LPPM) Universitas Nahdlatul Ulama Sunan Giri (UNUGIRI) atas pendampingan, arahan, serta fasilitas yang diberikan selama proses pelaksanaan penelitian hingga penyusunan laporan ini dapat diselesaikan dengan optimal.

REFERENCES

- [1] Yuli, "Produksi Kopi Indonesia dalam Dunia Internasional," Indonesia Baik, 2025, Accessed: Nov. 02, 2025. [Online]. Available: <https://indonesia.go.id/mediapublik/detail/2384>
- [2] N. Oktaviani et al., "Perubahan Iklim Mikro dan Produksi Kopi Arabika (*Coffea arabica* L.) pada Daerah Aktivitas Geothermal PLTP Kamojang di Kabupaten Bandung Microclimate change and arabica coffee (*Coffea arabica* L.) production in the geothermal activity area of the Kamojang GPP in Bandung Regency," Jurnal Agrikultura, vol. 2024, no. 3, pp. 400–412, 2024.
- [3] F. Adam, R. Agustina, and R. Fadhill, "Pengujian Cita Rasa Kopi Arabika Dengan Metode Cupping Test (Arabica Coffee Taste Testing With Cupping Test Method)," Jurnal Ilmiah Mahasiswa Pertanian, vol. 7, no. 1, 2022, [Online]. Available: www.jim.unsyiah.ac.id/JFP
- [4] B. Sumanto, Abelta Mika Setiari, Alfonzo Aruga Paripurna Barus, Iman Sabarisman, and Muhammad Arrofiq, "Pembelajaran Mesin Berbasis E-Nose Untuk Klasifikasi Daging Pada Produk Sosis," JST (Jurnal Sains dan Teknologi), vol. 13, no. 1, pp. 22–32, Apr. 2024, doi: 10.23887/jstundiksha.v13i1.70307.
- [5] P. Kecerdasan, B. Dalam, M. Higienitas, P. (Felix, K. Lie1, and A. Herwanto5, "Pemanfaatan Kecerdasan Buatan Dalam Meningkatkan Higienitas Pangan," 2023.
- [6] M. A. Barata, E. Noersasongko, Purwanto, and M. A. Soeleman, "Improving the Accuracy of C4.5 Algorithm with Chi-Square Method on Pure Tea Classification Using Electronic Nose," Jurnal RESTI, vol. 7, no. 2, pp. 226–235, Apr. 2023, doi: 10.29207/resti.v7i2.4687.
- [7] Wibowo Vadya Aryke, "STUDI PENGEMBANGAN E-TONGUE BERBASIS LIMA SENSOR ELEKTROKIMIA UNTUK KLASIFIKASI AIR MINUM," 2025.
- [8] D. Aulia, R. Sarno, S. C. Hidayati, A. N. Rosyid, and M. Rivai, "Identification of chronic obstructive pulmonary disease using graph convolutional network in Electronic Nose," Indonesian Journal of Electrical Engineering and Computer Science, vol. 34, no. 1, pp. 264–275, Apr. 2024, doi: 10.11591/ijeecs.v34i1.pp264-275.
- [9] A. Mujadin, S. Putra Rifaldi, O. Nur Samijayani, H. Suyono, A. Lastryanto, and P. Teknik Elektro, "Pengujian Kemurnian Minyak Kayu Putih Berbasis Electronic Nose Menggunakan Metode PCA Dan Neural Network," BULLET : Jurnal Multidisiplin Ilmu, 2024.
- [10] S. Wakhid, R. Sarno, S. I. Sabilla, and D. B. Maghfira, "Detection and classification of indonesian civet and non-civet coffee based on statistical analysis comparison using E-Nose," International Journal of Intelligent Engineering and Systems, vol. 13, no. 4, pp. 56–65, 2020, doi: 10.22266/IJIES2020.0831.06.
- [11] W. Harsono, R. Sarno, and S. Izza Sabilla, "Recognition of Original Arabica Civet Coffee based on Odor using Electronic Nose and Machine Learning," 2020.

- [12] D. Bayu Magfira and R. Sarno, "CLASSIFICATION OF ARABICA AND ROBUSTA COFFEE USING ELECTRONIC NOSE," 2018.
- [13] Maria Ulfa, Haryanto, and Kunto Aji Wibisono, "Desain Sistem Pengenalan dan Klasifikasi Kopi Bubuk Bermerek dengan Menggunakan Electronic Nose Berbasis Artifical Neural Network (ANN)," *J-Eltrik*, vol. 1, no. 2, p. 15, Nov. 2021, doi: 10.30649/j-eltrik.v1i2.15.
- [14] D. R. Wijaya, R. Sarno, and E. Zulaika, "Noise filtering framework for Electronic Nose signals: An application for beef quality monitoring," *Comput Electron Agric*, vol. 157, pp. 305–321, Feb. 2019, doi: 10.1016/j.compag.2019.01.001.
- [15] D. R. Prehanto, A. D. Indriyanti, I. K. D. Nuryana, and G. S. Permadi, "Classification based on K-Nearest Neighbor and Logistic Regression method of coffee using Electronic Nose," *IOP Conf Ser Mater Sci Eng*, vol. 1098, no. 3, p. 032080, Mar. 2021, doi: 10.1088/1757-899x/1098/3/032080.
- [16] S. Suhaila, D. Lelono, and Y. S. Sari, "Seleksi Fitur dengan Artificial Bee Colony untuk Optimasi Klasifikasi Data Teh menggunakan Support Vector Machine," *IJEIS (Indonesian Journal of Electronics and Instrumentation Systems)*, vol. 12, no. 1, p. 81, Apr. 2022, doi: 10.22146/ijeis.63902.
- [17] Nimatul Mamuriyah, Richard, and Haeruddin, "IMPLEMENTATION MEAN IMPUTATION AND OUTLIER DETECTION FOR LOAN PREDICTION USING THE RANDOM FOREST ALGORITHM," *JITK (Jurnal Ilmu Pengetahuan dan Teknologi Komputer)*, vol. 10, no. 4, pp. 937–944, Jun. 2025, doi: 10.33480/jitk.v10i4.6437.
- [18] S. Raharjo, L. Lestari, N. Hidayatullah, and M. Yaqub Rijal, "PENERAPAN METODE SIMPLE MOVING AVERAGE DALAM MEMPERHALUS DIFRAKTOGRAM XRD SAMPEL NIKEL," vol. 21, no. 01, pp. 5–13, 2025, doi: 10.62749/jaf.v21i01.p05-13.
- [19] Z. Zulkifli and R. Fajri, "Klasifikasi Tingkat Kematangan Buah Strawberry Menggunakan Algoritma Logistic Regression," *Data Sciences Indonesia (DSI)*, vol. 4, no. 2, pp. 50–59, Dec. 2024, doi: 10.47709/dsi.v4i2.4850.
- [20] D. I. Bagaskara, D. Syauqy, and B. H. Prasetyo, "Sistem Klasifikasi Kualitas Air untuk Budidaya Ikan Nila Hitam (*Oreochromis niloticus*) menggunakan Metode Support Vector Machine," 2022. [Online]. Available: <http://j-ptiik.ub.ac.id>
- [21] M. H. Madi, D. Syauqy, and W. Kurniawan, "Klasifikasi Kelayakan Konsumsi Susu Kedelai Berdasarkan PH, Warna, dan Aroma Menggunakan Metode Random Forest Berbasis Arduino," 2025. [Online]. Available: <http://j-ptiik.ub.ac.id>