

Pemilihan Fitur Untuk Prediksi Kematangan Kelapa Sawit Menggunakan *Explainable AI* (XAI)

Feature Selection for Oil Palm Ripeness Prediction Using Explainable AI (XAI)

Octo Bian Limet¹, Dwi Wahyu Prabowo^{*2}

^{1,2} Program Studi Sistem Informasi, Fakultas Ilmu Komputer, Universitas Darwan Ali

^{1,2} Jalan Batu Berlian No.10, Sampit, Kalimantan Tengah, Indonesia

email: ¹octobianlimet@gmail.com, ^{2*}dwi.wahyu9@unda.ac.id

Informasi Artikel

Dikirim, 18 Agustus 2025
Diterima, 17 September 2025
Diterbitkan, 5 Desember 2025

Kata Kunci :

Kelapa Sawit, *Machine Learning*, *Explainable AI*,
Seleksi Fitur, Prediksi

Keyword :

Palm Oil, Machine Learning,
Explainable AI, Feature
Selection, Prediction

ABSTRAK

Industri kelapa sawit merupakan sektor strategis di Indonesia, dengan kualitas tandan buah segar sebagai faktor utama yang menentukan mutu minyak. Penentuan tingkat kematangan kelapa sawit yang akurat sangat penting, namun metode manual yang masih digunakan saat ini rentan terhadap ketidakakuratan. Penelitian ini mengusulkan penerapan teknik pengolahan citra digital dan *machine learning* untuk klasifikasi tingkat kematangan kelapa sawit, dengan mengintegrasikan seleksi fitur menggunakan *recursive feature elimination* (RFE) serta metode *explainable AI* (*Shapley*) guna mengevaluasi potensi peningkatan performa sekaligus meningkatkan transparansi model. Dua model klasifikasi, yaitu *random forest* (RF) dan *support vector machine* (SVM), dievaluasi dengan dataset berisi 254 citra kelapa sawit yang terbagi ke dalam tiga kelas kematangan. Hasil eksperimen menunjukkan bahwa SVM mencapai akurasi tertinggi sebesar 72,78%, dengan nilai *precision* 73,83%, *recall* 71,53%, dan *F1-score* 71,08%, lebih unggul dibandingkan RF yang hanya memperoleh akurasi 61,45% tanpa seleksi fitur. Uji statistik *Mann-Whitney U* mengindikasikan bahwa seleksi fitur tidak memberikan peningkatan akurasi yang signifikan, tetapi SVM dengan kombinasi fitur warna (R, G, B) dan tekstur (GLCM) tetap mampu mengklasifikasikan kematangan kelapa sawit secara efektif. Analisis *Shapley* lebih lanjut mengungkapkan bahwa fitur warna, khususnya R dan G, memberikan kontribusi dominan terhadap hasil klasifikasi, dengan dukungan signifikan dari fitur tekstur pada kelas tertentu.

ABSTRACT

The palm oil industry is a strategic sector in Indonesia, where the quality of fresh fruit bunches (FFB) is a key determinant of oil yield. Accurate assessment of palm fruit ripeness is crucial; however, current manual methods remain prone to inaccuracy. This study proposes the application of digital image processing and machine learning for ripeness classification, integrating feature selection through recursive feature elimination (RFE) and explainable AI (Shapley) to evaluate both performance improvements and model transparency. Two classification models, Random Forest (RF) and Support Vector Machine (SVM), were evaluated using a dataset of 254 palm fruit images categorized into three ripeness classes. Experimental results show that SVM achieved the highest accuracy of 72,78%, with precision of 73,83%, recall of 71,53%, and F1-score of 71,08%, outperforming RF, which achieved only 61,45% accuracy without feature selection. Statistical analysis using the Mann-Whitney U test indicates that feature selection does not yield significant accuracy improvements; nevertheless, SVM utilizing color features (R, G, B) and texture features (GLCM) effectively classified ripeness levels. Further Shapley analysis reveals that color features, particularly R and G, contribute most dominantly to classification outcomes, with texture features providing additional discriminative support in certain classes

1. PENDAHULUAN

Kelapa sawit (*Elaeis guineensis* Jacq.) merupakan salah satu komoditas perkebunan strategis di Indonesia yang memiliki kontribusi signifikan terhadap perekonomian nasional [1]. Indonesia menempati posisi sebagai produsen dan eksportir minyak kelapa sawit terbesar di dunia, dengan nilai produksi yang terus meningkat setiap tahunnya. Keberhasilan industri kelapa sawit sangat bergantung pada kualitas tandan buah segar (TBS) yang dipanen, di mana tingkat kematangan buah menjadi faktor utama penentu kualitas minyak yang dihasilkan. Kesalahan dalam menentukan tingkat kematangan dapat menyebabkan penurunan rendemen minyak dan kualitas produk akhir, sehingga diperlukan metode penentuan kematangan yang akurat dan konsisten [2]. Selama ini, penentuan kematangan kelapa sawit di lapangan umumnya masih dilakukan secara manual berdasarkan pengamatan visual pekerja [3]. Metode ini bersifat subjektif, dipengaruhi oleh pengalaman individu, serta rentan terhadap kesalahan akibat kondisi pencahayaan dan kelelahan pengamat. Ketidakakuratan ini dapat berdampak langsung pada efisiensi panen dan kualitas hasil produksi. Oleh karena itu, diperlukan metode berbasis teknologi yang dapat membantu proses identifikasi kematangan secara objektif, cepat, dan dapat direplikasi di berbagai kondisi.

Sejumlah penelitian sebelumnya telah mengkaji penerapan pengolahan citra digital dan *machine learning* untuk memprediksi tingkat kematangan kelapa sawit. Teknik ekstraksi fitur yang umum digunakan meliputi nilai rata-rata kanal warna RGB, HSV, serta fitur tekstur yang diperoleh dari *gray level co-occurrence matrix* (GLCM) [4], [5]. Dalam proses klasifikasi, algoritma seperti *random forest* (RF) dan *support vector machine* (SVM) sering digunakan [6], [7]. Hasil dari penelitian terdahulu menunjukkan bahwa kombinasi fitur warna dan tekstur mampu mencapai akurasi prediksi di atas 85%. Namun demikian, beberapa keterbatasan masih ditemukan, di antaranya jumlah dataset yang relatif terbatas, kurangnya penerapan analisis keterjelasan model (*explainability*) [8]–[10], serta belum optimalnya pemilihan fitur yang benar-benar relevan terhadap target prediksi. Sebagian besar penelitian terdahulu juga cenderung berfokus pada peningkatan akurasi, tanpa memberikan penjelasan interpretatif mengenai alasan di balik keputusan model. Hal ini menyulitkan penerapan hasil penelitian secara langsung di lapangan, terutama bagi pihak yang membutuhkan transparansi dalam pengambilan keputusan panen. Transparansi model menjadi hal penting karena pekerja lapangan ataupun industri, memerlukan justifikasi yang jelas mengenai fitur apa yang berkontribusi dalam penentuan label kelas pada suatu dataset [11], [12]. Dengan adanya interpretabilitas, keputusan panen tidak hanya bergantung pada hasil prediksi, tetapi juga dapat diverifikasi dan dipertanggungjawabkan, sehingga meningkatkan kepercayaan serta adopsi teknologi di industri kelapa sawit.

Berdasarkan kesenjangan tersebut, penelitian ini menawarkan pendekatan yang mengintegrasikan teknik pemilihan fitur (*feature selection*) yaitu *recursive feature elimination* (RFE) dengan metode *explainable AI* (XAI) terhadap model klasifikasi RF dan SVM [13], [14]. Tujuannya adalah mengevaluasi potensi peningkatan performa sekaligus meningkatkan transparansi model berupa interpretasi yang jelas terkait kontribusi setiap fitur terhadap hasil klasifikasi. Dengan demikian, metode yang diusulkan diharapkan dapat menjadi solusi yang lebih akurat dan transparan untuk mendukung proses panen yang optimal serta meningkatkan efisiensi dan kualitas produksi kelapa sawit.

2. METODE PENELITIAN

Metode penelitian menjelaskan prosedur yang digunakan dalam penelitian, mencakup sumber dan karakteristik dataset, tahapan pemrosesan data, serta metode yang digunakan dalam pembangunan dan evaluasi model. Penjelasan disusun secara sistematis untuk memastikan keterulangan proses dan transparansi metodologi. Bagian awal menguraikan detail dataset yang menjadi dasar penelitian, diikuti penjelasan mengenai tahapan penelitian mulai dari pemuatan data, ekstraksi dan pemilihan fitur, pelatihan model, hingga analisis hasil. Seluruh langkah dirancang untuk mendukung tujuan penelitian dalam mengklasifikasikan tingkat kematangan buah kelapa sawit secara akurat dan terukur.

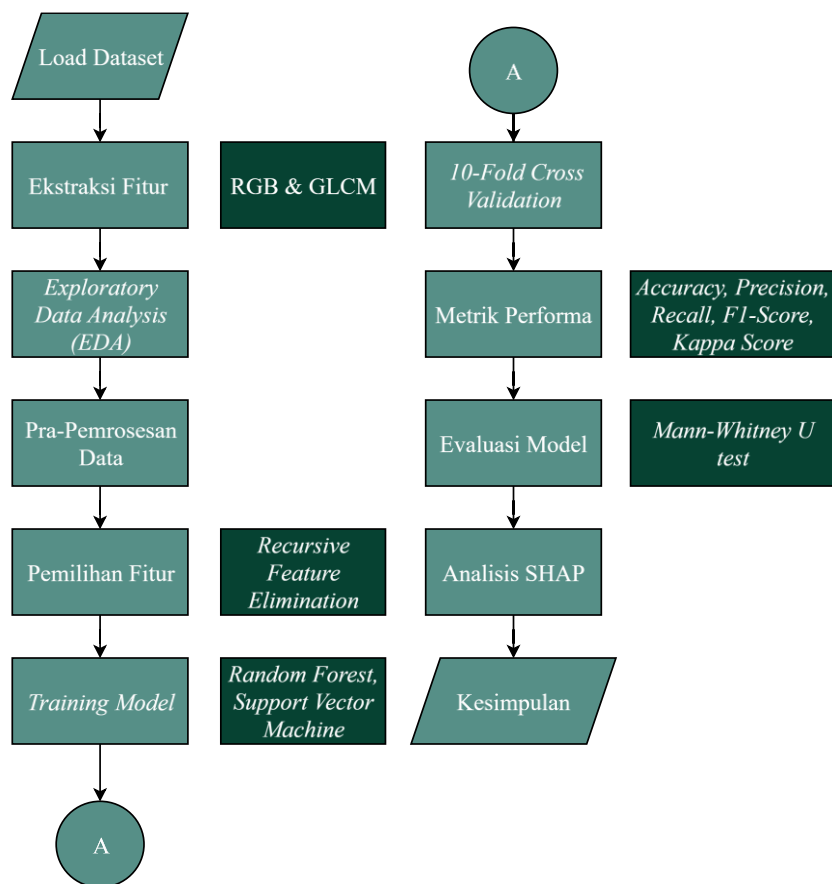
2.1. Dataset

Dataset yang digunakan pada penelitian ini bersumber dari repositori publik *Mendeley Data* (<https://data.mendeley.com/datasets/zrcf6v69y8/1>), yang terdiri atas 254 citra buah kelapa sawit dalam format JPEG. Citra tersebut dikelompokkan ke dalam tiga kelas tingkat kematangan, yaitu masak, mengkal, dan mentah, yang masing-masing merepresentasikan perbedaan visual pada warna dan tekstur buah. Proses pengumpulan dataset dilakukan oleh penyedia sumber data, sehingga penelitian ini memanfaatkan data sekunder tanpa melakukan pengambilan citra secara langsung. Selanjutnya, setiap citra diekstraksi menjadi dua jenis fitur utama, yaitu fitur warna berbasis nilai rata-rata kanal RGB, serta fitur tekstur yang dihitung menggunakan metode GLCM. Hasil ekstraksi fitur disimpan dalam format CSV untuk memudahkan pengolahan pada tahap pemodelan. Dataset ini dipilih karena memiliki kualitas visual yang memadai serta representasi yang seimbang pada setiap kelas, sehingga relevan untuk proses klasifikasi berbasis *machine learning*.

Dataset yang digunakan memiliki distribusi kelas yang relatif seimbang, yaitu 91 citra untuk kelas masak, 72 citra untuk kelas mengkal, dan 91 citra untuk kelas mentah. Distribusi ini membantu mengurangi risiko ketidakseimbangan kelas yang dapat memengaruhi proses pelatihan model, meskipun jumlah data yang terbatas tetap menjadi kendala. Karakteristik citra yang tersedia hanya mencakup kondisi visual tertentu, sehingga potensi bias dapat muncul apabila model diterapkan pada citra dengan pencahayaan, sudut pengambilan, atau kualitas kamera yang berbeda dari dataset ini. Keterbatasan tersebut berimplikasi pada kemampuan generalisasi model, sehingga uji pada dataset dari sumber berbeda diperlukan untuk menilai sejauh mana model mampu beradaptasi dengan variasi data yang lebih luas.

2.2. Tahapan Penelitian

Tahapan penelitian ini disusun untuk memberikan gambaran sistematis mengenai alur proses yang digunakan dalam pengembangan model prediksi kematangan buah kelapa sawit. Setiap tahap dirancang agar saling terhubung mulai dari pemuatan dataset, pengolahan dan pemilihan fitur, pelatihan serta validasi model, hingga analisis interpretasi hasil. Diagram tahapan penelitian pada Gambar 1 menyajikan urutan proses secara ringkas, sedangkan penjelasan rinci setiap tahap disajikan pada bagian berikut:



Gambar 1. Tahapan Penelitian

2.2.1. Load Dataset

Tahap pertama adalah memuat dataset citra kelapa sawit ke dalam penyimpanan *google drive* sehingga dapat diakses di *google colab*. Dataset kemudian dibaca dan dilakukan penyesuaian struktur folder untuk memastikan kesesuaian dengan kebutuhan proses ekstraksi fitur. Penataan struktur ini bertujuan mempermudah pengelolaan data dan mengurangi potensi kesalahan pada tahap berikutnya.

2.2.2. Ekstraksi Fitur

Setelah dataset siap, dilakukan ekstraksi fitur dari setiap citra menggunakan metode RGB untuk informasi warna dan GLCM untuk informasi tekstur. Nilai RGB mewakili intensitas warna *merah*, *hijau* dan *biru* sedangkan GLCM digunakan untuk menghitung atribut tekstur seperti *kontras*, *homogenitas*, *energi* dan *korelasi*. Hasil ekstraksi disimpan dalam bentuk numerik yang menjadi masukan untuk proses analisis dan klasifikasi.

2.2.3. Exploratory Data Analysis (EDA)

EDA dilakukan untuk memahami karakteristik data dan hubungan antar fitur. Analisis ini mencakup pemeriksaan distribusi nilai fitur, deteksi adanya *outlier* serta identifikasi korelasi antar atribut. Visualisasi yang digunakan untuk memberikan gambaran awal terhadap pola data.

2.2.4. Pra-Pemrosesan Data

Tahap pra-pemrosesan meliputi pengubahan label kelas (masak, mengkal, mentah) menjadi format numerik menggunakan *label encoder* dan normalisasi nilai fitur dengan *standard scaler*. *Encoding* diperlukan agar label dapat diproses oleh algoritma klasifikasi, sedangkan normalisasi memastikan semua fitur berada pada skala yang sama. Langkah ini membantu meningkatkan efisiensi dan akurasi model pada tahap pelatihan.

2.2.5. Pemilihan Fitur

Pemilihan fitur dilakukan menggunakan *recursive feature elimination* (RFE) yang bekerja dengan melatih model awal, mengevaluasi kontribusi setiap fitur, dan menghapus fitur dengan kontribusi terendah secara bertahap. Proses ini diulang hingga diperoleh subset fitur yang paling relevan terhadap target klasifikasi. Pada penelitian ini, algoritma *random forest* dan *support vector machine* digunakan sebagai estimator di dalam RFE untuk menilai pentingnya setiap fitur yang dipilih/diseleksi. Tujuan seleksi fitur adalah mengurangi kompleksitas model, mempercepat pelatihan, dan meningkatkan generalisasi hasil.

2.2.6. Training Model

Setelah diperoleh fitur terpilih, dua algoritma klasifikasi yaitu *random forest* dan *support vector machine* dilatih secara terpisah. Pemilihan kedua algoritma ini didasarkan pada kemampuannya menangani data berdimensi tinggi serta perbedaan pendekatan dalam proses klasifikasi, sehingga memungkinkan perbandingan kinerja yang komprehensif.

2.2.7. Validasi Model

Validasi dilakukan menggunakan metode *10-fold cross-validation*. *10-fold cross-validation* merupakan pendekatan evaluasi model yang sering dipakai dalam *machine learning*, terutama ketika ingin mendapatkan gambaran umum mengenai kinerja model yang stabil dan tidak terpengaruh bias [15]. Dataset dibagi menjadi sepuluh subset dengan proporsi yang seimbang. Pada setiap iterasi, sembilan subset digunakan untuk pelatihan dan satu subset digunakan untuk pengujian, lalu proses diulang hingga semua subset digunakan sebagai data uji.

2.2.8. Metrik Performa

Metrik performa dalam *machine learning* merupakan ukuran kuantitatif yang digunakan untuk mengevaluasi kualitas prediksi yang dihasilkan oleh model klasifikasi atau regresi. Performa model dievaluasi menggunakan lima metrik yaitu *accuracy*, *precision*, *recall*, *f1-score* dan *Cohen's Kappa Score*.

1. Accuracy

Akurasi merupakan metrik evaluasi dasar yang mengukur seberapa sering model menghasilkan prediksi yang benar terhadap total keseluruhan data [16]. Berikut persamaan dari *akurasi* ditunjukkan pada persamaan (1).

$$\text{Akurasi} = \frac{TP + TN}{TP + TN + FP + FN} \quad (1)$$

2. Precision

Presisi digunakan untuk mengukur ketepatan model dalam memprediksi kelas positif [17]. Berikut persamaan dari *presisi* bisa dilihat pada persamaan (2).

$$\text{Presisi} = \frac{TP}{TP + FP} \quad (2)$$

3. Recall

Recall (juga dikenal sebagai sensitivitas) mengukur sejauh mana model mampu menemukan seluruh data yang benar-benar positif [18]. Berikut persamaanya bisa dilihat pada persamaan (3).

$$\text{Recall} = \frac{TP}{TP + FN} \quad (3)$$

4. F1-Score

F1-Score adalah metrik yang menggabungkan *presisi* dan *recall* kedalam satu nilai tunggal [19]. Untuk persamaanya ditunjukkan pada persamaan (4).

$$F1 - Score = 2 * \frac{Presisi * Recall}{Presisi + Recall} \quad (4)$$

5. *Cohen's Kappa Score*

Kappa Score atau *Cohen's Kappa* adalah metrik statistik yang digunakan untuk mengukur tingkat kesepakatan antara dua pengklasifikasi, yaitu model dan label aktual dengan mempertimbangkan kemungkinan kesepakatan yang terjadi secara kebetulan [20]. Berikut persamaan dari *kappa* ditunjukkan pada persamaan (5)

$$k = \frac{p_o - p_e}{1 - p_e} \quad (5)$$

2.2.9. Evaluasi Model

Uji *Mann-Whitney*, yang juga dikenal sebagai *Mann-Whitney U Test* atau *Wilcoxon Rank-Sum Test*, merupakan salah satu metode statistik non-parametrik yang digunakan untuk membandingkan dua sampel independen dalam hal median atau distribusi secara umum [21]. Tujuannya adalah menentukan apakah perbedaan performa kedua model signifikan secara statistik. Pada penelitian ini, uji dilakukan dengan signifikansi sebesar 5% untuk membandingkan performa model klasifikasi yang akan di uji baik tanpa maupun dengan seleksi/pemilihan fitur. Hasil pengujian ini diharapkan dapat memberikan hasil statistik dari perbandingan performa model.

2.2.10. Analisis Shapley

Analisis model dilakukan menggunakan metode SHAP (*SHapley Additive exPlanations*) untuk mengukur kontribusi setiap fitur terhadap hasil prediksi. SHAP adalah salah satu metode *explainable AI* (XAI) yang digunakan untuk menjelaskan bagaimana sebuah model *machine learning* membuat prediksi [22]. Metode ini menghasilkan visualisasi yang menunjukkan fitur mana yang paling memengaruhi keputusan model. Tahap ini memberikan interpretasi terhadap kinerja model dan mendukung transparansi sistem prediksi.

3. HASIL DAN PEMBAHASAN

3.1. Hasil Eksperimen

Pada eksperimen ini, kedua model klasifikasi dengan semua fitur dan seleksi fitur dilatih dan dievaluasi menggunakan *stratified 10-Fold Cross Validation*. Eksperimen ini dilakukan dengan dataset kelapa sawit dengan tiga label/kelas tingkat kematangan, yaitu: masak, mengkal, dan mentah.

3.1.1. Random Forest Tanpa Seleksi Fitur

Tabel 1 menyajikan hasil rata-rata *10-Fold Cross Validation* pada algoritma *random forest* tanpa seleksi fitur. Nilai *akurasi* yang diperoleh sebesar 0,6145 dengan standar deviasi 0,1004, menunjukkan variasi kinerja yang moderat antar *fold*. Metrik *kappa* tercatat sebesar $0,4134 \pm 0,1492$, mengindikasikan tingkat kesepakatan yang cukup antara prediksi model dan label/kelas sebenarnya pada ketiga kelas kelapa sawit. Nilai *precision*, *recall* dan *F1-Score* masing-masing sebesar $0,6034 \pm 0,1191$, $0,5935 \pm 0,0988$ dan $0,5888 \pm 0,1028$, yang menggambarkan keseimbangan performa model dalam mengidentifikasi rata-rata performa pada semua kelas.

Tabel 1. Rata-rata Hasil 10-Fold CV RF Tanpa Seleksi Fitur

Metrik	Nilai (Mean $\pm$ Std)
<i>Akurasi</i>	$0,6145 \pm 0,1004$
<i>Kappa</i>	$0,4134 \pm 0,1492$
<i>Precision</i>	$0,6034 \pm 0,1191$
<i>Recall</i>	$0,5935 \pm 0,0988$
<i>F1-Score</i>	$0,5888 \pm 0,1028$

3.1.2. Random Forest Dengan Seleksi Fitur

Tabel 2 menunjukkan hasil rata-rata *10-Fold Cross Validation* pada algoritma *random forest* setelah dilakukan seleksi fitur. Nilai *akurasi* yang diperoleh sebesar 0,5788 dengan standar deviasi 0,1232, yang mengindikasikan adanya penurunan kinerja dibandingkan tanpa seleksi fitur. Metrik *kappa* tercatat sebesar $0,3610 \pm 0,1829$, menunjukkan tingkat kesepakatan yang lebih rendah antara prediksi model dan label/kelas sebenarnya pada ketiga kelas kelapa sawit. Nilai *precision*, *recall* dan *F1-Score* masing-masing sebesar $0,5705 \pm 0,1389$, $0,5594 \pm 0,1237$ dan $0,5545 \pm 0,1249$, yang merefleksikan penurunan konsistensi model dalam mengidentifikasi rata-rata performa pada semua kelas.

Tabel 2. Rata-rata Hasil 10-Fold CV RF Dengan Seleksi Fitur

Metrik	Nilai (Mean $\pm$ Std)
Akurasi	0,5788 $\pm$ 0,1232
Kappa	0,3610 $\pm$ 0,1829
Precision	0,5705 $\pm$ 0,1389
Recall	0,5594 $\pm$ 0,1237
F1-Score	0,5545 $\pm$ 0,1249

3.1.3. Support Vector Machine Tanpa Seleksi Fitur

Tabel 3 menunjukkan hasil rata-rata *10-fold cross validation* pada algoritma *support vector machine* tanpa seleksi fitur. Nilai *akurasi* yang diperoleh sebesar 0,6892 dengan standar deviasi 0,0739, menunjukkan kinerja yang lebih stabil dibandingkan *random forest* pada tabel sebelumnya. Metrik *kappa* tercatat sebesar 0,5305 $\pm$ 0,1088, yang mengindikasikan tingkat kesepakatan model dengan label sebenarnya pada ketiga kelas kelapa sawit berada pada kategori moderat. *Precision*, *recall* dan *F1-Score* masing-masing sebesar 0,6830 $\pm$ 0,0689, 0,6745 $\pm$ 0,0709 dan 0,6681 $\pm$ 0,0749, yang merepresentasikan rata-rata performa model dalam mengklasifikasikan semua kelas.

Tabel 3. Rata-rata Hasil 10-Fold CV SVM Tanpa Seleksi Fitur

Metrik	Nilai (Mean $\pm$ Std)
Akurasi	0,6892 $\pm$ 0,0739
Kappa	0,5305 $\pm$ 0,1088
Precision	0,6830 $\pm$ 0,0689
Recall	0,6745 $\pm$ 0,0709
F1-Score	0,6681 $\pm$ 0,0749

3.1.4. Support Vector Machine Dengan Seleksi Fitur

Tabel 4 menyajikan hasil rata-rata *10-Fold Cross Validation* pada algoritma *support vector machine* setelah dilakukan seleksi fitur. Nilai *akurasi* yang diperoleh sebesar 0,7278 dengan standar deviasi 0,0597, menunjukkan adanya peningkatan kinerja dibandingkan tanpa seleksi fitur. Metrik *kappa* sebesar 0,5901 $\pm$ 0,0857 mengindikasikan tingkat kesepakatan yang lebih tinggi antara prediksi model dan label sebenarnya pada ketiga kelas kelapa sawit. *Precision*, *recall* dan *F1-Score* masing-masing sebesar 0,7383 $\pm$ 0,0436, 0,7153 $\pm$ 0,0534 dan 0,7108 $\pm$ 0,0524, yang mencerminkan rata-rata performa model dalam mengklasifikasikan semua kelas secara konsisten.

Tabel 4. Rata-rata Hasil 10-Fold CV SVM Dengan Seleksi Fitur

Metrik	Nilai (Mean $\pm$ Std)
Akurasi	0,7278 $\pm$ 0,0597
Kappa	0,5901 $\pm$ 0,0857
Precision	0,7383 $\pm$ 0,0436
Recall	0,7153 $\pm$ 0,0534
F1-Score	0,7108 $\pm$ 0,0524

3.2. Pembahasan Hasil Eksperimen

Untuk mendapatkan pemahaman yang lebih mendalam mengenai perbedaan performa antar rata-rata hasil pada skema eksperimen yang didapatkan, dilakukan uji statistik yang bertujuan untuk mengidentifikasi apakah terdapat perbedaan yang signifikan secara statistik. Uji statistik dilakukan dengan metode uji statistik *Mann-Whitney* yang sangat penting dilakukan untuk memastikan bahwa perbedaan yang diamati bukanlah hasil dari kebetulan, melainkan mencerminkan perbedaan yang nyata dan dapat diandalkan antara setiap pendekatan yang diuji.

3.2.1. RF Tanpa Seleksi Fitur Vs RF Dengan Seleksi Fitur

Tabel 5 menyajikan hasil uji statistik *Mann-Whitney* untuk membandingkan kinerja *random forest* tanpa seleksi fitur dan dengan seleksi fitur pada lima metrik evaluasi. Nilai *p-value* pada seluruh metrik berada di atas 0,05, yang menunjukkan tidak terdapat perbedaan kinerja yang signifikan secara statistik antara kedua skenario. Nilai *U-statistic* yang sama pada setiap metrik mengindikasikan distribusi nilai yang relatif mirip antar kedua kelompok. Berdasarkan hasil ini, penerapan seleksi fitur pada *random forest* tidak memberikan peningkatan performa yang signifikan.

Tabel 5. Uji *Mann-Whitney* RF Tanpa Seleksi Fitur Vs RF Dengan Seleksi Fitur

Metrik	RF Tanpa Seleksi Fitur	RF Dengan Seleksi Fitur	Uji <i>Mann-Whitney</i>		
			<i>U-statistic</i>	<i>P-value</i>	Signifikan
<i>Akurasi</i>	0,6145	0,5788	58,0	0,569587208	Tidak
<i>Kappa</i>	0,4134	0,3610	58,0	0,570460473	Tidak
<i>Precision</i>	0,6034	0,5705	58,0	0,570460473	Tidak
<i>Recall</i>	0,5935	0,5594	58,0	0,520209006	Tidak
<i>F1-Score</i>	0,5888	0,5545	58,0	0,570460473	Tidak

3.2.2. SVM Tanpa Seleksi Fitur Vs SVM Dengan Seleksi Fitur

Tabel 6 menyajikan hasil uji statistik *Mann-Whitney* untuk membandingkan kinerja SVM tanpa seleksi fitur dan SVM dengan seleksi fitur pada lima metrik evaluasi. Seluruh nilai *p-value* berada di atas 0,05, termasuk pada *precision* yang memiliki nilai terendah sebesar 0,0539, sehingga tidak terdapat perbedaan kinerja yang signifikan secara statistik antara kedua skenario. Nilai *U-statistic* bervariasi antar metrik, menunjukkan adanya perbedaan peringkat nilai performa, namun tidak cukup kuat untuk mencapai signifikansi. Dengan demikian, penerapan seleksi fitur pada SVM tidak memberikan peningkatan performa yang signifikan secara statistik.

Tabel 6. Uji *Mann-Whitney* SVM Tanpa Seleksi Fitur Vs SVM Dengan Seleksi Fitur

Metrik	SVM Tanpa Seleksi Fitur	SVM Dengan Seleksi Fitur	Uji <i>Mann-Whitney</i>		
			<i>U-statistic</i>	<i>P-value</i>	Signifikan
<i>Akurasi</i>	0,6892	0,7278	33,5	0,222433829	Tidak
<i>Kappa</i>	0,5305	0,5901	34,0	0,240258166	Tidak
<i>Precision</i>	0,6830	0,7383	24,0	0,053902557	Tidak
<i>Recall</i>	0,6745	0,7153	32,5	0,198426825	Tidak
<i>F1-Score</i>	0,6681	0,7108	34,0	0,241321593	Tidak

3.2.3. RF Tanpa Seleksi Fitur Vs SVM Tanpa Seleksi Fitur

Tabel 7 menyajikan hasil uji statistik *Mann-Whitney* untuk membandingkan kinerja *random forest* dan *support vector machine* tanpa seleksi fitur pada lima metrik evaluasi. Seluruh nilai *p-value* berada di atas 0,05, termasuk *recall* dan *F1-Score* yang memiliki nilai terendah masing-masing sebesar 0,0887 dan 0,0757, sehingga tidak terdapat perbedaan kinerja yang signifikan secara statistik antara kedua algoritma klasifikasi. Variasi nilai *U-statistic* menunjukkan perbedaan peringkat performa pada tiap metrik, namun tidak cukup untuk mencapai tingkat signifikansi. Hasil ini mengindikasikan bahwa pada kondisi tanpa seleksi fitur, perbedaan kinerja antara *random forest* dan *support vector machine* tidak signifikan secara statistik.

Tabel 7. Uji Mann-Whitney RF Tanpa Seleksi Fitur Vs SVM Tanpa Seleksi Fitur

Metrik	RF Tanpa Seleksi Fitur	SVM Tanpa Seleksi Fitur	<i>U-statistic</i>	Uji Mann-Whitney <i>P-value</i>	Signifikan
<i>Akurasi</i>	0,6145	0,6892	28,0	0,102286298	Tidak
<i>Kappa</i>	0,4134	0,5305	25,0	0,063922135	Tidak
<i>Precision</i>	0,6034	0,6830	28,0	0,10410989	Tidak
<i>Recall</i>	0,5935	0,6745	27,0	0,088732801	Tidak
<i>F1-Score</i>	0,5888	0,6681	26,0	0,075661572	Tidak

3.2.4. RF Dengan Seleksi Fitur Vs SVM Dengan Seleksi Fitur

Tabel 8 menyajikan hasil uji statistik *Mann-Whitney* untuk membandingkan kinerja *random forest* dan *support vectot machine* setelah dilakukan seleksi fitur pada lima metrik evaluasi. Seluruh nilai *p-value* berada di bawah 0,05, sehingga perbedaan kinerja antara kedua algoritma klasifikasi signifikan secara statistik pada semua metrik. Nilai *U-statistic* yang relatif rendah mengindikasikan bahwa nilai performa SVM secara konsisten lebih tinggi dibandingkan RF pada skenario ini. Hasil ini menunjukkan bahwa setelah seleksi fitur, *support vector machine* memberikan kinerja yang secara signifikan lebih baik dibandingkan *random forest*.

Tabel 8. Uji Mann-Whitney RF Dengan Seleksi Fitur Vs SVM Dengan Seleksi Fitur

Metrik	RF Dengan Seleksi Fitur	SVM Dengan Seleksi Fitur	<i>U-statistic</i>	Uji Mann-Whitney <i>P-value</i>	Signifikan
<i>Akurasi</i>	0,5788	0,7278	11,0	0,003510222	Ya
<i>Kappa</i>	0,3610	0,5901	9,0	0,002185327	Ya
<i>Precision</i>	0,5705	0,7383	14,0	0,007284557	Ya
<i>Recall</i>	0,5594	0,7153	10,0	0,002796242	Ya
<i>F1-Score</i>	0,5545	0,7108	12,0	0,004586392	Ya

Berdasarkan Pembahasan hasil eksperimen diatas, jika dibandingkan dengan penelitian sebelumnya pada bagian pendahuluan, di mana kombinasi fitur warna dan tekstur mampu mencapai akurasi prediksi di atas 85%, hasil terbaik pada penelitian ini yaitu SVM dengan seleksi fitur hanya mencapai akurasi sebesar 72.78%. Perbedaan ini kemungkinan disebabkan oleh variasi dalam ukuran dan kualitas dataset, perbedaan metode ekstraksi fitur, serta perbedaan strategi evaluasi model yang digunakan. Selain itu, keterbatasan jumlah citra dalam dataset dapat memengaruhi kemampuan model untuk belajar pola yang lebih beragam, sehingga berdampak pada performa yang lebih rendah. Keterbatasan ini juga berimplikasi pada kemampuan generalisasi model, sehingga penerapannya pada kondisi nyata di lapangan perlu dilakukan dengan hati-hati. Untuk penelitian selanjutnya, penggunaan dataset yang lebih besar dan beragam sangat diperlukan agar hasil prediksi lebih representatif.

3.2.5 Analisis Effect Size dan Power Analysis

Analisis menggunakan *effect size* dan *power analysis* yang merujuk kepada Tabel 9 memberikan perspektif yang lebih komprehensif terhadap hasil uji Mann–Whitney. Pada perbandingan RF tanpa dan dengan seleksi fitur, nilai *effect size* sangat kecil ($r \approx 0,135$) dengan *observed power* hanya sekitar 0,09, sehingga meskipun tidak signifikan secara statistik, hasil ini juga tidak memiliki arti praktis yang kuat. Sebaliknya, pada perbandingan SVM tanpa dan dengan seleksi fitur, efek yang ditunjukkan cenderung kecil hingga sedang ($r \approx 0,27–0,44$), dengan *power* bervariasi antara 0,23–0,50; hal ini mengindikasikan adanya potensi perbedaan praktis, khususnya pada metrik *precision* ($r \approx 0,44$; $power \approx 0,50$), meskipun kekuatan uji masih terbatas. Perbandingan RF dan SVM tanpa seleksi fitur menghasilkan *effect size* sedang ($r \approx 0,37–0,42$) dengan *power* yang moderat (0,38–0,47), yang menegaskan keunggulan SVM meskipun dengan risiko kesalahan tipe II yang masih ada. Sementara itu, hasil yang paling menonjol adalah perbandingan RF dan SVM setelah seleksi fitur menunjukkan *effect size* besar ($r \approx 0,61–0,69$) dengan *observed power* tinggi (0,78–0,87), sehingga keunggulan SVM atas RF pada kondisi ini dapat dikatakan kuat baik secara statistik maupun praktis. Dengan demikian, *power analysis* mendukung kesimpulan bahwa hasil penelitian ini andal hanya untuk efek besar, sementara deteksi efek kecil hingga sedang masih memerlukan ukuran sampel yang lebih besar atau desain eksperimental yang lebih kompleks.

Tabel 9. Hasil Uji Mann–Whitney dengan Effect Size dan Power

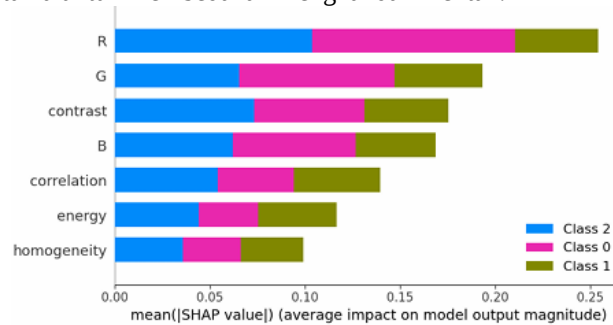
Perbandingan Model	Metrik	U-statistic	Z-value	Effect Size (r)	Kategori Efek	Power ($\approx$)
RF tanpa vs dengan seleksi	Accuracy	58,0	0,605	0.135	Small	0,09
	Kappa	58,0	0,605	0.135	Small	0,09
	Precision	58,0	0,605	0.135	Small	0,09
	Recall	58,0	0,605	0.135	Small	0,09
	F1-score	58,0	0,605	0.135	Small	0,09
SVM tanpa vs dengan seleksi	Accuracy	33,5	-1,247	0.279	Small–Medium	0,24
	Kappa	34,0	-1,209	0.270	Small–Medium	0,23
	Precision	24,0	-1,965	0.439	Medium	0,50
	Recall	32,5	-1,323	0.296	Small–Medium	0,26
	F1-score	34,0	-1,209	0.270	Small–Medium	0,23
RF tanpa vs SVM tanpa seleksi	Accuracy	28,0	-1,663	0.372	Medium	0,38
	Kappa	25,0	-1,890	0.423	Medium	0,47
	Precision	28,0	-1,663	0.372	Medium	0,38
	Recall	27,0	-1,739	0.389	Medium	0,41
	F1-score	26,0	-1,814	0.406	Medium	0,44
RF dengan vs SVM dengan seleksi	Accuracy	11,0	-2,948	0.659	Large	0,84
	Kappa	9,0	-3,099	0.693	Large	0,87
	Precision	14,0	-2,721	0.609	Large	0,78
	Recall	10,0	-3,024	0.676	Large	0,86
	F1-score	12,0	-2,873	0.642	Large	0,82

3.3. Analisis Shapley Untuk Skema Yang Bersesuaian

Analisis *Shapley* dilakukan untuk mengevaluasi kontribusi relatif dari setiap fitur terhadap keputusan model, baik pada skema tanpa seleksi fitur maupun dengan seleksi fitur. Pendekatan ini memberikan pemahaman yang lebih mendalam mengenai variabel mana yang paling berpengaruh dalam proses klasifikasi tingkat kematangan kelapa sawit. Dengan membandingkan hasil pada model *random forest* dan *support vector machine*, analisis ini diharapkan dapat mengungkap perbedaan pola ketergantungan model terhadap fitur warna dan tekstur serta menilai dampak seleksi fitur terhadap interpretabilitas prediksi.

3.3.1. RF Tanpa Seleksi Fitur

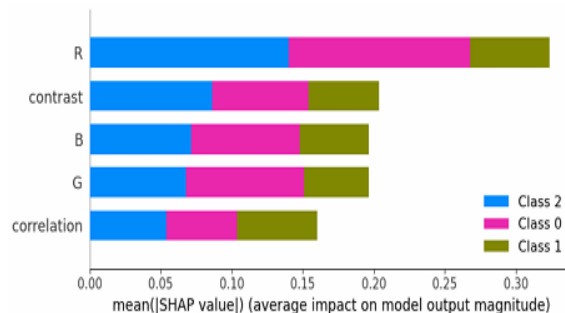
Gambar 2 memperlihatkan bahwa fitur warna R memiliki kontribusi rata-rata terbesar terhadap output model *random forest* tanpa seleksi fitur, diikuti oleh fitur G, fitur *contrast* dan fitur B. Kontribusi fitur warna (R, G, B) relatif dominan pada prediksi **Kelas 0 (Masak)**, sedangkan fitur tekstur seperti *correlation* dan *energy* memiliki kontribusi yang lebih besar pada **Kelas 2 (Mentah)**. Fitur *homogeneity* menunjukkan pengaruh relatif lebih tinggi terhadap **Kelas 1 (Mengkal)** dibandingkan fitur tekstur lainnya. Secara keseluruhan, model tampak mengandalkan informasi fitur warna untuk membedakan buah masak, sementara fitur tekstur memberikan kontribusi tambahan terutama untuk membedakan mengkal dan mentah.



Gambar 2. Analisis Shapley RF Tanpa Seleksi Fitur

3.3.2. RF Dengan Seleksi Fitur

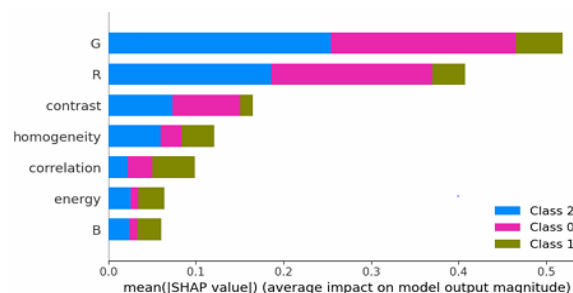
Gambar 3 menunjukkan hasil analisis *Shapley* pada model *random forest* setelah dilakukan seleksi fitur. Fitur **R** memiliki kontribusi paling besar terhadap prediksi ketiga kelas, diikuti oleh fitur **contrast**, **B**, **G** dan **correlation**. Distribusi nilai SHAP memperlihatkan bahwa fitur R secara konsisten memberikan pengaruh dominan pada pemisahan kelas 0 (Masak), kelas 1 (Mengkal), dan kelas 2 (Mentah). Hal ini mengindikasikan bahwa pemilihan subset fitur meningkatkan fokus model pada variabel yang paling relevan terhadap klasifikasi tingkat kematangan.



Gambar 3. Analisis Shapley RF Dengan Seleksi Fitur

3.3.3. SVM Tanpa Seleksi Fitur

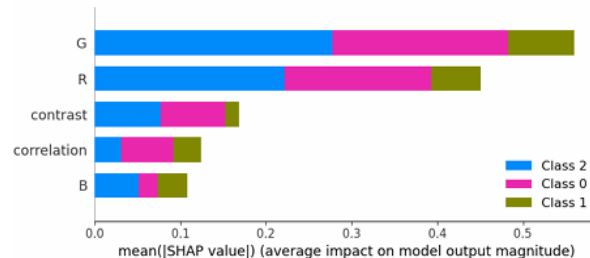
Gambar 4 menunjukkan hasil analisis *Shapley* pada model *support vector machine* tanpa seleksi fitur. Fitur **G** memiliki kontribusi terbesar terhadap prediksi ketiga kelas, diikuti oleh fitur **R**, yang menunjukkan bahwa variasi warna hijau dan merah lebih signifikan dalam menentukan kelas. **Contrast** juga memiliki pengaruh penting, terutama pada **Kelas 1 (Mengkal)**, sementara fitur-fitur tekstur seperti **correlation** dan **energy** memberikan kontribusi lebih kecil pada pemisahan kelas. Hasil ini menunjukkan bahwa meskipun seleksi fitur tidak diterapkan, warna dan tekstur tetap memainkan peran kunci dalam prediksi tingkat kematangan kelapa sawit.



Gambar 4. Analisis Shapley SVM Tanpa Seleksi Fitur

3.3.4. SVM Dengan Seleksi Fitur

Gambar 5 menunjukkan hasil analisis *Shapley* pada model *support vector machine* setelah seleksi fitur. Fitur **G** memiliki kontribusi terbesar dalam memprediksi ketiga kelas, diikuti oleh **R** yang memiliki kontribusi signifikan terutama pada **Kelas 2 (Mentah)**. Fitur **contrast** dan **correlation** berperan lebih kecil namun tetap memberikan dampak pada pemisahan kelas, dengan **contrast** lebih dominan untuk **Kelas 0 (Masak)**. Secara keseluruhan, fitur warna **G** dan **R** menunjukkan pengaruh terbesar pada pemisahan kelas-kelas tingkat kematangan kelapa sawit, dengan **contrast** dan **correlation** membantu dalam memisahkan kelas yang lebih rendah kematangannya dan memperkuat klasifikasi pada kelas-kelas tertentu.



Gambar 5. Analisis Shapley SVM Dengan Seleksi Fitur

Dominasi fitur warna, khususnya kanal R dan G, dapat dijelaskan oleh perubahan fisiologis dan biokimia yang terjadi selama proses pematangan buah kelapa sawit. Pada tahap awal, kandungan klorofil yang tinggi menghasilkan spektrum hijau yang kuat, sehingga nilai kanal G menjadi lebih dominan. Seiring dengan meningkatnya kematangan, degradasi klorofil terjadi bersamaan dengan peningkatan akumulasi *karotenoid*, terutama β -karoten, yang memperkuat intensitas kanal R. Pergeseran ini menyebabkan fitur warna lebih representatif dalam membedakan tingkat kematangan dibandingkan fitur tekstur. Dengan demikian, hasil analisis *shapley* yang menunjukkan dominasi kanal R dan G konsisten dengan mekanisme biologis yang mendasari perubahan warna buah kelapa sawit [23]–[25].

3.4. Keterbatasan Penelitian

Penelitian ini memiliki beberapa keterbatasan yang perlu diperhatikan dalam interpretasi hasil. Pertama, jumlah dataset yang digunakan relatif terbatas, yaitu hanya 254 citra buah kelapa sawit, sehingga kemampuan model dalam mengenali pola yang lebih beragam masih kurang optimal dan berpotensi menurunkan generalisasi pada kondisi nyata di lapangan. Kedua, citra yang digunakan bersumber dari repositori publik dengan karakteristik visual yang homogen, sehingga variasi kondisi pencahayaan, sudut pengambilan gambar, maupun kualitas kamera belum sepenuhnya terwakili. Hal ini dapat menimbulkan bias ketika model diterapkan pada data dengan kondisi berbeda. Ketiga, penelitian ini hanya memanfaatkan fitur warna berbasis kanal RGB dan fitur tekstur GLCM, sehingga aspek visual lain seperti bentuk, ukuran, atau spektrum warna yang lebih kompleks (misalnya HSV atau LAB) belum dieksplorasi lebih lanjut. Keempat, meskipun analisis *Shapley* memberikan interpretabilitas yang baik, penelitian ini belum menguji integrasi metode XAI lain yang mungkin memberikan perspektif tambahan dalam menjelaskan keputusan model. Keterbatasan ini menunjukkan bahwa masih diperlukan pengembangan lebih lanjut melalui pemanfaatan dataset yang lebih besar dan bervariasi, penambahan jenis fitur, serta eksplorasi teknik interpretabilitas yang lebih beragam agar model prediksi kematangan kelapa sawit dapat lebih andal dan aplikatif di lingkungan industri yang kompleks.

3.5 Tantangan Implementasi

Tantangan implementasi model di lapangan sangat berkaitan dengan kondisi pencahayaan dan kualitas kamera yang digunakan dalam proses pengambilan citra. Perubahan intensitas cahaya, bayangan, serta perbedaan resolusi kamera dapat memengaruhi distribusi nilai warna dan tekstur yang menjadi dasar ekstraksi fitur. Hal ini berpotensi menurunkan akurasi prediksi karena model telah dilatih pada data dengan kondisi visual yang relatif seragam. Selain itu, faktor lingkungan seperti posisi buah pada tandan atau keberadaan kotoran di permukaan kulit buah dapat menambah kompleksitas dalam proses klasifikasi. Oleh karena itu, diperlukan strategi penyesuaian, misalnya melalui teknik augmentasi data atau normalisasi citra, untuk meningkatkan ketahanan model terhadap variasi kondisi lapangan. Upaya ini penting agar model tidak hanya efektif, tetapi juga andal dalam penerapan di industri kelapa sawit.

3.6 Arah Penelitian Lanjutan

Penelitian ini masih dapat dikembangkan lebih lanjut melalui beberapa arah. Pertama, penerapan metode *deep learning*, khususnya *convolutional neural networks* (CNN), berpotensi meningkatkan akurasi klasifikasi

dengan mengekstraksi fitur secara otomatis tanpa perlu rekayasa fitur manual. Kedua, eksplorasi fitur spektral lain seperti ruang warna HSV, LAB, atau bahkan data citra hiperspektral dapat memperkaya representasi visual perubahan fisiologis buah selama proses pematangan. Ketiga, pengumpulan dataset yang lebih besar dan bervariasi, mencakup kondisi pencahayaan, sudut pengambilan gambar, dan kualitas kamera yang beragam, diperlukan untuk meningkatkan kemampuan generalisasi model di lapangan. Selain itu, integrasi metode *explainable AI* lain di luar *Shapley*, seperti LIME atau Grad-CAM pada model *deep learning*, juga dapat memberikan interpretabilitas yang lebih mendalam. Dengan demikian, penelitian lanjutan diharapkan tidak hanya mampu menghasilkan model dengan performa yang lebih tinggi, tetapi juga lebih adaptif dan dapat diandalkan dalam mendukung pengambilan keputusan panen kelapa sawit di lingkungan industri yang kompleks.

4. KESIMPULAN

Berdasarkan hasil eksperimen yang dilakukan, penelitian ini menunjukkan bahwa penerapan teknik seleksi fitur menggunakan *recursive feature elimination* (RFE) tidak memberikan peningkatan kinerja yang signifikan pada model *random forest* (RF) maupun *support vector machine* (SVM) dalam klasifikasi kematangan kelapa sawit. Meskipun demikian, SVM menunjukkan kinerja yang lebih stabil dan konsisten dibandingkan dengan RF, dengan nilai *akurasi*, *precision*, *recall* dan *F1-Score* yang lebih tinggi pada semua skenario. Hasil uji statistik *Mann-Whitney* juga mengonfirmasi bahwa tidak ada perbedaan signifikan antara model dengan dan tanpa seleksi fitur, baik pada RF maupun SVM. Namun, perbandingan antara kedua algoritma klasifikasi setelah seleksi fitur menunjukkan bahwa SVM secara signifikan lebih unggul dibandingkan RF, dengan peningkatan yang terlihat pada semua metrik evaluasi. Analisis *Shapley* lebih lanjut menunjukkan bahwa fitur warna khususnya fitur R dan G memiliki kontribusi terbesar terhadap prediksi tingkat kematangan kelapa sawit, meskipun peran fitur tekstur juga signifikan pada kelas tertentu. Hasil penelitian ini memberikan pemahaman bahwa meskipun seleksi fitur tidak selalu meningkatkan akurasi, pemilihan fitur yang tepat dapat membantu memperbaiki interpretasi model dan meningkatkan transparansi dalam proses pengambilan keputusan panen kelapa sawit.

UCAPAN TERIMA KASIH

Terima kasih kami sampaikan kepada semua pihak yang telah memberikan dukungan dalam penelitian ini. Terutama kepada pembimbing yang telah memberikan arahan dan bimbingan yang sangat berharga sepanjang proses penelitian. Terima kasih juga kepada pengelola dataset di *Mendeley Data* yang memungkinkan akses data yang digunakan dalam penelitian ini. Selain itu, kami mengapresiasi dukungan teknis yang diberikan oleh *platform Google Colaboratory* dan pustaka-pustaka yang digunakan dalam pengolahan data dan pelatihan model. Semoga penelitian ini dapat memberikan kontribusi positif dalam pengembangan aplikasi teknologi dalam sektor pertanian, khususnya dalam industri kelapa sawit. Demikian, kami mengucapkan terima kasih kepada keluarga dan teman-teman yang selalu memberikan dukungan baik materi maupun moral selama penelitian ini berlangsung.

DAFTAR PUSTAKA

- [1] F. Saqdiyah, H. Mulyati, and A. Setiawan Slamet, "Analisis Pemilihan Pemasok Kelapa Sawit yang Berkelanjutan dengan Menggunakan Metode PROMETHEE (Studi Kasus pada PT Perkebunan Nusantara III)," *J. Manaj. dan Organ.*, vol. 13, no. 2, pp. 124–133, Jun. 2022, doi: 10.29244/jmo.v13i2.37539.
- [2] N. Evitarina and K. Kusriani, "Metode Klasifikasi Kematangan Tandan Buah Segar Kelapa Sawit: Sebuah Tinjauan Sistematis," *G-Tech J. Teknol. Terap.*, vol. 8, no. 4, pp. 2324–2333, Oct. 2024, doi: 10.70609/gtech.v8i4.5050.
- [3] J. A. Sinambela, D. Cherie, A. Andasuryani, and M. Makky, "Prediksi Tingkat Kematangan Tandan Buah Segar (TBS) Kelapa Sawit Berbasis Sifat Optis," *J. Teknol. Pertan. Andalas*, vol. 29, no. 1, pp. 33–40, Mar. 2025, doi: 10.25077/jtpa.29.1.33-40.2025.
- [4] M. Afriansyah, J. Saputra, Y. Sa'adati, and Valian Yoga Pudya Ardhana, "Optimasi Algoritma Naïve Bayes Untuk Klasifikasi Buah Apel Berdasarkan Fitur Warna RGB," *Bull. Comput. Sci. Res.*, vol. 3, no. 3, pp. 242–249, Apr. 2023, doi: 10.47065/bulletincsr.v3i3.251.
- [5] H. P. Hadi and E. H. Rachmawanto, "Ekstraksi Fitur Warna dan GLCM Pada Algoritma KNN untuk Klasifikasi Kematangan Rambutan," *J. Inform. Polinema*, vol. 8, no. 3, pp. 63–68, Jun. 2022, doi: 10.33795/jip.v8i3.949.
- [6] N. Aini, M. Arif, I. T. Agustin, and Z. B. Toyibah, "Implementasi Algoritma Random Forest untuk Klasifikasi Bidang MSIB di Prodi Pendidikan Informatika," *J. Inform.*, vol. 11, no. 1, pp. 11–16, Apr. 2024, doi: 10.31294/inf.v11i1.20637.
- [7] N. Astrianda, "Klasifikasi Kematangan Buah Tomat Dengan Variasi Model Warna Menggunakan Support Vector Machine," *VOCATECH Vocat. Educ. Technol. J.*, vol. 1, no. 2, pp. 45–52, Apr. 2020, doi: 10.38038/vocatech.v1i2.27.
- [8] A. Adadi and M. Berrada, "Peeking Inside the Black-Box: A Survey on Explainable Artificial Intelligence (XAI)," *IEEE Access*, vol. 6, pp. 52138–52160, Sep. 2018, doi: 10.1109/ACCESS.2018.2870052.
- [9] T. Spinner, U. Schlegel, H. Schafer, and M. El-Assady, "explAiner: A Visual Analytics Framework for Interactive

- p and Explainable Machine Learning,”
- IEEE Trans. Vis. Comput. Graph.*
- , p. 1, Aug. 2019, doi: 10.1109/tvcg.2019.2934629.
- [10] A. Y. Al Hammadi *et al.*, “Explainable artificial intelligence to evaluate industrial internal security using EEG signals in IoT framework,” *Ad Hoc Networks*, vol. 123, no. August, p. 102641, Dec. 2021, doi: 10.1016/j.adhoc.2021.102641.
 - [11] D. V. Carvalho, E. M. Pereira, and J. S. Cardoso, “Machine learning interpretability: A survey on methods and metrics,” *Electron.*, vol. 8, no. 8, pp. 1–34, 2019, doi: 10.3390/electronics8080832.
 - [12] S. Coulibaly, B. Kamsu-Foguem, D. Kamissoko, and D. Traore, “Deep learning for precision agriculture: A bibliometric analysis,” *Intell. Syst. with Appl.*, vol. 16, no. April, p. 200102, 2022, doi: 10.1016/j.iswa.2022.200102.
 - [13] A. R. I. Pratama, S. A. Latipah, and B. N. Sari, “Optimasi Klasifikasi Curah Hujan Menggunakan Support Vector Machine (SVM) Dan Recursive Feature Elimination (RFE),” *JUPI (Jurnal Ilm. Penelit. dan Pembelajaran Inform.*, vol. 7, no. 2, pp. 314–324, May 2022, doi: 10.29100/jupi.v7i2.2675.
 - [14] N. M. Putri, M. Praseptiawan, and M. C. Untoro, “Analisis Model Sistem Rekomendasi Kursus Mooc Dengan Metode Collaborative Filtering Dan Integrasi Explainable Ai,” *Indones. J. Bus. Intell.*, vol. 7, no. 1, pp. 26–36, 2024.
 - [15] M. A. Aprihartha and I. Idham, “Optimization of Classification Algorithms Performance with k-Fold Cross Validation,” *Eig. Math. J.*, vol. 7, no. 2, pp. 61–66, Sep. 2024, doi: 10.29303/emj.v7i2.212.
 - [16] H. Hajaroh, T. Suprpti, and R. Narasati, “Implementasi Algoritma Naive Bayes untuk Analisis Sentimen Ulasan Produk Makanan dan Minuman di Tokopedia,” *JATI (Jurnal Mhs. Tek. Inform.*, vol. 8, no. 1, pp. 111–118, Feb. 2024, doi: 10.36040/jati.v8i1.8237.
 - [17] N. Nurzaman, N. Suarna, and W. Prihartono, “Analisis Sentimen Ulasan Aplikasi Threads di Google Playstore Menggunakan Algoritma Naïve Bayes,” *JATI (Jurnal Mhs. Tek. Inform.*, vol. 8, no. 1, pp. 967–974, Mar. 2024, doi: 10.36040/jati.v8i1.8708.
 - [18] M. A. Saddam, E. Kurniawan D, and I. Indra, “Analisis Sentimen Fenomena PHK Massal Menggunakan Naive Bayes dan Support Vector Machine,” *J. Inform. J. Pengemb. IT*, vol. 8, no. 3, pp. 226–233, Sep. 2023, doi: 10.30591/jpit.v8i3.4884.
 - [19] N. Cahyono and Anggista Oktavia Praneswara, “Analisis Sentimen Ulasan Aplikasi TikTok Shop Seller Center di Google Playstore Menggunakan Algoritma Naive Bayes,” *Indones. J. Comput. Sci.*, vol. 12, no. 6, Dec. 2023, doi: 10.33022/ijcs.v12i6.3473.
 - [20] I. Amal and J. Jayanta, “Perbandingan Pelabelan Otomatis Dan Manual Untuk Analisis Sentimen Terhadap Kenaikan Harga BBM Pertamina Pada Twitter Menggunakan Algoritma Support Vector Machine,” in *Prosiding Seminar Nasional Mahasiswa Bidang Ilmu Komputer dan Aplikasinya*, 2023, vol. 4, no. 2, pp. 473–487.
 - [21] M. Susilawati, D. Selpia, M. Fathurrahman, N. Pratiwi, and R. Purnami, “Penerapan uji mann-whitney dalam perbandingan prestasi akademik mahasiswa statistika Universitas Hamzanwadi angkatan 2022 dan 2023,” *J. Eksbar*, vol. 1, no. 2, pp. 19–28, 2024.
 - [22] H. Gani, “An Explainable Machine Learning Model to Explain the Influential Climate Parameters Based on Rainfall Prediction,” *J. IT*, vol. 15, no. 2, pp. 98–110, 2024.
 - [23] J. J. J. Martin *et al.*, “Lipidomic Profiles of Lipid Biosynthesis in Oil Palm during Fruit Development,” *Metabolites*, vol. 13, no. 6, p. 727, Jun. 2023, doi: 10.3390/metabo13060727.
 - [24] H. F. Teh *et al.*, “Hormones, Polyamines, and Cell Wall Metabolism during Oil Palm Fruit Mesocarp Development and Ripening,” *J. Agric. Food Chem.*, vol. 62, no. 32, pp. 8143–8152, Aug. 2014, doi: 10.1021/jf500975h.
 - [25] T. J. Tranbarger *et al.*, “Regulatory Mechanisms Underlying Oil Palm Fruit Mesocarp Maturation, Ripening, and Functional Specialization in Lipid and Carotenoid Metabolism,” *Plant Physiol.*, vol. 156, no. 2, pp. 564–584, Jun. 2011, doi: 10.1104/pp.111.175141.

