

Klasifikasi Kualitas Daun Teh Menggunakan Metode Support Vector Machine

Bayu Fadjar Dwi Puta^{a1}, I Ketut Gede Suhartana^{a2}, Putu Praba Santika^{a3}

^aProgram Studi Informatika, Fakultas Matematika dan Ilmu Pengetahuan Alam,
Universitas Udayana, Bali
Jln. Raya Kampus UNUD, Bukit Jimbaran, Kuta Selatan, Badung, Bali, Indonesia
¹bayufadjar.business@gmail.com
²ikg.suhartana@unud.ac.id
³praba@unud.ac.id

Abstract

Tea leaf quality serves as the fundamental determinant of both sensory characteristics and commercial competitiveness in the global tea market. However, manual assessment of tea leaf quality remains limited by observer subjectivity and inconsistent classification. This study aims to develop an automatic tea leaf quality classification system based on leaf maturity using a digital image processing approach. The method employed is Support Vector Machine (SVM) with a combination of three feature extraction techniques: color histogram for color features, Gabor filter for texture features, and Histogram of Oriented Gradients (HOG) for shape features. The dataset consists of 4,272 tea leaf images classified into four quality classes: Premium Grade, Standard Grade, Basic Grade, and Reject Grade. Principal Component Analysis (PCA) was applied for dimensionality reduction while maintaining 95% data variance. Testing results show an accuracy of 84.53% with an F1-score of 84.56%, demonstrating the effectiveness of the system in automatically classifying tea leaf quality.

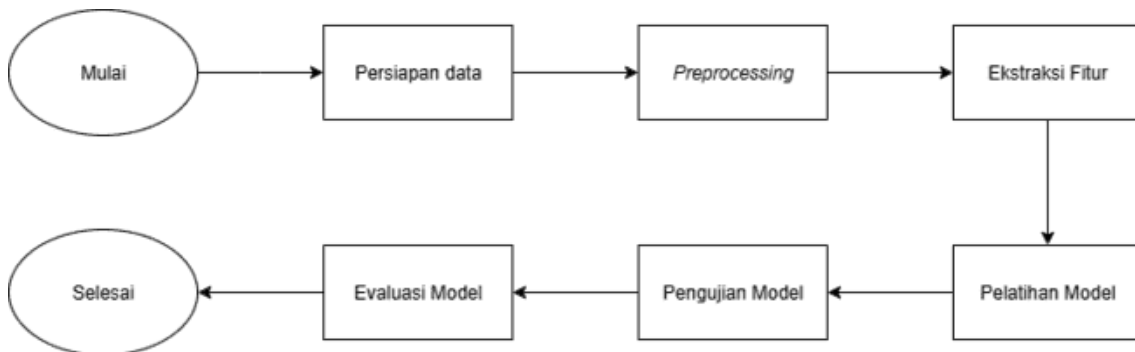
Keywords: Support Vector Machine, Tea Leaf Quality Classification, Gabor filter, Histogram of Oriented Gradients, Color Histogram, Machine Learning

1. Pendahuluan

Teh merupakan komoditas agrikultur yang berasal dari Tiongkok dan saat ini menempati posisi kedua sebagai minuman yang paling banyak dikonsumsi di dunia setelah air, dengan estimasi konsumsi mencapai 3 miliar cangkir per hari [1]. Dalam menghadapi dinamika persaingan di pasar global, diperlukan strategi untuk membangun keunggulan kompetitif, salah satunya melalui upaya menjaga dan meningkatkan kualitas produk secara konsisten. Konsistensi cita rasa teh merupakan aspek penting yang berkorespondensi dengan kualitas bahan baku, sehingga pemilihan daun teh terbaik menjadi aspek krusial dalam proses produksinya. Penilaian tingkat kematangan daun teh secara manual oleh manusia memiliki sejumlah keterbatasan yang signifikan. Perbedaan persepsi subjektif antar individu seperti tingkat pengalaman, dan kesehatan penglihatan yang berbeda dapat menyebabkan inkonsistensi dalam proses pemilihan daun teh. Selain itu, metode manual memerlukan waktu yang lama dan tenaga yang cukup besar, terutama ketika harus menangani volume daun teh dalam jumlah besar. Oleh karena itu, dibutuhkan suatu sistem otomatis yang mampu mengklasifikasikan tingkat kualitas daun teh secara objektif dan akurat untuk meningkatkan efektivitas dan efisiensi proses produksi. Sistem semacam ini dapat memberikan manfaat nyata bagi para petani dalam memilih daun teh yang masih layak digunakan dalam proses produksi dari yang sudah tidak memenuhi standar kualitas. Sebelumnya telah ada beberapa penelitian yang meneliti terkait klasifikasi kualitas daun teh dengan menggunakan dataset dan metode yang berbeda [2][3][4]. Penelitian ini memberikan kontribusi baru dengan memanfaatkan dataset citra yang dikembangkan berdasarkan klasifikasi usia daun teh. Metode klasifikasi yang diusulkan pada penelitian ini menggunakan Support Vector Machine (SVM), dengan menggabungkan beberapa teknik ekstraksi fitur, yaitu Gabor Filter, Histogram of Oriented Gradients (HOG), dan Color Histogram. Penelitian ini bertujuan untuk

mengembangkan sebuah sistem klasifikasi otomatis tingkat kualitas daun teh berdasarkan usia daun menggunakan metode pengolahan citra digital dengan tingkat akurasi yang baik.

2. Metode Penelitian



Gambar 1. Alur Penelitian

Pada Gambar 1 digambarkan alur penelitian yang dilakukan melalui serangkaian tahapan sistematis yang dimulai dengan persiapan data, yang mencakup proses pengumpulan data citra daun teh serta redistribusi data ke dalam subset pelatihan, validasi, dan pengujian guna memastikan proporsi yang seimbang dan representatif. Setelah data terkumpul, dilakukan tahap *preprocessing* yang terdiri dari konversi warna gambar dari BGR ke RGB, ekstraksi Region of Interest (ROI) berdasarkan informasi bounding box dari label YOLO, kemudian menormalisasi ukuran ROI dengan mengubah ukuran resolusinya menjadi 64×64 piksel. Data citra hasil *preprocessing* digunakan dalam ekstraksi fitur yang menerapkan tiga teknik utama, yaitu histogram warna untuk merepresentasikan informasi warna, Histogram of Oriented Gradients (HOG) untuk mendeskripsikan bentuk, serta Gabor Filter untuk menangkap karakteristik tekstur. Fitur-fitur yang telah diekstraksi kemudian digunakan dalam tahap pelatihan model menggunakan algoritma SVM. Model yang telah dilatih selanjutnya diuji untuk menghasilkan prediksi terhadap data uji, dan kinerjanya dievaluasi menggunakan matriks evaluasi seperti akurasi, presisi, recall, dan F1-score guna menilai efektivitas sistem klasifikasi yang dikembangkan.

2.1. Persiapan Data

Dataset yang digunakan dalam penelitian ini merupakan dataset sekunder yang diunduh dari repositori Mendeley Data dan dikembangkan oleh Kabir et al. [5]. Dataset tersebut dipublikasikan melalui artikel “Tea leaf dataset for classification of different tea qualities based on tea grading standards” [6], dan disusun berdasarkan standar pengelompokan kualitas daun teh yang mengacu pada usia serta karakteristik morfologis daun. Dataset ini terbagi ke dalam tiga *folder* utama, yaitu raw, annotated, dan annotated and augmented. *Folder* raw berisi citra daun teh mentah yang belum diberi label, sementara annotated berisi citra yang telah diberikan label menggunakan format YOLO. *Folder* annotated and augmented, yang digunakan dalam penelitian ini, mencakup citra yang telah diberi label dan diperluas melalui augmentasi berupa dua rotasi acak untuk setiap gambar, sehingga menghasilkan total 5.961 citra beranotasi.

Meskipun *folder* annotated and augmented telah terbagi secara bawaan ke dalam tiga subset (*training*, *validation*, dan *testing*), distribusi jumlah citra antar kelas masih belum merata. Oleh karena itu, dilakukan proses redistribusi data secara acak untuk menyamakan jumlah data per kelas. Proses redistribusi ini mengacu pada jumlah data dari kelas dengan frekuensi terkecil, yaitu sebanyak 1.068 citra per kelas, dan dibagi dengan rasio 70% untuk data pelatihan, 10% untuk data validasi, dan 20% untuk data pengujian. Empat kelas kualitas yang digunakan dalam penelitian ini adalah Premium Grade, Standard Grade, Basic Grade, dan Reject Grade, sesuai dengan klasifikasi yang ditetapkan dalam dataset asli.

2.2. Preprocessing

Tahap *preprocessing* dilakukan untuk menyiapkan citra daun teh sebelum masuk ke proses ekstraksi fitur. Proses ini dimulai dengan penyesuaian *color space* citra agar sesuai dengan standar representasi warna digital yang umum digunakan yaitu RGB. Selanjutnya, dilakukan ekstraksi ROI, yaitu bagian gambar yang merepresentasikan objek utama berupa daun teh, berdasarkan koordinat anotasi yang telah tersedia. Tahap ini bertujuan untuk menghilangkan bagian latar belakang yang tidak relevan agar sistem fokus pada area objek.

Setelah ROI diperoleh, seluruh citra diperkecil atau diperbesar ke dalam ukuran standar yang seragam menjadi 64×64 piksel. Hasil dari tahap *preprocessing* ini adalah citra daun teh dengan latar belakang yang diminimalkan, fokus pada objek utama, dan dalam ukuran serta format warna yang seragam siap untuk dilakukan proses ekstraksi fitur.

2.3. Ekstraksi Fitur

Proses ekstraksi fitur pada penelitian ini bertujuan untuk menangkap informasi karakteristik warna, tekstur, dan bentuk dari setiap jenis kualitas daun teh. Ekstraksi fitur warna dilakukan menggunakan teknik color histogram yang diterapkan pada *color space* RGB, HSV, dan LAB. Pendekatan ekstraksi multi *color space* dipilih untuk memperbanyak informasi warna sehingga meningkatkan sensitivitas sistem untuk mengenali variasi warna setiap jenis kualitas daun teh yang berbeda. Teknik dilakukan dengan menghitung histogram setiap *channel* warna yang kemudian dinormalisasi agar menghasilkan vektor fitur yang stabil.

Selanjutnya, fitur tekstur diekstraksi menggunakan Gabor filter, yang dikenal efektif dalam mendeskripsikan pola tekstur lokal pada citra. Filter ini diterapkan pada setiap channel dari ketiga *color space*, yaitu RGB, HSV, dan LAB di berbagai orientasi dan skala, lalu dari hasil filter tersebut diekstraksi statistik deskriptif seperti nilai rata-rata, simpangan baku, nilai maksimum, dan minimum. Pendekatan ini memungkinkan sistem untuk mengenali variasi halus dalam tekstur permukaan daun, yang berkaitan erat dengan tingkat kematangan dan kualitas daun teh.

Pengambilan informasi fitur bentuk daun teh dipilih metode Histogram of Oriented Gradients (HOG) yang diterapkan pada citra grayscale. Metode ini memetakan arah dominan dari gradasi piksel dan membaginya ke dalam blok-blok untuk menangkap informasi struktur tepi serta kontur objek secara lokal. Informasi bentuk dari HOG membantu sistem untuk membedakan morfologi daun teh yang setiap kelas yang berbeda.

Ketiga jenis fitur tersebut kemudian digabungkan menjadi satu vektor fitur yang merepresentasikan keseluruhan informasi citra daun teh. Vektor gabungan tersebut akan memiliki dimensi yang besar dan berkemungkinan besar memiliki reduksi dimensi yang tidak diperlukan, sehingga dilakukan opsi reduksi dimensi menggunakan Principal Component Analysis (PCA) dengan tingkat variansi kumulatif standar sebesar 95% [7].

2.4. Support Vector Machine

Metode klasifikasi yang digunakan dalam penelitian ini adalah SVM, yang merupakan salah satu metode *supervised learning* yang dapat menangani masalah klasifikasi dan regresi dengan data berdimensi tinggi.

SVM bekerja dengan mencari hyperplane pemisah optimal untuk memaksimalkan margin antara kelas-kelas yang berbeda dalam ruang fitur. SVM memiliki fungsi kernel untuk memetakan data ke ruang berdimensi lebih tinggi yang dapat digunakan untuk memudahkan pemisahan ketika data tidak dapat dipisahkan secara linear [8].

Dalam implementasi ini, digunakan metode Grid Search untuk mengoptimalkan parameter-parameter utama pada model SVM. Parameter yang disesuaikan meliputi parameter C (parameter regularisasi), gamma (koefisien fungsi kernel), dan kernel (linear, polynomial, RBF). Melalui pengoptimalan parameter, model SVM yang dihasilkan diharapkan mampu

mengklasifikasikan semua kelas dengan optimal dan tidak terjadi bias pada salah satu kelas. Model ini akan dilatih menggunakan hasil ekstraksi fitur dari data training yang telah digabungkan dalam satu vektor fitur sebagai representasi akhir dari setiap citra.

2.5. Pengujian dan Evaluasi

Setelah pelatihan, model SVM digunakan untuk melakukan prediksi pada data uji. Untuk menilai kinerja model secara menyeluruh, dilakukan evaluasi menggunakan beberapa matriks, yaitu precision, recall, dan f1-score, yang dihitung untuk setiap kelas. Selain itu, digunakan juga confusion matrix untuk memvisualisasikan distribusi prediksi model terhadap label, serta mengidentifikasi pola kesalahan klasifikasi antar kelas. Evaluasi ini bertujuan untuk memastikan bahwa model tidak hanya memiliki akurasi tinggi, tetapi juga konsisten dalam mengenali keseluruhan kelas.

3. Hasil dan Diskusi

3.1. Deskripsi Data



Gambar 2. Contoh Dataset Citra

Dataset yang digunakan dalam penelitian ini terdiri dari citra daun teh yang telah dikelompokkan ke dalam empat kelas kualitas, yaitu Premium Grade, Standard Grade, Basic Grade, dan Reject Grade yang dapat dilihat sampel citra untuk masing-masing kelas pada Gambar 2. Setiap citra disimpan dalam format JPEG (.jpg) dengan resolusi tetap sebesar 64×64 piksel. Pada setiap sampel citra pada Gambar 2, terdapat kotak merah yang merepresentasikan bounding box atau wilayah objek utama daun teh yang diamati. Kotak tersebut diperoleh dari informasi label yang disimpan dalam file .txt berformat YOLO, yang menyertakan koordinat posisi dan ukuran objek dalam citra. bounding box ini digunakan untuk mengekstraksi ROI dalam proses preprocessing, sehingga fitur yang diambil fokus pada area daun teh yang ingin diamati.

Dilakukan proses redistribusi data terhadap seluruh dataset untuk memastikan pemerataan data antar kelas. Setiap kelas ditetapkan memiliki 1.068 citra, yang kemudian dibagi ke dalam tiga subset dengan rasio 70% untuk data pelatihan, 10% untuk data validasi, dan 20% untuk data pengujian. Dengan demikian, total dataset citra daun teh yang digunakan berjumlah 4.272 citra. Proses redistribusi ini bertujuan untuk menjaga keseimbangan antar kelas pada setiap subset data sehingga model yang dibangun dapat melakukan pembelajaran secara adil dan tidak bias terhadap kelas tertentu. Hasil rincian dari pendistribusian ulang dataset daun teh yang digunakan telah disajikan pada Tabel 1.

Tabel 1. Distribusi Data

Pembagian Kelas	Test	Train	Valid
Premium Grade	215	747	106
Standard Grade	215	747	106
Basic Grade	215	747	106
Reject Grade	215	747	106

3.2. Hasil Preprocessing dan Ekstraksi ROI



Gambar 3. Data *Preprocessing*

Seluruh data citra yang digunakan, baik untuk pelatihan, validasi, maupun pengujian, terlebih dahulu melalui tahap *preprocessing* sebelum digunakan untuk ekstraksi fitur. Proses ini meliputi konversi warna citra ke format RGB dan ekstraksi ROI berdasarkan koordinat pada *file* labels. Setiap ROI yang diperoleh kemudian disesuaikan ukurannya menjadi 64×64 piksel untuk memastikan keseragaman dimensi ukuran citra. Langkah ini bertujuan agar hanya bagian citra yang relevan (daun teh) yang digunakan untuk proses selanjutnya, sekaligus menjaga konsistensi data di seluruh tahapan pemrosesan. Hasil dari proses *preprocessing* ini dapat dilihat pada Gambar 3, yang menampilkan contoh citra daun teh yang telah dipotong dan disesuaikan ukurannya.

3.3. Hasil Ekstraksi Fitur

Citra yang telah dilakukan *preprocessing* kemudian digunakan untuk ekstraksi fitur untuk mengubah representasi visual citra menjadi vektor numerik yang dapat digunakan oleh algoritma klasifikasi. Proses ekstraksi ini dirancang untuk menangkap tiga karakteristik utama dari citra daun teh, yaitu warna, tekstur, dan bentuk, melalui 3 teknik ekstraksi fitur.

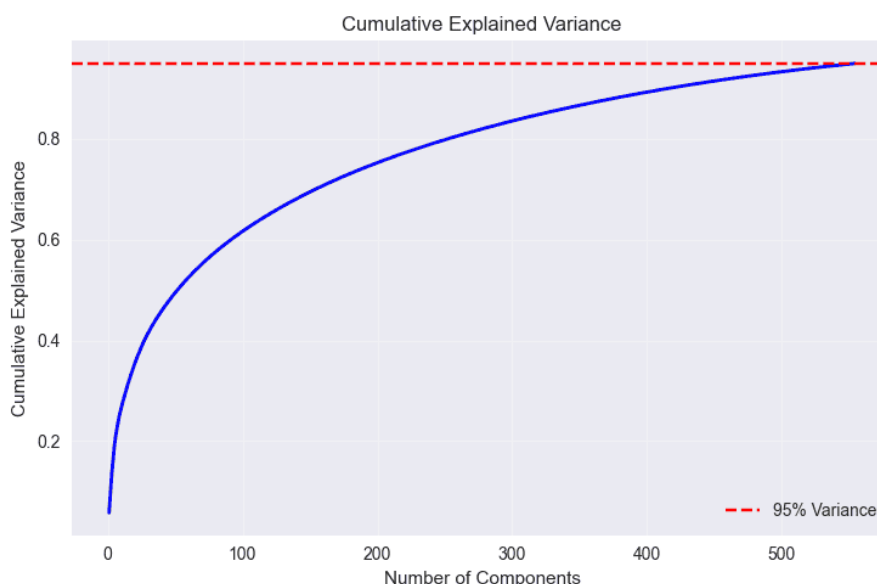
Pertama, fitur warna diekstraksi menggunakan teknik color histogram yang diterapkan pada tiga ruang warna berbeda, yaitu RGB, HSV, dan LAB. Setiap channel warna dari ketiga ruang tersebut dihitung histogramnya dan dinormalisasi. Fitur yang dihasilkan dari teknik ini adalah 288 fitur. Selanjutnya, fitur tekstur diperoleh menggunakan Gabor filter, yang diterapkan pada setiap channel dari RGB, HSV, LAB dalam berbagai orientasi dan skala. Total fitur yang dihasilkan dari Gabor filter mencapai 432 fitur, yang secara khusus merepresentasikan pola tekstur lokal permukaan daun teh. Digunakan metode HOG pada citra grayscale Untuk menangkap informasi

bentuk. Fitur bentuk yang dihasilkan dari HOG mencapai 1.764 fitur. Rincian dari dimensi fitur yang dihasilkan telah disajikan pada Tabel 2.

Tabel 2. Hasil Ekstraksi Fitur

Dimensi Fitur	Jumlah
Color Histogram	288
Gabor Filter	432
HOG	1764
Total	2484

Secara keseluruhan, jumlah fitur gabungan yang dihasilkan dari ketiga teknik ini adalah 2.484 fitur. Untuk mengurangi kompleksitas data dan menghindari redudansi informasi, dilakukan proses reduksi dimensi menggunakan Principal Component Analysis (PCA) dengan tingkat variansi kumulatif yang dipertahankan sebesar 95%. Hasil dari proses PCA ini mereduksi 2.484 fitur citra daun teh menjadi 554 fitur akhir yang masih mempertahankan 95% informasi dari jumlah fitur asli. Visualisasi hasil reduksi fitur ini ditampilkan pada Gambar 4, yang menunjukkan jumlah komponen utama yang diperlukan untuk mencapai tingkat informasi tersebut.



Gambar 4. Grafik Reduksi Fitur dengan PCA

3.4. Pelatihan Model SVM

Setelah tahap ekstraksi fitur selesai, proses pelatihan model klasifikasi dilakukan menggunakan algoritma SVM. Data yang digunakan dalam pelatihan merupakan hasil ekstraksi fitur dari subset data pelatihan, yang telah direduksi dimensinya menggunakan PCA. Model SVM dilatih untuk mengklasifikasikan citra daun teh ke dalam empat kelas kualitas berdasarkan representasi numerik dari fitur warna, tekstur, dan bentuk. Untuk memperoleh konfigurasi parameter terbaik pada model SVM, digunakan teknik optimasi Grid Search yang dikombinasikan dengan 5-Fold Stratified Cross-Validation. Proses tuning ini dirancang untuk mengevaluasi total 76 kombinasi parameter berdasarkan tiga komponen utama, yaitu C, gamma, kernel. Pemilihan parameter terbaik dilakukan berdasarkan metrik F1-score makro, yang memperhitungkan performa model secara seimbang pada semua kelas. Hasil dari Grid Search menunjukkan bahwa kombinasi parameter terbaik adalah C dengan nilai 10, gamma scale, kernel RBF.

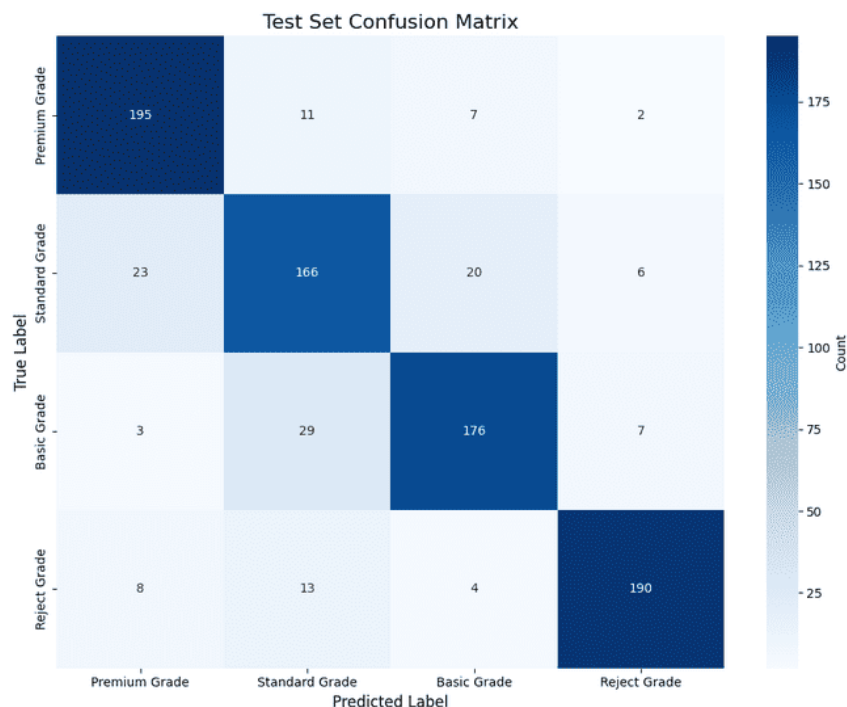
Dengan parameter tersebut, model menghasilkan best cross-validation score terbaik sebesar 0.6645, dan best cross-validation accuracy sebesar 0.6645. Setelah konfigurasi terbaik diperoleh, model SVM dilatih ulang menggunakan seluruh data pelatihan dan kemudian dievaluasi terhadap subset data validasi. Hasil evaluasi pada data validasi menunjukkan peningkatan performa yang signifikan, yaitu akurasi validasi sebesar 0.8396 dan f1-score sebesar 0.8361. Perbedaan skor antara hasil cross-validation dan evaluasi pada data validasi mengindikasikan bahwa model berhasil mempelajari pola dari data pelatihan dengan baik dan mampu melakukan generalisasi yang baik pada data yang belum pernah dilihat sebelumnya.

3.5. Hasil Prediksi dan Evaluasi Model

Detailed Classification Report:				
	precision	recall	f1-score	support
Premium Grade	0.8515	0.9070	0.8784	215
Standard Grade	0.7580	0.7721	0.7650	215
Basic Grade	0.8502	0.8186	0.8341	215
Reject Grade	0.9268	0.8837	0.9048	215
accuracy			0.8453	860
macro avg	0.8466	0.8453	0.8456	860
weighted avg	0.8466	0.8453	0.8456	860

Gambar 5. Hasil Pengujian

Berdasarkan hasil pengujian pada data uji pada Gambar 5, model SVM yang telah dilatih menunjukkan performa yang baik dalam mengklasifikasikan keempat kelas kualitas daun teh. Hasil pengujian f1-score rata-rata (macro average) sebesar 84.56% dan akurasi keseluruhan sebesar 84.53%. Namun, kelas Standar Grade memiliki nilai prediksi f1-score yang jauh lebih rendah dibandingkan dengan kelas lainnya, yaitu 76.5%. Rendahnya f1-score pada kelas Standard Grade disebabkan oleh nilai precision dan recall yang relatif lebih rendah, masing-masing sebesar 75.8% dan 77.21%. Hal ini menunjukkan bahwa model cenderung salah mengklasifikasikan citra Standard Grade ke kelas lain, dan juga terkadang salah memasukkan citra kelas lain ke dalam Standard Grade.



Gambar 6. Confusion Matrix

Gambar 6 merupakan visualisasi confusion matrix dari hasil pengujian model. Berdasarkan Gambar 6, hanya 166 citra Standard Grade yang diklasifikasikan dengan benar dari total 215 citra. 23 citra salah diklasifikasikan sebagai Premium Grade, 20 citra salah diklasifikasikan sebagai Basic Grade, dan 6 citra salah diklasifikasikan ke Reject Grade. Selain itu, sebagian besar kesalahan klasifikasi dari kelas lain juga cenderung terprediksi sebagai kelas Standard Grade. Distribusi kesalahan ini menunjukkan bahwa model mengalami kesulitan dalam membedakan Standard Grade dari kelas-kelas lain, terutama Premium dan Basic Grade, yang secara visual mungkin memiliki kemiripan dalam bentuk, warna, atau tekstur daun.

4. Kesimpulan

Penelitian ini bertujuan untuk merancang dan mengembangkan sistem klasifikasi otomatis tingkat kualitas daun teh berdasarkan usia daun dengan pendekatan pengolahan citra digital menggunakan algoritma Support Vector Machine (SVM). Model dibangun dengan menggabungkan tiga teknik ekstraksi fitur utama, yaitu color histogram untuk fitur warna, Gabor filter untuk fitur tekstur, dan Histogram of Oriented Gradients (HOG) untuk fitur bentuk. Hasil pengujian menunjukkan bahwa model yang diusulkan mampu mencapai F1-score sebesar 84,53%, yang mengindikasikan performa klasifikasi yang cukup baik dalam mengenali keempat kelas kualitas daun teh. Meskipun demikian, model masih memiliki keterbatasan dalam membedakan citra dengan kelas Standard Grade, yang memiliki kemiripan visual dengan kelas lainnya. Oleh karena itu, penelitian ini membuka peluang untuk pengembangan lebih lanjut, seperti penerapan metode pra-pemrosesan tambahan seperti penghapusan latar belakang, atau eksplorasi menggunakan algoritma klasifikasi lainnya untuk meningkatkan akurasi dan generalisasi model.

Daftar Pustaka

- [1] H. Nursodik, S. Santoso, and S. Nurfadillah, "Competitiveness and Determining Factors of Indonesian Tea Export Volume in the World Market," *HABITAT*, vol. 32, no. 3, pp. 163–172, Dec. 2021, doi: 10.21776/ub.habitat.2021.032.3.18.
- [2] A. Bakhshipour, H. Zareiforoush, and I. Bagheri, "Application of decision trees and fuzzy inference system for quality classification and modeling of black and green tea based on

- visual features,” *Journal of Food Measurement and Characterization*, vol. 14, no. 3, pp. 1402–1416, Feb. 2020, doi: 10.1007/s11694-020-00390-8.
- [3] M. Xu, J. Wang, and L. Zhu, “Tea quality evaluation by applying E-nose combined with chemometrics methods,” *Journal of Food Science and Technology*, vol. 58, no. 4, pp. 1549–1561, Aug. 2020, doi: 10.1007/s13197-020-04667-0.
- [4] X. Zhao, Y. He, H. Zhang, Z. Ding, C. Zhou, and K. Zhang, “A quality grade classification method for fresh tea leaves based on an improved YOLOv8x-SPPCSPC-CBAM model,” *Scientific Reports*, vol. 14, no. 1, Feb. 2024, doi: 10.1038/s41598-024-54389-y.
- [5] M. M. Kabir, “TeaLeafAgeQuality: Age-Stratified Tea Leaf Quality Classification Dataset,” Mendeley Data. [Online]. Available: <https://data.mendeley.com/datasets/7t964jmmy3/1>
- [6] M. M. Kabir, M. S. Hafiz, S. Bandyopadhyay, J. R. Jim, and M. F. Mridha, “Tea leaf age quality: Age-stratified tea leaf quality classification dataset,” *Data in Brief*, vol. 54, p. 110462, Jun. 2024, doi: 10.1016/j.dib.2024.110462.
- [7] M. Singh, S. Bansal, S. Ahuja, R. K. Dubey, B. K. Panigrahi, and N. Dey, “Transfer learning-based ensemble support vector machine model for automated COVID-19 detection using lung computerized tomography scan data,” *Medical & Biological Engineering & Computing*, vol. 59, no. 4, pp. 825–839, Mar. 2021, doi: 10.1007/s11517-020-02299-2.
- [8] M. A. Chandra and S. S. Bedi, “Survey on SVM and their application in image classification,” *International Journal of Information Technology*, vol. 13, no. 5, pp. 1–11, Jan. 2018, doi: 10.1007/s41870-017-0080-1.

