

Klasifikasi Jenis Sampah Menggunakan Metode Random Forest

Diana Sava Salsabila¹, Arimbi Puspitasari², Dwi Rolliawati³

¹ Sistem Informasi, Universitas Islam Negeri Sunan Ampel, Surabaya

² Sistem Informasi, Universitas Islam Negeri Sunan Ampel, Surabaya

³ Sistem Informasi, Universitas Islam Negeri Sunan Ampel, Surabaya

¹dianasavasalsabila@gmail.com, ²arimbipuspitasari69068@gmail.com, ³dwi_roll@uinsa.ac.id

INTISARI

Pemilahan sampah secara manual kerap menghadapi kendala, terutama di wilayah dengan infrastruktur dan kesadaran rendah. Penelitian ini mengembangkan model klasifikasi jenis sampah menggunakan algoritma Random Forest yang diimplementasikan dengan framework PySpark. Dataset yang digunakan adalah Real Waste dengan 4.752 gambar dari berbagai kategori sampah. Setelah preprocessing dan pembagian data, model menghasilkan akurasi sebesar 44,64%, dengan precision 48,62% dan recall 44,64%, menunjukkan performa yang masih perlu ditingkatkan. Kategori plastik dan vegetasi memiliki akurasi terbaik, sedangkan kategori makanan organik dan tekstil mengalami kesulitan akibat kemiripan fitur visual. Tujuan utama penelitian ini adalah memberikan solusi praktis yang mendukung pengelolaan sampah berbasis teknologi. Studi ini memberikan dasar awal bagi pemanfaatan teknologi berbasis pembelajaran mesin untuk mendukung pengelolaan sampah yang lebih efektif dan berkelanjutan.

Kata kunci: Klasifikasi Sampah, Random Forest, PySpark, Pembelajaran Mesin

ABSTRACT

Manual waste segregation often faces challenges, especially in areas with limited infrastructure and low public awareness. This research develops a waste classification model using the Random Forest algorithm, implemented with the PySpark framework. The dataset used is Real Waste, consisting of 4,752 images across various waste categories. Following preprocessing and data splitting, the model achieved an accuracy of 44.64%, with a precision of 48.62% and a recall of 44.64%, demonstrating performance that requires further enhancement. The categories of plastic and vegetation achieved the best accuracy, while organic food and textile categories faced difficulties due to similarities in visual features. The main goal of this study is to provide a practical solution that supports technology-based waste management. This research establishes an initial foundation for utilizing machine learning technologies to promote more effective and sustainable waste management systems.

Keywords: Waste Classification, Random Forest, PySpark, Machine Learning

1. PENDAHULUAN

Sampah merupakan salah satu tantangan utama yang dihadapi masyarakat modern seiring dengan meningkatnya populasi, aktivitas ekonomi, dan perkembangan industri. Peningkatan jumlah penduduk menyebabkan produksi sampah semakin meningkat dari tahun ke tahun (Danang Aji Kurniawan & Ahmad Zaenal Santoso, 2021). Berdasarkan data dari Sistem Informasi Pengelolaan Sampah Nasional (SIPSN) menunjukkan bahwa pada tahun 2023, jumlah timbulan sampah di Indonesia mencapai 39,74 juta (ton/tahun)

dari 373 Kabupaten/kota se-Indonesia, dengan sampah yang terkelola sebesar 60,85% dan sampah yang tidak terkelola sebesar 39,15%. Kondisi ini dapat menimbulkan berbagai masalah serius, seperti pencemaran tanah, air, dan udara, hingga ancaman terhadap kesehatan masyarakat. Mencampur jenis-jenis sampah juga dapat menyebabkan bahaya apabila digabungkan menjadi satu (Sutanty & Kusuma Astuti, 2023). Berdasarkan jenis sampah, komposisi sampah terbanyak dihasilkan oleh sampah sisa makanan sebesar 39,67%, sedangkan berdasarkan sumber sampah, komposisi sampah terbanyak dihasilkan oleh rumah tangga sebesar 60,44% (*Sistem Informasi Pengelolaan Sampah Nasional: SIPSN*, 2024).

Pemilahan jenis sampah menjadi langkah awal yang penting untuk mengurangi dampak negatif dan meningkatkan pengelolaan limbah secara berkelanjutan. Pemilahan ini penting untuk mendukung upaya daur ulang, mengurangi beban tempat pembuangan akhir (TPA), serta meningkatkan efisiensi proses pengolahan limbah. Namun, pemilahan sampah secara manual masih menjadi tantangan besar, terutama di wilayah dengan tingkat kesadaran dan infrastruktur pengelolaan limbah yang rendah. Sampai saat ini, peran masyarakat secara umum hanya sebatas mengumpulkan dan membuang sampah saja (Ristya, 2020). Hal ini memicu perlunya solusi berbasis teknologi untuk membantu proses pemilahan secara cepat, tepat, dan efisien.

Di era teknologi digital, pembelajaran mesin (*machine learning*) telah membuka peluang besar untuk mengatasi masalah ini. Dengan menggunakan *algoritma* pembelajaran mesin, klasifikasi jenis sampah dapat dilakukan secara otomatis melalui analisis citra. Salah satu *algoritma* yang telah terbukti andal dalam klasifikasi data kompleks adalah *Random Forest*. Keunggulan utama dari *algoritma* ini adalah kemampuannya untuk meningkatkan akurasi apabila terdapat data yang hilang, untuk *resisting outliers*, dan menyimpan data dengan efisien (Devella et al., 2020). Selain itu, *Random Forest* memiliki proses seleksi fitur yang dapat mengambil fitur terbaik, sehingga dapat meningkatkan kinerja model klasifikasi (Supriyadi et al., 2020).

Beberapa penelitian sebelumnya telah menunjukkan efektivitas *algoritma* dan teknologi modern dalam pemilahan sampah. Misalnya, penggunaan model *ResNet-50* untuk klasifikasi jenis sampah berdasarkan citra visual menghasilkan akurasi tinggi hingga 98,7% dengan 7 kelas jenis sampah seperti kardus, kaca, logam, plastik, dan kompos (Nuariputri & Sukmasetya, 2023). Penelitian lain menggunakan *algoritma*

Random Forest untuk memprediksi atribut yang memengaruhi kualitas produk, seperti dalam kasus pengklasifikasian kualitas anggur merah, menunjukkan bahwa *algoritma* ini unggul dengan akurasi tertinggi dibandingkan *Decision Tree* dan *Support Vector Machine* (SVM) [4]. Selain itu, penelitian lain menggunakan *Random Forest* untuk klasifikasi citra bunga dengan metode segmentasi dan pengambilan fitur seperti *eccentricity*, *perimeter*, *metric*, dan *area* berhasil mencapai akurasi 100% dalam pengujian dengan skenario *K-fold cross-validation* (Rosyani et al., 2021).

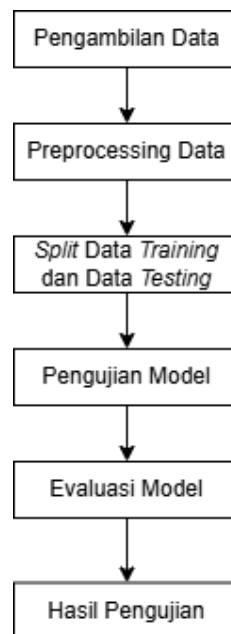
Dalam penelitian ini, kami memanfaatkan *PySpark*, *framework* berbasis *Apache Spark*, untuk membangun sistem klasifikasi jenis sampah menggunakan *algoritma Random Forest*. *PySpark* memberikan keunggulan dalam skalabilitas dan kinerja yang lebih baik untuk data besar dan memungkinkan pemrosesan yang terdistribusi (Handijono & Suhatman, 2024). Penelitian ini bertujuan untuk memberikan solusi praktis yang mendukung pengelolaan sampah berbasis teknologi dengan hasil akurasi yang andal. Hasil dari penelitian ini dapat menjadi bagian dari solusi untuk pengelolaan sampah yang lebih efektif, serta berkontribusi pada upaya global untuk menghasilkan lingkungan yang lebih bersih dan sehat.

2. METODOLOGI

Penelitian ini dilakukan melalui beberapa tahapan yang disajikan dalam bentuk diagram alur metode penelitian pada Gambar 1. Terdapat enam tahap dalam proses penelitian terkait klasifikasi sampah, yaitu pengambilan data, preprocessing data, split data training dan data testing, pengujian model, evaluasi model, dan hasil pengujian.

2.1. Dataset

Dataset yang digunakan dalam penelitian ini merupakan *open dataset* dengan nama *Real Waste* yang diperoleh dari platform *Kaggle*. Dataset ini berisi total 4.752 gambar yang telah dikelompokkan ke dalam beberapa kategori, yaitu *Cardboard*, *Food Organics*, *Glass*, *Metal*, *Miscellaneous Trash*, *Paper*, *Plastic*, *Textile Trash*, *Vegetation*. Dengan variasi yang mencakup jenis sampah organik maupun anorganik, dataset ini menjadi sumber data yang kaya untuk mengeksplorasi pendekatan berbasis pembelajaran mesin, pengenalan pola, atau pengolahan citra.



Gambar 1. Alur Penelitian

Keberagaman kategori dalam dataset ini memberikan peluang besar untuk mendukung penelitian yang mendalam di bidang pengolahan sampah, terutama dalam membangun model klasifikasi yang akurat dan efisien. Dengan jumlah data yang relatif besar dan distribusi kategori yang cukup beragam, dataset ini sangat ideal untuk digunakan dalam pengembangan teknologi berbasis pembelajaran mesin, baik untuk pengenalan pola, segmentasi gambar, maupun klasifikasi otomatis.

2.2. *Preprocessing Data*

Proses pengolahan data merupakan langkah penting dalam persiapan dataset gambar untuk digunakan dalam model pembelajaran mesin. Proses pengolahan data sangat penting untuk mencegah *algoritma* menjadi kurang akurat dalam memprediksi kualitas kopi. Karena itu, sebelum mengolah data, kita harus memastikan bahwa data yang akan kita gunakan bersih (Ciptady et al., 2022). Berikut ini merupakan langkah-langkah proses pengolahan data yang kami lakukan:

a. Koversi Format Gambar

Proses pertama dalam tahap preprocessing adalah konversi semua gambar ke format RGB. Hal ini dilakukan untuk memastikan konsistensi dalam pemrosesan gambar, menghindari perbedaan dalam saluran warna yang dapat mempengaruhi performa model.

b. *Resizing*

Langkah selanjutnya adalah mengubah ukuran gambar yang awalnya memiliki dimensi 524x524 piksel menjadi dimensi 64x64 piksel. Ukuran yang lebih kecil ini juga mempercepat proses pelatihan model dan memungkinkan pengolahan batch gambar dengan lebih efisien, yang sangat penting ketika bekerja dengan dataset besar.

c. *Label Ekstraksi*

Dalam konteks pengolahan gambar untuk tugas klasifikasi, label gambar diekstraksi dari nama file gambar. Misalnya, gambar dengan nama *Paper_001.jpg* akan memiliki label *Paper*. Label ini digunakan untuk melatih model, menghubungkan setiap gambar dengan kelas yang relevan. Dalam kasus pengolahan gambar dalam jumlah besar, teknik seperti pemrosesan paralel atau penggunaan pipeline otomatis dapat digunakan untuk mengekstraksi label ini dengan efisien.

Dengan struktur data ini, setiap gambar siap digunakan dalam model pembelajaran mesin yang dapat mengidentifikasi pola dan fitur dalam gambar untuk tugas klasifikasi atau deteksi objek.

2.3. *Split Data*

Dalam pengembangan model pembelajaran mesin, pembagian dataset memainkan peran yang sangat penting untuk memastikan bahwa model yang dikembangkan mampu belajar dengan baik serta dapat dievaluasi secara akurat. Dataset biasanya dibagi menjadi dua bagian utama, yaitu data pelatihan (*training data*) dan data pengujian (*testing data*). Data pelatihan adalah kumpulan data yang digunakan untuk melatih model menggunakan data yang benar, sedangkan data pengujian adalah kumpulan data yang digunakan untuk menguji kemampuan model untuk mengklasifikasikan data dengan benar (Baiq Nurul Azmi et al., 2023). Proses pembagian dataset dilakukan dengan proporsi tertentu yang umumnya ditentukan berdasarkan kebutuhan dan sifat dataset. Dalam penelitian ini, kami menggunakan pendekatan pembagian dataset dengan rasio 80% untuk data pelatihan dan 20% untuk data pengujian. Proporsi ini dipilih karena memberikan keseimbangan yang baik antara jumlah data yang cukup untuk melatih model secara optimal dan jumlah data yang memadai untuk mengevaluasi kinerja model secara objektif. Pembagian ini memastikan bahwa model memiliki cukup banyak data untuk belajar sekaligus diuji pada

data yang tidak termasuk dalam proses pelatihan, sehingga memberikan gambaran yang lebih akurat tentang kemampuan model untuk melakukan generalisasi pada data baru.

2.4. Pelatihan Model

Setiap pohon keputusan dalam *Random Forest* bekerja secara independen dan menghasilkan prediksi masing-masing, kemudian hasil prediksi dari seluruh pohon digabungkan untuk menghasilkan hasil akhir yang lebih akurat dan stabil. Berikut ini merupakan parameter yang kami gunakan dalam menggunakan metode *Random Forest*:

a. Jumlah Pohon

Parameter ini menentukan berapa banyak pohon keputusan yang akan dibangun dalam *Random Forest*. Pada penelitian ini, jumlah pohon yang kami gunakan adalah 10. Setiap pohon akan dilatih pada subset yang berbeda dari data pelatihan menggunakan teknik sampling dengan pengembalian. Ini berarti bahwa setiap pohon hanya akan melihat sebagian dari data pelatihan yang sama, memberikan variasi yang dapat meningkatkan generalisasi model.

b. Kedalaman Maksimum

Parameter ini membatasi kedalaman setiap pohon keputusan. Kedalaman pohon keputusan adalah jumlah lapisan atau level pembagian yang dilakukan pada data. Dengan menetapkan kedalaman maksimum menjadi 5, kami membatasi jumlah pembagian fitur yang dapat dilakukan oleh setiap pohon keputusan. Pembatasan ini bertujuan untuk mencegah model dari *overfitting*, yaitu ketika model terlalu kompleks dan terlalu cocok dengan data pelatihan, sehingga kinerjanya menurun pada data yang belum terlihat (data uji).

c. Pelatihan Model dengan Data Pelatihan

Model *Random Forest* dilatih menggunakan data pelatihan, yang merupakan data yang telah dipisahkan sebelumnya untuk proses pelatihan. Setiap pohon keputusan dilatih secara independen pada subset dari data pelatihan, yang dapat berisi data yang terduplikasi atau tidak ada. Pada setiap node dalam pohon keputusan, *algoritma* memilih fitur terbaik untuk melakukan pemisahan berdasarkan kriteria tertentu untuk meminimalkan kesalahan klasifikasi. Setelah pohon dibangun, proses ini diulang untuk setiap pohon dalam *Random Forest*. Meskipun masing-masing pohon hanya

melihat sebagian dari data, mereka bekerja sama untuk menghasilkan prediksi akhir yang lebih kuat.

3. HASIL DAN PEMBAHASAN

Hasil klasifikasi yang diperoleh dengan mengimplementasikan metode *Random Forest* pada dataset jenis sampah dapat memberikan wawasan yang mendalam mengenai kinerja model dalam mengidentifikasi dan mengklasifikasikan berbagai jenis sampah. Akurasi yang dicapai menjadi salah satu tolok ukur utama dalam mengevaluasi seberapa efektif metode ini dalam melakukan klasifikasi. Namun, untuk mendapatkan pemahaman yang lebih komprehensif mengenai kinerja model, perlu juga dilihat metrik lainnya seperti precision dan recall.

3.1. Evaluasi Model

Hasil evaluasi terhadap model klasifikasi menggunakan metode *Random Forest* ditunjukkan pada Tabel 1.

Tabel 1. Hasil Evaluasi Model

<i>Accuracy (%)</i>	<i>Precision (%)</i>	<i>Recall (%)</i>
44,64	48,62	44,64

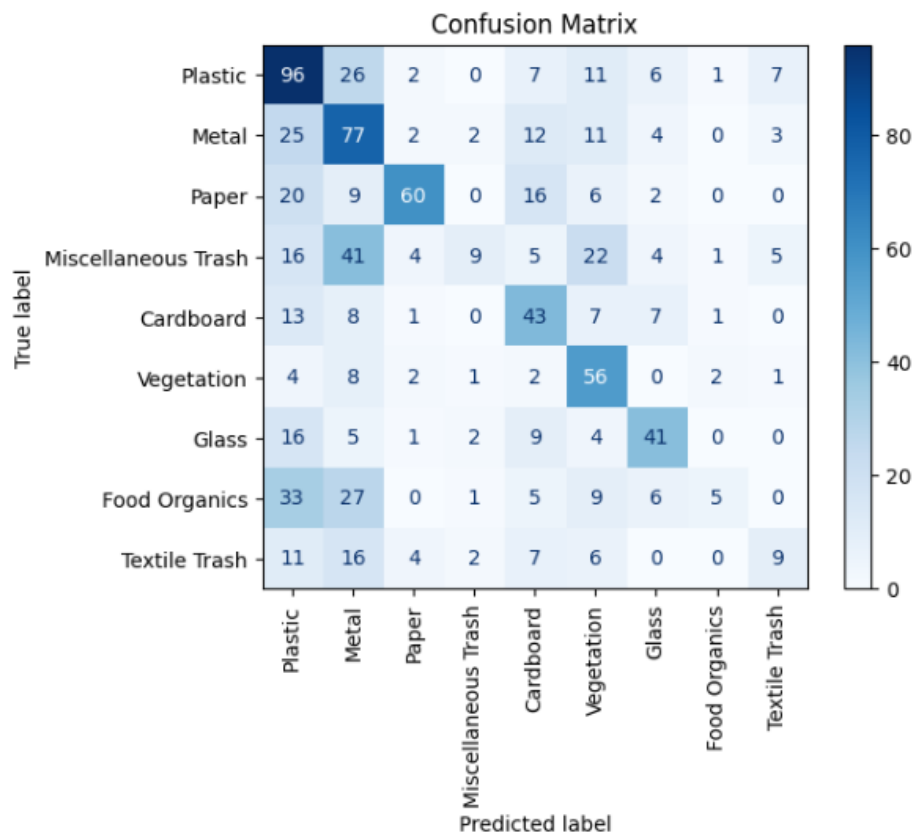
Berdasarkan Tabel 1. akurasi keseluruhan sebesar 44,64% pada data uji menunjukkan bahwa model yang digunakan masih memiliki performa yang belum optimal dalam mengklasifikasikan data sampah dengan tepat. Nilai akurasi ini mencerminkan proporsi prediksi yang benar dari keseluruhan data yang diuji, namun angka ini terbilang cukup rendah, yang mengindikasikan bahwa model belum mampu mencapai tingkat presisi yang tinggi dalam menentukan kategori sampah dengan benar.

Selain itu, nilai presisi yang tercatat sebesar 48,62% juga menunjukkan bahwa meskipun model mampu menghasilkan prediksi yang positif, ada sekitar 51,38% prediksi positif yang sebenarnya tidak akurat atau salah dalam klasifikasinya. Hal ini menunjukkan bahwa model lebih sering menghasilkan hasil positif palsu daripada yang seharusnya. Sementara itu, nilai *recall* yang sebesar 44,64% mengindikasikan bahwa model hanya berhasil menangkap sebagian kecil sampah yang seharusnya diklasifikasikan dalam kategori tertentu, meninggalkan banyak sampah yang sebenarnya milik kategori tersebut namun tidak terdeteksi. Secara keseluruhan, nilai presisi dan *recall*

yang rendah ini menunjukkan bahwa model masih menghadapi kesulitan yang signifikan dalam membedakan kategori sampah dengan akurat.

3.2. Confusion Matrix

Untuk mengevaluasi performa model secara mendalam, kami menggunakan *Confusion Matrix* untuk melihat distribusi prediksi model terhadap label sebenarnya. *Confusion matrix* membantu untuk melihat dengan jelas akurasi dari hasil klasifikasi yang sudah dilakukan. Gambar 2. menunjukkan *Confusion Matrix* dari hasil prediksi model.



Gambar 2. Hasil *Confusion Matrix*

Beberapa poin penting dari hasil *Confusion Matrix* adalah sebagai berikut:

a. Kategori dengan Akurasi Tinggi

Plastic, *Metal*, dan *Vegetation* merupakan tiga kategori dengan akurasi tertinggi diantara kategori yang lain. Kategori *Plastic* menunjukkan performa terbaik dengan 96 prediksi benar. Kategori *Metal* menunjukkan kinerja yang cukup baik dengan 77 prediksi benar. Kategori *Vegetation* berhasil dikenali dengan 56 prediksi benar, dengan kesalahan minimal ke kategori lain.

b. Kategori dengan Akurasi Menengah

Paper dan *Glass* merupakan dua kategori dengan akurasi menengah. Kategori *Paper* memperoleh 60 prediksi benar. Kategori *Glass* memperoleh 41 prediksi benar, kategori ini memiliki akurasi moderat.

c. Kategori dengan Akurasi Rendah

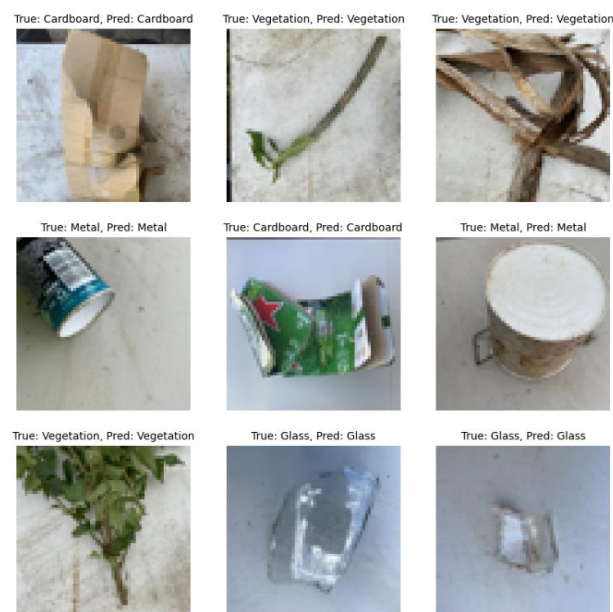
Food Organics, *Textile Trash*, dan *Miscellaneous Trash* merupakan tiga kategori dengan akurasi rendah. Kategori *Food Organics* hanya memperoleh 9 prediksi benar. Kategori *Textile Trash* memperoleh 8 prediksi benar. Kategori *Miscellaneous Trash* memperoleh 41 prediksi benar.

3.3. Hasil Pengujian

Dalam pengujian yang dilakukan terdapat dua hasil pengujian, yaitu hasil pengujian yang benar dan hasil pengujian yang salah. Hasil yang benar menunjukkan bahwa model mampu mengklasifikasikan objek dengan tepat sesuai dengan kategori yang seharusnya. Namun, ada juga hasil yang menunjukkan kesalahan klasifikasi, terutama ketika objek yang diuji memiliki kemiripan visual yang signifikan dengan objek lain.

a. Hasil Pengujian Benar

Pada hasil pengujian yang benar, model berhasil mengklasifikasikan objek dengan akurat sesuai dengan kategori yang seharusnya. Gambar 3. yang merupakan visualisasi hasil pengujian yang benar.



Gambar 3. Visualisasi Hasil Pengujian Benar

Terdapat gambar botol plastik berhasil dikenali sebagai plastik, dan sampah organik terdeteksi dengan benar sebagai sampah organik. Keberhasilan ini menandakan bahwa model dapat secara efektif mengenali fitur visual yang membedakan objek-objek ini, dan dapat mengklasifikasikannya dengan tingkat presisi yang tinggi. Hal ini mengindikasikan bahwa model telah belajar dengan baik dari data pelatihan, serta proses ekstraksi fitur yang diterapkan telah berjalan dengan optimal. Dalam konteks ini, model dapat menjadi alat yang sangat berguna dalam otomatisasi klasifikasi sampah, memberikan solusi yang efisien untuk pengelolaan sampah yang lebih baik dan ramah lingkungan.

b. Hasil Pengujian Salah

Pada hasil pengujian yang salah, terlihat bahwa model sering kali gagal dalam mengidentifikasi objek dengan benar ketika objek tersebut memiliki kemiripan fitur visual yang cukup tinggi dengan objek lain. Gambar 4. merupakan visualisasi hasil pengujian yang salah.



Gambar 4. Visualisasi Hasil Pengujian Salah

Salah klasifikasi semacam ini menunjukkan adanya kelemahan dalam proses ekstraksi fitur, yang mungkin dipengaruhi oleh resolusi gambar yang terlalu kecil. Resolusi rendah dapat menyebabkan hilangnya detail penting, sehingga model tidak mampu menangkap perbedaan halus antara objek yang sangat mirip. Penyebab lain dari kesalahan ini bisa jadi terkait dengan keterbatasan dalam dataset pelatihan, di

mana model belum terpapar cukup banyak variasi objek serupa. Untuk memperbaiki hasil pengujian yang salah ini, perlu dilakukan peningkatan resolusi gambar saat proses resize agar detail objek tetap terjaga. Dengan resolusi yang lebih tinggi, model dapat dilatih untuk lebih sensitif dalam membedakan objek yang serupa, sehingga meningkatkan akurasi dalam pengklasifikasian sampah dan objek lainnya. Di samping itu, peningkatan teknik ekstraksi fitur dan augmentasi data juga dapat membantu memperbaiki performa model.

4. KESIMPULAN

Berdasarkan hasil evaluasi model klasifikasi menggunakan metode *Random Forest*, kinerja model dalam mengidentifikasi berbagai jenis sampah masih tergolong belum optimal. Dengan akurasi sebesar 44,64%, model menunjukkan performa yang relatif rendah. Nilai presisi yang hanya mencapai 48,62% menunjukkan banyaknya kesalahan klasifikasi positif, sementara nilai *recall* sebesar 44,64% mengindikasikan bahwa model gagal mengenali sebagian besar sampah yang seharusnya masuk ke dalam kategori tertentu. Sedangkan, hasil dari *Confusion Matrix* memperlihatkan bahwa kategori seperti *Plastic*, *Metal*, dan *Vegetation* memiliki akurasi yang lebih baik dibandingkan kategori lain. Sebaliknya, kategori seperti *Food Organics*, *Textile Trash*, dan *Miscellaneous Trash* memiliki akurasi yang rendah. Hal ini menandakan bahwa model kurang efektif dalam menangani sampah dengan ciri visual yang kurang spesifik jika resolusi gambar rendah. Oleh karena itu, untuk meningkatkan performa, perbaikan pada ekstraksi fitur dan peningkatan resolusi dataset menjadi langkah yang penting agar model dapat lebih akurat dalam mengenali dan mengklasifikasikan berbagai jenis sampah.

5. DAFTAR PUSTAKA

- Baiq Nurul Azmi, Arief Hermawan, & Donny Avianto. (2023). Analisis Pengaruh Komposisi Data Training dan Data Testing pada Penggunaan PCA dan *Algoritma Decision Tree* untuk Klasifikasi Penderita Penyakit Liver. *JTIM : Jurnal Teknologi Informasi Dan Multimedia*, 4(4), 281–290. <https://doi.org/10.35746/jtim.v4i4.298>
- Ciptady, K., Harahap, M., Jonvin, J., Ndruru, Y., & Ibadurrahman, I. (2022). Prediksi Kualitas Kopi Dengan *Algoritma Random Forest* Melalui Pendekatan Data Science. *Data Sciences Indonesia (DSI)*, 2(1). <https://doi.org/10.47709/dsi.v2i1.1708>
- Danang Aji Kurniawan, D. A. K., & Ahmad Zaenal Santoso, A. Z. S. (2021). Pengelolaan Sampah di daerah Sepatan Kabupaten Tangerang. *ADI Pengabdian Kepada*

Masyarakat, 1(1), 31–36. <https://doi.org/10.34306/adimas.v1i1.247>

Devella, S., Yohannes, Y., & Rahmawati, F. N. (2020). Implementasi Random Forest Untuk Klasifikasi Motif Songket Palembang Berdasarkan SIFT. *JATISI (Jurnal Teknik Informatika Dan Sistem Informasi)*, 7(2), 310–320. <https://doi.org/10.35957/jatisi.v7i2.289>

Handijono, A., & Suhatman, Z. (2024). Meningkatkan Deduplikasi Data melalui Kesamaan Teks dalam Pembelajaran Mesin: Pendekatan Komprehensif. *AKADEMIK: Jurnal Mahasiswa Humanis*, 4(2), 602–615. <https://doi.org/10.37481/jmh.v4i2.955>

Nuariputri, J., & Sukmasetya, P. (2023). Klasifikasi Jenis Sampah Menggunakan Base ResNet-50. *Jurnal Ilmiah Komputasi*, 22(3), 379–386. <https://doi.org/10.32409/jikstik.22.3.3380>

Ristya, T. O. (2020). Penyuluhan Pengelolaan Sampah Dengan Konsep 3R Dalam Mengurangi Limbah Rumah Tangga. *Cakrawala: Jurnal Manajemen Pendidikan Islam Dan Studi Sosial*, 4(2), 30–41. <https://doi.org/10.33507/cakrawala.v4i2.250>

Rosyani, P., Saprudin, S., & Amalia, R. (2021). Klasifikasi Citra Menggunakan Metode Random Forest dan Sequential Minimal Optimization (SMO). *Jurnal Sistem Dan Teknologi Informasi (Justin)*, 9(2), 132. <https://doi.org/10.26418/justin.v9i2.44120>

Sistem Informasi Pengelolaan Sampah Nasional: SIPSN. (2024). 2024. Website: <https://sipsn.menlhk.go.id/>, diakses tanggal 12 Desember 2024.

Supriyadi, R., Gata, W., Maulidah, N., & Fauzi, A. (2020). Penerapan *Algoritma* Random Forest Untuk Menentukan Kualitas Anggur Merah. *E-Bisnis: Jurnal Ilmiah Ekonomi Dan Bisnis*, 13(2), 67–75. <https://doi.org/10.51903/e-bisnis.v13i2.247>

Sutanty, E., & Kusuma Astuti, D. (2023). Penerapan Model Arsitektur VGG16 untuk Klasifikasi Jenis Sampah. *DECODE: Jurnal Pendidikan Teknologi Informasi*, 3(2), 407–419. <https://doi.org/10.51454/decode.v3i2.331>