

Perbandingan Strategi Implementasi Model *MobileNetV2* untuk Pengenalan Jenis Sampah Otomatis dari Citra Digital

Alexandra Jennifer Matahurila¹, Elizabeth Anndini Shayna Putri², Windra Swastika³

Program Studi Teknik Informatika, Universitas Ma Chung, Villa Puncak Tidar Blok N-1, Malang, Jawa Timur, Indonesia, 65151

Correspondence: Alexandra Jennifer Matahurila (312310004@student.machung.ac.id)

Received: 27 05 2026 – Revised: 29 07 2026 - Accepted: 04 08 2026 - Published: 05 08 2026

Abstrak. Pengelolaan sampah yang belum terpilah dari sumbernya merupakan tantangan lingkungan global yang menghambat pencapaian target *Sustainable Development Goals (SDGs)* serta efektivitas ekonomi sirkular. Inefisiensi dan risiko pada proses pemilahan manual mendorong pengembangan solusi otomatis, sehingga penelitian ini menyajikan analisis komparatif mendalam untuk mengevaluasi dampak strategi penanganan data terhadap kinerja model *deep learning* MobileNetV2 dalam klasifikasi citra sampah. Dua pendekatan utama diinvestigasi secara sistematis. Pendekatan pertama berfungsi sebagai *baseline*, menerapkan *transfer learning* pada dataset asli yang distribusinya tidak seimbang. Sementara itu, pendekatan kedua dirancang sebagai strategi teroptimasi yang secara aktif mengintegrasikan augmentasi data (mencakup rotasi, *flipping*, dan penyesuaian kecerahan) untuk meningkatkan variasi, serta penyeimbangan kelas melalui metode *oversampling* untuk mengatasi bias. Proses pelatihan pada kedua model dioptimalkan menggunakan *callbacks* penting seperti *early stopping* untuk mencegah *overfitting* dan reduksi *learning rate* otomatis. Hasil evaluasi menunjukkan perbedaan performa yang signifikan. Model *baseline* mencapai akurasi uji 95,23%, namun dengan kinerja yang tidak stabil dan bias terhadap kelas mayoritas. Sebaliknya, model teroptimasi menunjukkan superioritas jelas dengan raihan akurasi uji 97,92%, didukung oleh kemampuan generalisasi yang kuat dan performa yang adil serta merata di semua kategori sampah. Dengan demikian, disimpulkan bahwa penerapan augmentasi dan penyeimbangan kelas adalah faktor fundamental yang secara signifikan meningkatkan keandalan dan akurasi model MobileNetV2, menjadikannya kandidat yang jauh lebih kuat untuk diimplementasikan pada sistem manajemen sampah cerdas.

Kata kunci: Augmentasi Data, Klasifikasi Sampah, MobileNetV2, Penyeimbangan Kelas, *Transfer Learning*

PENDAHULUAN

Volume sampah global terus meningkat seiring dengan pertumbuhan populasi dan urbanisasi. Menurut Bank Dunia, produksi sampah dunia diperkirakan akan mencapai 3,4 miliar ton per tahun pada tahun 2050, naik dari 2,01 miliar ton pada tahun 2016. Pengelolaan sampah yang tidak efektif, khususnya kegagalan dalam memilah sampah dari sumbernya, telah menjadi salah satu tantangan lingkungan paling mendesak saat ini. Masalah ini tidak hanya menyebabkan polusi tanah, air, dan udara, tetapi juga menghambat upaya daur ulang yang efisien, yang merupakan komponen kunci dalam ekonomi sirkular dan pencapaian target *Sustainable Development Goals (SDGs)*, terutama SDG 11 (*Make cities and human settlements inclusive, safe, resilient and sustainable*) dan SDG 12 (*Ensure sustainable consumption and production patterns*).

Secara tradisional, proses pemilahan sampah bergantung pada tenaga manusia, yang tidak hanya lambat dan padat karya, tetapi juga membawa risiko kesehatan bagi para pekerja. Sebagai respons terhadap tantangan ini, inovasi teknologi berbasis Kecerdasan Buatan (*Artificial Intelligence - AI*) menawarkan solusi yang menjanjikan. Salah satu pendekatan yang paling relevan adalah penggunaan *computer vision* dan *deep learning* untuk mengotomatiskan proses identifikasi dan klasifikasi sampah berdasarkan citra digital.

Model *Convolutional Neural Network (CNN)* telah menunjukkan kinerja yang luar biasa dalam tugas klasifikasi citra. Di antara berbagai arsitektur CNN, MobileNetV2 dikenal karena efisiensinya yang tinggi, menjadikannya pilihan ideal untuk aplikasi yang memerlukan kecepatan inferensi cepat dengan sumber daya komputasi terbatas, seperti pada perangkat seluler atau sistem tersemat (*embedded systems*). Arsitektur ini menggunakan *inverted residual blocks* dan *linear bottlenecks* untuk mengurangi jumlah parameter dan komputasi tanpa mengorbankan akurasi secara signifikan.

Meskipun demikian, keberhasilan implementasi model *deep learning* sangat bergantung pada kualitas dan kuantitas data pelatihan. Dalam konteks klasifikasi sampah, dataset sering kali mengalami dua masalah utama: ketidakseimbangan kelas (beberapa jenis sampah lebih umum daripada yang lain) dan kurangnya variasi dalam citra. Oleh karena itu, penelitian ini terinspirasi untuk mengeksplorasi bagaimana strategi penanganan data dapat memaksimalkan potensi arsitektur MobileNetV2.

Tujuan utama dari penelitian ini adalah untuk melakukan analisis komparatif terhadap dua strategi implementasi model MobileNetV2: (1) strategi dasar tanpa penanganan data khusus, dan (2) strategi yang diperkaya dengan teknik augmentasi data dan penyeimbangan kelas. Dengan membandingkan kedua pendekatan ini, penelitian ini bertujuan untuk mengidentifikasi strategi optimal yang dapat menghasilkan model klasifikasi sampah yang tidak hanya akurat tetapi juga andal dan dapat digeneralisasi dengan baik untuk aplikasi dunia nyata.

MASALAH

Persoalan utama yang mendasari penelitian ini berakar pada tantangan teknis dalam pengembangan model klasifikasi citra yang andal untuk objek dunia nyata yang kompleks seperti sampah. Meskipun arsitektur MobileNetV2 menawarkan efisiensi, kinerjanya dapat terdegradasi secara signifikan jika tidak didukung oleh strategi data yang tepat. Secara spesifik, penelitian ini berfokus pada dua masalah keilmuan utama:

1. **Ketidakseimbangan Kelas dalam Dataset Sampah:** Dataset citra sampah yang tersedia secara publik sering kali tidak seimbang. Sebagai contoh, dalam dataset yang digunakan pada penelitian ini, jumlah citra untuk kategori 'clothes' (pakaian) jauh lebih banyak dibandingkan kategori 'trash' (sampah umum) atau 'battery' (baterai). Jika model dilatih pada data yang tidak seimbang ini tanpa penanganan khusus, ia akan cenderung menjadi bias dan sangat mahir dalam mengenali kelas mayoritas, namun memiliki performa yang buruk pada kelas minoritas. Hal ini sangat merugikan dalam aplikasi praktis, di mana kesalahan klasifikasi pada jenis sampah tertentu (misalnya, baterai yang berbahaya) dapat memiliki konsekuensi serius.
2. **Keterbatasan Generalisasi Model Akibat Variasi Citra:** Citra sampah di dunia nyata memiliki variasi yang sangat tinggi dalam hal pencahayaan, sudut pengambilan gambar, oklusi (tertutup sebagian), dan deformasi bentuk. Model yang dilatih pada dataset dengan variasi terbatas akan kesulitan untuk melakukan generalisasi pada data baru yang belum pernah dilihat sebelumnya (*unseen data*). Fenomena ini, yang dikenal sebagai *overfitting*, menyebabkan model memiliki akurasi sangat tinggi pada data pelatihan tetapi performanya menurun drastis pada data validasi dan pengujian. Hal ini menunjukkan bahwa model hanya "menghafal" data pelatihan, bukan "mempelajari" fitur-fitur intrinsik dari setiap kategori sampah.

Oleh karena itu, penelitian ini menguraikan perbandingan dua skenario untuk menjawab masalah tersebut. Skenario pertama menjadi tolok ukur (*baseline*) performa model dalam kondisi data yang "apa adanya", sementara skenario kedua secara aktif mencoba menyelesaikan masalah ketidakseimbangan dan generalisasi melalui rekayasa data.

METODE PELAKSANAAN

Penelitian ini menggunakan pendekatan kuantitatif eksperimental dengan metode simulasi Ipteks. Metode ini dipilih karena kegiatan utamanya adalah membangun dan menguji model komputasi (*deep learning*) untuk mensimulasikan tugas kognitif manusia (mengenali jenis sampah), yang bertujuan untuk menghasilkan pengetahuan baru mengenai strategi implementasi yang paling efektif.

Langkah-langkah pelaksanaan penelitian diuraikan sebagai berikut:

a. Dataset dan Sumber Data

Penelitian ini menggunakan dataset publik "Garbage Classification V2" yang tersedia di platform Kaggle. Dataset ini berisi citra-citra yang telah dikelompokkan ke dalam 10 kelas

yang berbeda, yaitu: *battery, biological, cardboard, clothes, glass, metal, paper, plastic, shoes, dan trash*.

b. Desain Eksperimen dan Pra-pemrosesan Data

Untuk membandingkan dampak penanganan data, penelitian ini dirancang dengan dua strategi eksperimental yang berbeda:

- **Strategi 1: Model Baseline (Tanpa Augmentasi & Penyeimbangan)**
 - Dataset asli digunakan tanpa modifikasi. Data dibagi menjadi set pelatihan (80%), validasi (10%), dan pengujian (10%).
 - Pembagian datanya adalah: 13.827 gambar latih, 2.961 gambar validasi, dan 2.974 gambar uji.
 - Distribusi kelas dalam set ini dibiarkan tidak seimbang, mencerminkan kondisi data asli.
- **Strategi 2: Model yang Dioptimalkan (Dengan Augmentasi & Penyeimbangan)**
 - **Penyeimbangan Kelas:** Teknik *oversampling* diterapkan untuk menyeimbangkan distribusi kelas. Setiap kelas dalam dataset pelatihan, validasi, dan pengujian disamakan jumlahnya. Berdasarkan program yang dijalankan, setiap kelas di-set menjadi 4000 gambar, menghasilkan total 28.000 gambar latih, 6.000 gambar validasi, dan 6.000 gambar uji.
 - **Augmentasi Data:** Pada set pelatihan, teknik augmentasi data diterapkan secara *on-the-fly*. Teknik ini bertujuan untuk meningkatkan variasi data secara artifisial. Augmentasi yang digunakan meliputi:
 - **Rotasi:** Memutar gambar secara acak dalam rentang derajat tertentu.
 - **Flipping:** Membalik gambar secara horizontal dan vertikal.
 - **Brightness Adjustment:** Mengubah tingkat kecerahan gambar secara acak.

c. Arsitektur dan Pelatihan Model

- **Arsitektur Model:** Kedua strategi menggunakan arsitektur **MobileNetV2** dengan pendekatan *transfer learning*. Model MobileNetV2 yang telah dilatih sebelumnya pada dataset ImageNet digunakan sebagai basis (*base model*). Lapisan-lapisan konvolusi pada basis model ini "dibekukan" (*frozen*) agar bobotnya tidak berubah selama pelatihan awal. Di atas basis model ini, ditambahkan beberapa lapisan baru (*classification head*), yang terdiri dari GlobalAveragePooling2D, Dropout untuk regularisasi, dan sebuah lapisan Dense dengan 10 unit neuron dan fungsi aktivasi softmax untuk klasifikasi 10 kelas sampah.
- **Proses Pelatihan:** Model pada kedua strategi dilatih dengan konfigurasi berikut:

- **Optimizer:** Adam, dengan *learning rate* awal $1.00e-04$.
- **Fungsi Kerugian (*Loss Function*):** CategoricalCrossentropy, karena ini adalah masalah klasifikasi multi-kelas.
- **Callbacks:** Untuk mengoptimalkan pelatihan dan mencegah *overfitting*, digunakan tiga *callbacks*:
 1. ModelCheckpoint: Menyimpan bobot model dengan akurasi validasi (*val_accuracy*) terbaik secara otomatis.
 2. EarlyStopping: Menghentikan proses pelatihan jika *val_accuracy* tidak menunjukkan peningkatan setelah beberapa *epoch* tertentu (misalnya, 10 *epoch*), dan mengembalikan bobot terbaik yang tersimpan. Ini mencegah model dari *overfitting* yang berlebihan.
 3. ReduceLROnPlateau: Mengurangi *learning rate* secara otomatis jika metrik *val_accuracy* berhenti meningkat, untuk membantu model menemukan titik minimum yang lebih baik.

d. Metrik Evaluasi

Kinerja model dari kedua strategi dievaluasi menggunakan metrik standar untuk masalah klasifikasi:

- **Akurasi (*Accuracy*):** Persentase prediksi yang benar dari keseluruhan data. Metrik ini memberikan gambaran umum performa model.
- **Laporan Klasifikasi (*Classification Report*):** Laporan ini memberikan metrik yang lebih rinci untuk setiap kelas, yaitu:
 - **Presisi (*Precision*):** Dari semua yang diprediksi sebagai kelas X, berapa persen yang benar-benar kelas X.
 - **Daya Ingat (*Recall*):** Dari semua data yang seharusnya kelas X, berapa persen yang berhasil diprediksi dengan benar oleh model.
 - **Skor F1 (*F1-Score*):** Rata-rata harmonik dari presisi dan *recall*, memberikan ukuran tunggal yang menyeimbangkan keduanya. Metrik ini sangat penting untuk mengevaluasi kinerja pada dataset yang tidak seimbang.
- **Matriks Kebingungan (*Confusion Matrix*):** Sebuah tabel yang memvisualisasikan performa model dengan menunjukkan jumlah prediksi benar dan salah untuk setiap kelas. Ini membantu mengidentifikasi kelas mana yang sering tertukar oleh model.
- **Grafik Akurasi dan Kerugian:** Plot yang menunjukkan perubahan akurasi dan *loss* pada data latih dan validasi di setiap *epoch*, digunakan untuk menganalisis perilaku belajar model.

HASIL DAN PEMBAHASAN

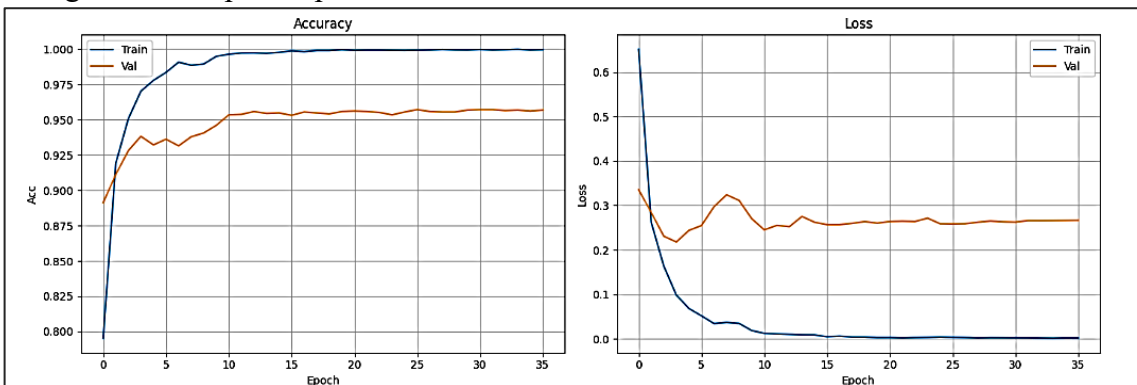
Pada bagian ini, disajikan temuan-temuan dari kedua eksperimen yang telah dilakukan. Hasil kuantitatif dari log pelatihan dan pengujian dianalisis dan dibahas secara komparatif untuk menarik kesimpulan yang valid.

Hasil Eksperimen 1: Model Baseline (Tanpa Augmentasi & Penyeimbangan)

Model pertama dilatih pada dataset asli yang tidak seimbang. Proses pelatihan dihentikan oleh *callback EarlyStopping* pada *epoch* ke-36, dengan performa terbaik pada *epoch* ke-26. Hasil akhir yang dicapai adalah sebagai berikut:

- Akurasi Pelatihan Puncak: **99,95%**
- Akurasi Validasi Puncak: **95,71%**
- Akurasi Akhir pada Data Uji: **95,23%**

Perilaku belajar model ini dapat diamati secara visual melalui grafik akurasi dan kerugian selama proses pelatihan.

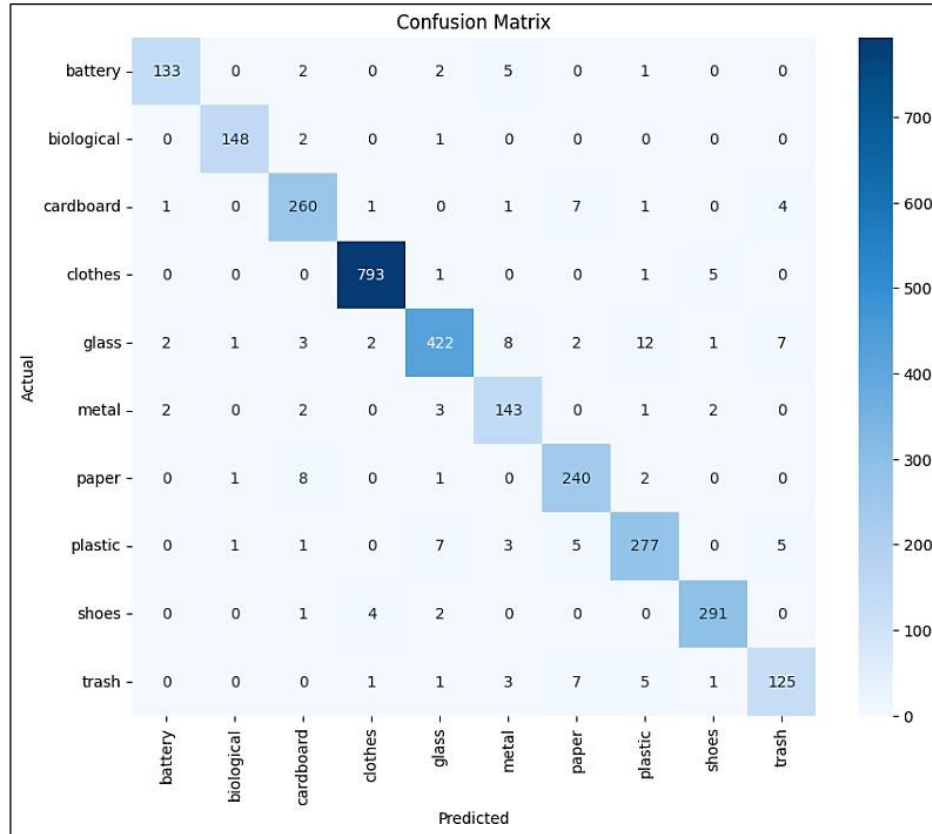


Gambar 1. Grafik Akurasi (*Accuracy*) dan Kerugian (*Loss*) pada Pelatihan Model Baseline (Tanpa Augmentasi & Penyeimbangan). Garis biru menunjukkan metrik pada data pelatihan, sedangkan garis orange menunjukkan metrik pada data validasi.

Analisis Gambar 1:

Dari grafik pada Gambar 1, terlihat jelas adanya fenomena *overfitting*. Kurva akurasi pelatihan (biru) menanjak dengan cepat dan mendekati 100%, sementara kurva akurasi validasi (oranye) mengalami stagnasi di sekitar 95% dan cenderung tidak stabil. Celah (gap) yang signifikan antara kedua kurva ini mengindikasikan bahwa model menjadi terlalu "hafal" terhadap data pelatihan dan kesulitan melakukan generalisasi pada data baru yang belum pernah dilihatnya (data validasi). Hal serupa terlihat pada grafik *loss*, di mana *loss* pelatihan terus menurun sementara *loss* validasi berhenti menurun dan mulai berfluktuasi.

Untuk mengevaluasi performa model pada setiap kelas secara spesifik, digunakan *confusion matrix* dan *classification report*.



Gambar 2. *Confusion Matrix* Kinerja Model Baseline pada Data Uji. Sumbu vertikal (y-axis) merepresentasikan kelas sebenarnya (*True Label*), dan sumbu horizontal (x-axis) merepresentasikan kelas yang diprediksi (*Predicted Label*). Nilai pada diagonal utama menunjukkan jumlah prediksi yang benar.

Analisis Gambar 2:

Confusion matrix pada Gambar 2 secara visual mengonfirmasi adanya ketidakseimbangan performa. Angka-angka pada diagonal utama menunjukkan jumlah prediksi yang benar, sedangkan angka di luar diagonal menunjukkan kesalahan klasifikasi. Terlihat bahwa kelas 'clothes' (dengan 800 sampel uji) memiliki nilai diagonal yang sangat tinggi. Sebaliknya, dapat diamati adanya kesalahan klasifikasi yang signifikan pada kelas-kelas minoritas. Sebagai contoh, sejumlah gambar 'trash' dan 'metal' salah diklasifikasikan sebagai kelas lain, yang ditunjukkan oleh nilai yang relatif tinggi di luar diagonal utama pada baris-baris tersebut. Temuan visual ini diperkuat oleh metrik kuantitatif dalam laporan klasifikasi berikut.

Meskipun akurasi uji 95,23% terlihat tinggi, analisis lebih dalam pada *Classification Report* menunjukkan adanya kelemahan:

Kelas	Presisi	Recall	F1-Score
clothes	0.99	0.99	0.99
biological	0.98	0.98	0.98
shoes	0.97	0.98	0.97
battery	0.96	0.93	0.95
glass	0.96	0.92	0.94
paper	0.92	0.95	0.94
cardboard	0.93	0.95	0.94
plastic	0.92	0.93	0.92
metal	0.88	0.93	0.91
trash	0.89	0.87	0.88

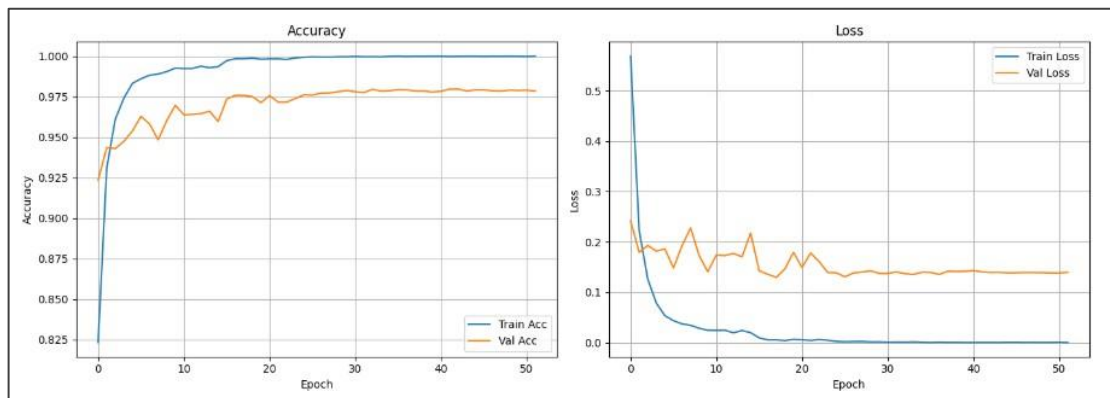
Tabel 1. *Classification Report* pada Model Baseline
(Tanpa Augmentasi & Penyeimbangan)

Terlihat bahwa kelas mayoritas seperti 'clothes' mencapai F1-Score hampir sempurna (0.99), sementara kelas minoritas seperti 'trash' dan 'metal' memiliki F1-Score yang jauh lebih rendah (0.88 dan 0.91). Ini mengonfirmasi hipotesis bahwa model menjadi bias terhadap kelas dengan jumlah data yang lebih banyak. Selain itu, terdapat celah (gap) yang cukup signifikan antara akurasi pelatihan (99,95%) dan akurasi validasi (95,71%), yang mengindikasikan adanya *overfitting*.

Hasil Eksperimen 2: Model yang Dioptimalkan (Dengan Augmentasi & Penyeimbangan)

Model kedua dilatih pada dataset yang telah diseimbangkan dan diaugmentasi. Pelatihan berhenti pada *epoch* ke-52, dengan performa terbaik tercatat pada *epoch* ke-42. Hasil yang diperoleh menunjukkan peningkatan signifikan:

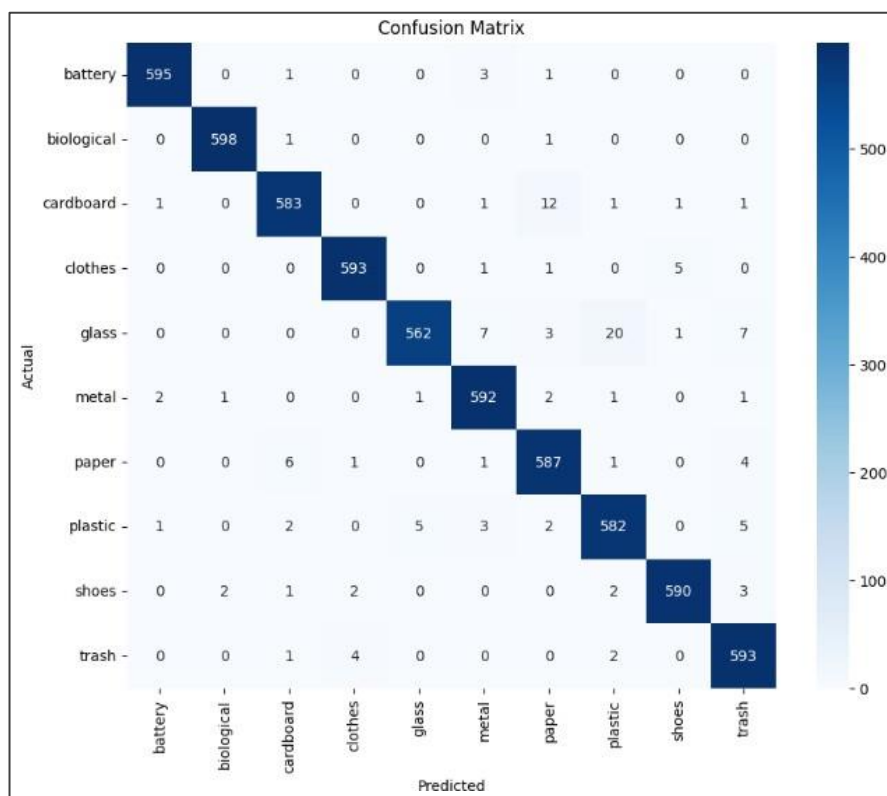
- Akurasi Pelatihan Puncak: **100,00%**
- Akurasi Validasi Puncak: **97,97%**
- Akurasi Akhir pada Data Uji: **97,92%**



Gambar 3 Grafik Akurasi (*Accuracy*) dan Kerugian (*Loss*) pada Pelatihan Model yang Dioptimalkan (Dengan Augmentasi & Penyeimbangan).menunjukkan jumlah prediksi yang benar.

Analisis Gambar 3:

Berbeda secara drastis dengan Gambar 1, grafik pada Gambar 3 menunjukkan perilaku belajar yang jauh lebih sehat. Kurva akurasi pelatihan dan validasi bergerak naik secara bersamaan dengan celah yang sangat kecil di antara keduanya. Ini menandakan bahwa augmentasi data berhasil memaksa model untuk belajar fitur yang lebih general, sehingga mengurangi *overfitting* secara signifikan. Kedua kurva sama-sama mencapai tingkat performa yang tinggi dan stabil, menunjukkan bahwa model mampu menggeneralisasi pengetahuannya dengan sangat baik dari data latih ke data validasi.



Gambar 4. *Confusion Matrix* Kinerja Model yang Dioptimalkan pada Data Uji.

Analisis Gambar 4:

Confusion matrix untuk model yang dioptimalkan pada Gambar 4 menunjukkan peningkatan yang sangat jelas dibandingkan Gambar 2. Diagonal utama terlihat jauh lebih "dominan", dengan nilai-nilai yang tinggi dan konsisten di semua kelas. Angka-angka di luar diagonal (yang merepresentasikan kesalahan) menjadi jauh lebih kecil. Secara khusus, kesalahan klasifikasi untuk kelas-kelas yang sebelumnya bermasalah seperti 'trash' dan 'metal' telah berkurang secara drastis. Visualisasi ini menegaskan bahwa strategi penyeimbangan kelas berhasil membuat model mengenali setiap kategori sampah dengan tingkat keandalan yang setara.

Peningkatan akurasi uji sebesar **2,69%** (dari 95,23% menjadi 97,92%) adalah pencapaian yang penting. Namun, keunggulan utama strategi ini terlihat pada *Classification Report* yang jauh lebih merata:

Kelas	Presisi	Recall	F1-Score
biological	1.00	1.00	1.00
clothes	0.99	0.99	0.99
battery	0.99	0.99	0.99
shoes	0.99	0.98	0.99
cardboard	0.98	0.97	0.98
metal	0.97	0.99	0.98
trash	0.97	0.99	0.98
paper	0.96	0.98	0.97
plastic	0.96	0.97	0.96
glass	0.99	0.94	0.96

Tabel 2. *Classification Report* pada Model yang Dioptimalkan
(Dengan Augmentasi & Penyeimbangan)

Semua kelas kini memiliki F1-Score di atas 0.96, termasuk kelas 'metal' dan 'trash' yang sebelumnya berkinerja rendah. Ini membuktikan bahwa teknik penyeimbangan kelas berhasil mengatasi bias model, memaksanya untuk belajar fitur dari semua kelas secara setara.

Pembahasan Komparatif

Perbandingan langsung antara kedua strategi menunjukkan keunggulan mutlak dari Strategi 2. Temuan ini menegaskan pentingnya penanganan data sebagai fondasi untuk membangun model *deep learning* yang kuat.

1. **Dampak Penyeimbangan Kelas:** Dengan menyamakan jumlah sampel untuk setiap kelas, Strategi 2 berhasil menghilangkan bias yang terlihat pada Strategi 1. Model tidak lagi "malas" dengan hanya menebak kelas mayoritas. Hal ini sejalan dengan banyak penelitian di bidang klasifikasi citra yang menunjukkan bahwa penanganan ketidakseimbangan data, baik melalui *oversampling* (seperti pada penelitian ini), *undersampling*, atau metode sintetik seperti SMOTE, adalah langkah krusial untuk keadilan dan keandalan model.
2. **Dampak Augmentasi Data:** Augmentasi data secara artifisial memperkaya dataset pelatihan, menyajikan model dengan variasi gambar yang lebih luas. Hal ini memaksa model untuk belajar fitur yang lebih robust dan invarian terhadap perubahan posisi, orientasi, dan pencahayaan. Akibatnya, model dari Strategi 2 menunjukkan kemampuan generalisasi yang lebih baik, dibuktikan dengan celah yang lebih kecil antara performa pada data validasi dan data uji. Temuan ini mendukung penelitian oleh Chlap et al. (2021) yang menyimpulkan bahwa augmentasi adalah teknik regularisasi yang sangat efektif untuk meningkatkan kinerja model pada data medis.

Singkatnya, sementara Strategi 1 menghasilkan model yang tampak bagus di permukaan, kinerjanya tidak dapat diandalkan dan tidak adil. Sebaliknya, Strategi 2, dengan investasi tambahan dalam pra-pemrosesan data, menghasilkan model yang secara kuantitatif lebih akurat dan secara kualitatif lebih andal dan siap untuk aplikasi dunia nyata.

KESIMPULAN

Penelitian ini berhasil melakukan analisis perbandingan antara dua strategi implementasi model MobileNetV2 untuk klasifikasi sampah digital. Hasilnya secara tegas menunjukkan bahwa strategi implementasi yang mengintegrasikan teknik penyeimbangan kelas melalui *oversampling* dan augmentasi data (Strategi 2) secara signifikan lebih unggul daripada strategi dasar (Strategi 1). Model yang dioptimalkan mencapai akurasi uji sebesar 97,92%, dengan kinerja yang kuat dan merata di semua 10 kelas sampah. Peningkatan ini membuktikan bahwa penanganan masalah fundamental seperti ketidakseimbangan kelas dan kurangnya variasi data adalah langkah yang lebih kritis daripada sekadar memilih arsitektur model yang canggih. Sebagai rekomendasi, implementasi sistem pemilahan sampah otomatis berbasis AI di masa depan harus memprioritaskan penggunaan dataset yang seimbang dan beragam. Untuk penelitian selanjutnya, model yang dihasilkan dapat diuji pada data video *real-time* atau diimplementasikan pada perangkat keras *edge computing* untuk menguji kinerjanya dalam skenario operasional yang sebenarnya.

UCAPAN TERIMA KASIH

Penulis mengucapkan terima kasih kepada Windra Swastika, Ph.D, selaku dosen pengampu mata kuliah Deep Learning di Universitas Ma Chung, yang telah memberikan bimbingan dan arahan yang sangat berharga selama proses pengerjaan proyek dan penulisan makalah ini. Ucapan terima kasih juga ditujukan kepada rekan-rekan satu tim dan semua pihak yang telah memberikan dukungan.

DAFTAR PUSTAKA

- Al-Taff, R., & Al-Nuaimy, W. (2024). Garbage Classification from Visual Footprints: Using Transfer Learning Strategy. In Proceedings of the 13th International Conference on Pattern Recognition Applications and Methods - ICPRAM (pp. 571–578). SciTePress.
- Huta Julu, D. I. S., & Nurdiah, D. (2025). KLASIFIKASI SAMPAH ORGANIK DAN NON ORGANIK MENGGUNAKAN TRANSFER LEARNING. Jurnal Transformatika, 23(1).
- Valle, E., Tso, G. Y. F., & Yavari, A. (2024). A data augmentation methodology to reduce the class imbalance in histopathology images. Journal of Digital Imaging, 37, 1767–1782.
- World Bank. (2018). What a Waste 2.0: A global snapshot of solid waste management to 2050. World Bank. <https://doi.org/10.1596/978-1-4648-1329-0>
- Yong, L., Ma, L., Sun, D., & Du, L. (2023). Application of MobileNetV2 to waste classification. PLOS ONE, 18(3), e0282336. <https://doi.org/10.1371/journal.pone.0282336>

