

K-Nearest Neighbor dengan Jarak Euclidean, Manhattan, dan Minkowski pada klasifikasi sampah

Siti Rahmi Kelilauw^a, Purnawansyah^b, Abdul Rachman Manga^c, Rahma Puspitasari

Universitas Muslim Indonesia, Makassar, Indonesia

Email: ^akelilauwsitirahmi@gmail.com; ^bpurnawansyah@umi.ac.id; ^cabdulrachman.manga@umi.ac.id; rahmapuspitasari.iclabs@umi.ac.id

Received: 17-02-2025 | Revised: 02-08-2025 | Accepted: 12-09-2025 | Published: 29-09-2025

Abstrak

Sampah menjadi permasalahan besar di Indonesia seiring dengan meningkatnya populasi dan konsumsi masyarakat. Untuk mengatasi permasalahan ini, klasifikasi sampah yang akurat diperlukan guna mendukung pengelolaan dan daur ulang yang lebih efektif. Penelitian ini menerapkan algoritma *K-Nearest Neighbors* (KNN) untuk mengklasifikasikan sampah berdasarkan jenisnya, dengan menggunakan tiga metrik jarak: Euclidean, Manhattan, dan Minkowski. Dataset yang digunakan terdiri dari 997 gambar yang terbagi menjadi dua kelas, yaitu kertas dan kardus. Data dibagi menjadi 80% untuk pelatihan dan 20% untuk pengujian. Proses *preprocessing* meliputi *resize* gambar ke ukuran 128x128 piksel serta normalisasi data. Model dievaluasi menggunakan metrik akurasi, presisi, *recall*, dan *F1-score*. Hasil penelitian menunjukkan bahwa jarak Manhattan memberikan performa terbaik dengan akurasi 83%, diikuti oleh Euclidean dengan 75,50%, dan Minkowski dengan 66,50%. Temuan ini menegaskan bahwa pemilihan metrik jarak dalam algoritma KNN sangat memengaruhi kinerja model. Dengan pendekatan ini, diharapkan dapat meningkatkan efisiensi sistem pengelolaan sampah berbasis teknologi, sehingga mendukung upaya daur ulang yang lebih baik dan ramah lingkungan.

Kata kunci: klasifikasi sampah, KNN, Euclidean, Manhattan, Minkowski

Pendahuluan

Sampah menjadi masalah besar di Indonesia. Seiring dengan pertambahan jumlah penduduk dan peningkatan konsumsi, volume sampah yang dihasilkan semakin banyak. Sampah ini tidak hanya mencakup sampah rumah tangga, tetapi juga sampah industri, plastik, dan sampah organik yang belum dikelola dengan baik [1]. Sampah menjadi masalah serius yang dihadapi oleh banyak negara, terutama Indonesia. Volume dan variasi jenis sampah di Indonesia terus meningkat setiap tahunnya seiring dengan pertumbuhan jumlah penduduk. Menurut data dari Sistem Informasi Penanggulangan Sampah Nasional (SIPSN) yang dikutip dari situs resmi SIPSN per 2021, total timbunan sampah di Indonesia mencapai 24,67 juta ton per tahun. Terdapat pengurangan sampah sebesar 13,38% atau sekitar 3,3 juta ton dibandingkan tahun sebelumnya [2]. Namun, penanganan sampah di Indonesia hanya dapat mengelola sekitar 50,43% atau 12,44 juta ton per tahun. Dengan jumlah tersebut, Indonesia menghasilkan sekitar 67.590 ton sampah atau setara dengan 0,25 kg per orang per hari. Angka-angka ini menunjukkan bahwa Indonesia saat ini berada dalam situasi darurat sampah. Masalah ini merupakan tantangan besar yang harus ditangani bersama. Berdasarkan Undang-Undang Republik Indonesia Nomor 18 Tahun 2008 tentang Pengelolaan Sampah, Pasal 1 (1) menyatakan bahwa "sampah adalah sisa kegiatan sehari-hari manusia atau proses alam yang berbentuk padat" [3]. Sampah dibedakan menjadi dua jenis, yaitu sampah anorganik dan sampah organik, yang didasarkan pada sifatnya. Sampah anorganik adalah sampah yang tidak dapat terurai secara alami, seperti logam, kaca, dan plastik. Sedangkan sampah organik biasanya berupa sampah yang sudah membusuk, seperti sisa makanan, daun-daunan, dan buah-buahan. Saat ini, sampah di lingkungan sekitar kita seringkali tercampur dan tidak dipilah dengan benar. Hal ini disebabkan oleh kurangnya pemahaman dan kesadaran masyarakat dalam membuang sampah sesuai dengan jenisnya.

Salah satu langkah awal dalam pengelolaan sampah adalah klasifikasi sampah berdasarkan jenisnya. Dengan klasifikasi yang tepat, proses pengolahan dan daur ulang dapat dilakukan secara lebih efektif. Dalam konteks ini, teknologi machine learning dapat memberikan solusi yang inovatif untuk meningkatkan akurasi dan efisiensi klasifikasi sampah [4],[5]. Algoritma *K-Nearest Neighbors* (KNN) merupakan salah satu metode machine learning yang sederhana namun efektif untuk menyelesaikan masalah klasifikasi, termasuk dalam pengelompokan sampah berdasarkan jenisnya [6],[7].

Penelitian ini berfokus pada penggunaan algoritma KNN untuk mengklasifikasikan dua jenis sampah, yaitu kertas dan kardus. Pemilihan algoritma KNN didasarkan pada kemampuannya untuk menangani data yang

beragam dan kesederhanaannya dalam implementasi. Dalam penelitian ini, nilai parameter K ditetapkan sebesar 3, yang merupakan nilai optimal berdasarkan pengujian awal [8],[9].

Tujuan dari penelitian ini adalah untuk mengevaluasi performa algoritma KNN dalam klasifikasi sampah serta memberikan wawasan tentang potensi penerapannya dalam sistem pengelolaan sampah berbasis teknologi. Dengan demikian, penelitian ini diharapkan dapat berkontribusi dalam upaya pengelolaan sampah yang lebih efisien dan ramah lingkungan.

Metode

K-Nearest Neighbors (KNN) adalah algoritma supervised learning yang digunakan untuk klasifikasi dan regresi. Algoritma ini bekerja berdasarkan prinsip kesamaan (similarity), yaitu menentukan kelas atau nilai berdasarkan sejumlah tetangga terdekat K [10],[11]. KNN adalah pengklasifikasi pembelajaran non-parametrik dan terawasi, yang menggunakan kedekatan untuk membuat klasifikasi atau prediksi tentang pengelompokan titik data individu [12],[13]. Ini adalah salah satu klasifikasi dan regresi yang populer dan paling sederhana yang digunakan dalam machine learning. Untuk menentukan tetangga terdekat, digunakan fungsi jarak seperti Euclidean, jarak Manhattan, dan jarak Minkowski seperti pada persamaan (1) – (3).

$$d(x_i, y_i) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2} \quad (1)$$

$d(x, y) = \sum ((x_i - y_i)^2)$, di mana x_i dan y_i adalah koordinat titik x dan y pada dimensi ke- i^{th} . Proses perhitungannya melibatkan selisih antara setiap koordinat yang kemudian dikuadratkan, dijumlahkan, dan diambil akar kuadratnya.

$$d(x, y) = \sum_{i=1}^n |x_i - y_i| \quad (2)$$

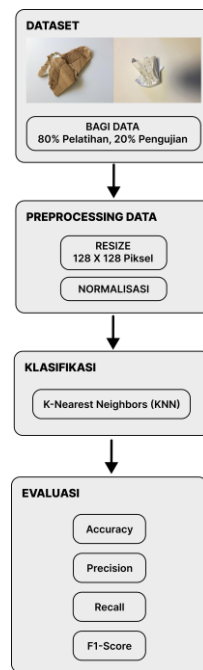
Jarak Manhattan $d(x, y)$ antara dua titik x dan y didefinisikan sebagai jumlah dari nilai absolut perbedaan koordinat mereka pada setiap dimensi. Jika data memiliki n dimensi, maka setiap titik x dan y memiliki koordinat x_i dan y_i pada dimensi ke- i^{th} . Nilai absolut dari selisih koordinat antara kedua titik pada dimensi ke- i^{th} , yaitu $|x_i - y_i|$, dihitung untuk setiap dimensi, kemudian dijumlahkan untuk mendapatkan jarak Manhattan secara keseluruhan.

$$d(x, y) = \left(\sum_{i=1}^n |x_i - y_i|^p \right)^{\frac{1}{p}} \quad (3)$$

Jarak Minkowski adalah metrik umum yang digunakan untuk mengukur jarak antara dua titik dalam ruang vektor dengan jumlah dimensi tertentu. Metrik ini merupakan generalisasi dari jarak Euclidean dan Manhattan, yang dikendalikan oleh parameter p .

Perancangan

Pada Gambar 1 menunjukkan alur penelitian. Dimulai dari pembagian dataset 80 untuk pelatihan dan 20 untuk pengujian kemudian dilakukan resize pada gambar dengan ukuran 128 x 128 pixel dan dilakukan normalisasi selanjutnya di klasifikasikan menggunakan model KNN dan diukur menggunakan metrik evaluasi seperti akurasi, presisi, recall dan f1-score.



Gambar 1. Alur Penelitian

A. Dataset

Pada Penelitian ini menggunakan dataset publik dari website Kaggle yang terdiri dari 997 gambar dimana dibagi menjadi 2 kelas yaitu kelas kertas yang terdiri dari 594 gambar kertas dan 403 gambar karton. Data tersebut kemudian dibagi menjadi data pelatihan 80% dan pengujian 20%.

B. Preprocessing

1. Resize

Dalam proses preprocessing, gambar diubah ukurannya menjadi 128×128 piksel untuk memastikan konsistensi ukuran input sebelum dimasukkan ke dalam model [14].

2. Normalisasi

Normalisasi ukuran ini penting agar model dapat mengenali pola dengan lebih efektif dan mengurangi kompleksitas komputasi [15].

C. Klasifikasi dalam model KNN

Data yang telah diproses ini kemudian digunakan sebagai input untuk algoritma KNN, yang akan mengklasifikasikan gambar berdasarkan kesamaan dengan sampel data terdekat dalam ruang fitur [16],[17]. KNN bekerja dengan menghitung jarak, pada penelitian ini menggunakan Jarak Euclidean, Manhattan, dan Minkowski antara gambar yang diuji dan data pelatihan, kemudian menentukan kelas berdasarkan mayoritas dari K tetangga terdekat.

D. Evaluation

Berbagai metrik kinerja digunakan untuk mengevaluasi kemampuan model secara lebih mendalam. Metrik ini mencakup akurasi, yang menunjukkan seberapa sering model membuat prediksi yang benar; presisi, yang menilai ketepatan model dalam mengklasifikasikan suatu kelas tertentu serta F1 Score, yang merupakan kombinasi dari presisi dan recall untuk memberikan keseimbangan antara keduanya.

Untuk memberikan gambaran yang lebih jelas mengenai perhitungan metrik dalam evaluasi model, persamaan (2) - (5) digunakan sebagai rumus untuk menghitung F1-score, recall, akurasi, dan presisi.

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (2)$$

$$\text{Precision} = \frac{TP}{TP + FP} \quad (3)$$

$$\text{Recall} = \frac{TP}{TP + FN} \quad (4)$$

$$\text{F1-Score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (5)$$

True Positive (TP) adalah jumlah prediksi positif yang benar, sedangkan True Negative (TN) adalah jumlah prediksi negatif yang benar. False Positive (FP) mengacu pada jumlah prediksi positif yang salah, dan False Negative (FN) adalah jumlah prediksi negatif yang salah.

Pemodelan

Pada penelitian ini, model K-Nearest Neighbors (KNN) dievaluasi menggunakan tiga jenis metrik jarak, yaitu Euclidean, Manhattan, dan Minkowski. Evaluasi dilakukan berdasarkan accuracy, precision, recall, dan F1-score untuk menilai performa dari masing-masing metode jarak.

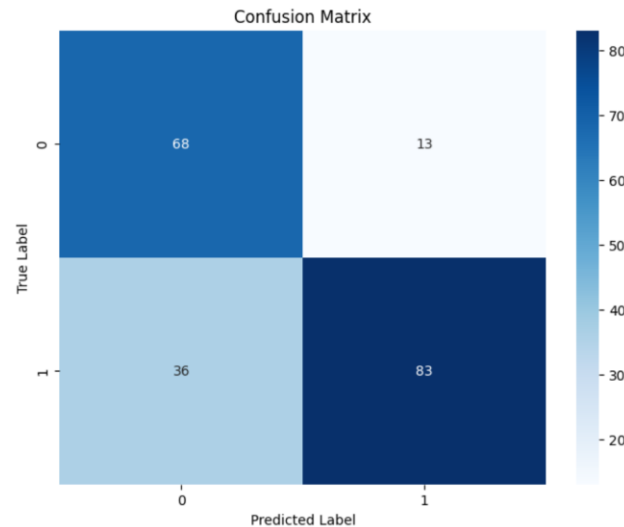
Hasil evaluasi ditampilkan dalam Tabel 1, yang menunjukkan bahwa penggunaan jarak Manhattan menghasilkan kinerja terbaik dibandingkan dengan dua metode lainnya, dengan accuracy sebesar 83%. Sementara itu, jarak Euclidean menghasilkan accuracy sebesar 75.50%, dan jarak Minkowski memberikan hasil paling rendah dengan accuracy sebesar 66.50%.

Selain itu, metrik precision, recall, dan F1-score juga menunjukkan tren yang serupa, di mana metode jarak Manhattan tetap unggul dibandingkan dengan dua metode lainnya. Hal ini menunjukkan bahwa pemilihan fungsi jarak dalam algoritma KNN memiliki dampak yang signifikan terhadap performa model.

Tabel 1. Hasil dari Model Algoritma KNN

Jarak	Accuracy	Precision	Recall	F1 - Score
Euclidean	75.50 %	78 %	76%	76 %
Manhattan	83.00 %	84%	83%	83%
Minkowski	66.50%	71%	67%	67%

Berdasarkan hasil evaluasi model KNN dengan jarak Euclidean, diperoleh bahwa model memiliki accuracy sebesar 75.50%, yang menunjukkan proporsi prediksi yang benar dari keseluruhan data. Precision model tercatat 78%, menunjukkan bahwa sebagian besar prediksi kertas adalah benar. Recall model tercatat 76%, yang berarti model mampu mengenali sebagian besar data kertas dengan baik, meskipun ada beberapa kesalahan dalam memprediksi data ini sebagai karton. Nilai F1-Score sebesar 76% menunjukkan keseimbangan antara precision dan recall, mengindikasikan performa model yang stabil dalam membedakan antara kelas kertas dan karton. Hasil ini menunjukkan bahwa jarak Euclidean memberikan performa yang cukup baik pada dataset yang digunakan. Gambar 2 menunjukkan confusion matrix pada model KNN dengan jarak euclidean.

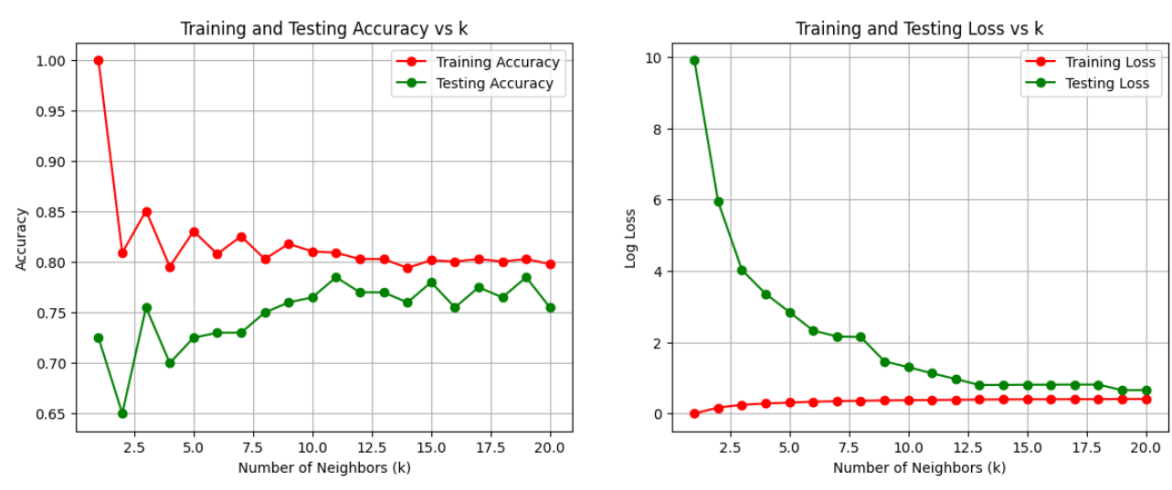


Gambar 2. *Confusion Matrix* dari Model KNN dengan jarak Eucledian

Gambar 3 menunjukkan hubungan antara jumlah tetangga terdekat K dengan akurasi serta loss pada model KNN. Pada grafik Akurasi terlihat bahwa akurasi pelatihan lebih tinggi dibandingkan dengan akurasi pengujian untuk berbagai nilai k. Ketika K bernilai kecil, model memiliki akurasi pelatihan yang sangat tinggi, mendekati 100%, tetapi akurasi pengujian lebih rendah. Hal ini menunjukkan bahwa model mengalami overfitting pada nilai K yang kecil. Seiring bertambahnya K, akurasi pengujian meningkat secara bertahap dan mencapai nilai yang lebih stabil, sementara akurasi pelatihan sedikit menurun.

Pada grafik loss, ditampilkan loss untuk pelatihan dan pengujian terhadap nilai K. Pada kecil, loss pengujian sangat tinggi, yang menunjukkan bahwa model tidak mampu menggeneralisasi dengan baik. Namun, seiring bertambahnya nilai K, loss pengujian menurun dan akhirnya mendekati loss pelatihan, menandakan bahwa model menjadi lebih stabil dan generalisasi meningkat.

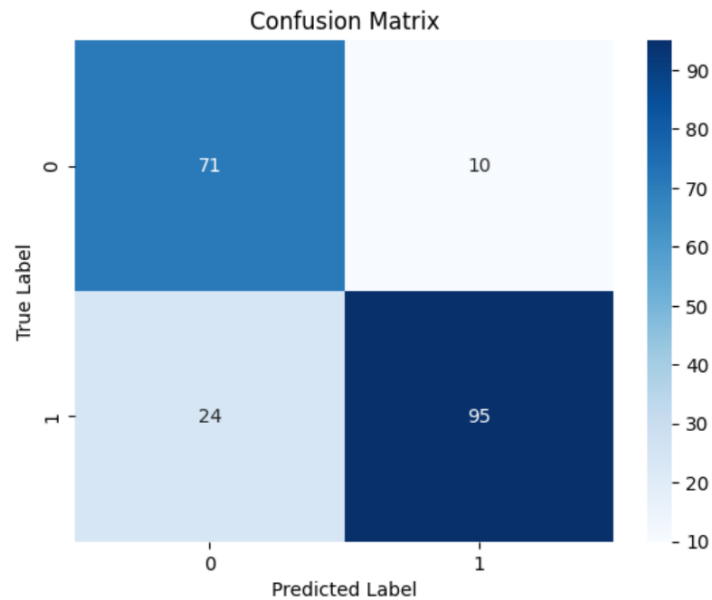
Dari kedua grafik ini, dapat disimpulkan bahwa pemilihan nilai k sangat berpengaruh terhadap performa model KNN. Nilai K yang terlalu kecil menyebabkan overfitting, sedangkan K yang lebih besar membantu model menjadi lebih generalisasi dengan akurasi lebih stabil dan loss lebih rendah.



Gambar 3. Grafik Accuracy dan Loss dari Model KNN dengan jarak euclidean

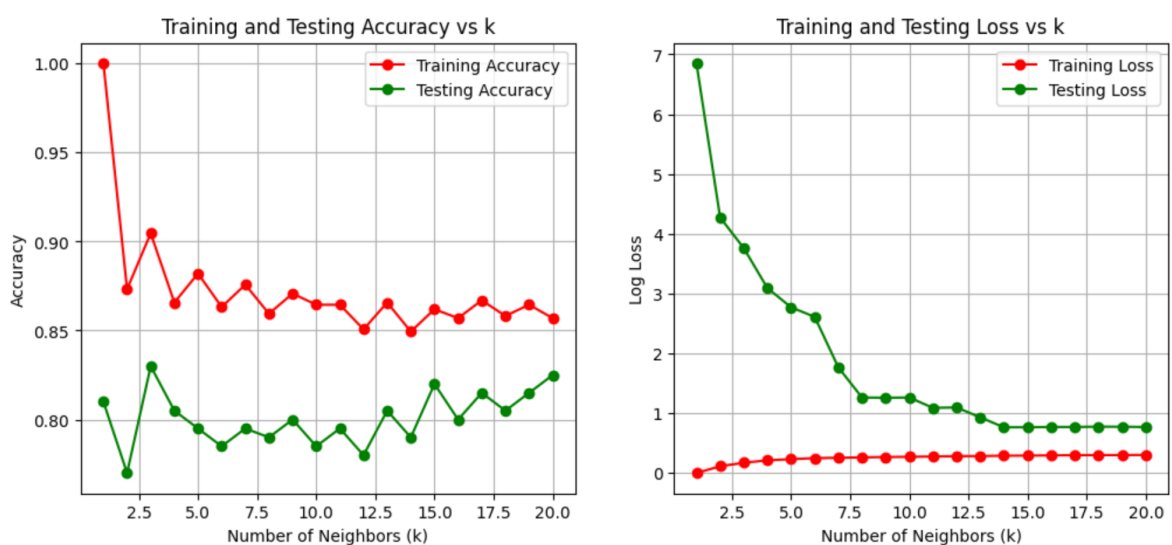
Berdasarkan Confusion Matrix untuk model KNN dengan jarak Manhattan, performa model sesuai dengan evaluasi di tabel. Model berhasil memprediksi 71 data karton dengan benar dan 95 data kertas dengan benar. Namun, terdapat 10 data karton yang salah diprediksi sebagai kertas dan 24 data kertas yang salah diprediksi sebagai karton. Dengan accuracy 83.00%, model menunjukkan proporsi prediksi yang benar secara

keseluruhan. Precision sebesar 84% menunjukkan bahwa sebagian besar prediksi kertas adalah benar, dan recall sebesar 83% menandakan kemampuan model yang cukup baik dalam mendeteksi data kertas. F1-Score sebesar 83% mengindikasikan keseimbangan antara precision dan recall, memperkuat bahwa jarak Manhattan memberikan performa terbaik pada model KNN ini. Gambar 4 Menunjukkan Confusion Matrix dari model KNN dengan jarak manhattan.

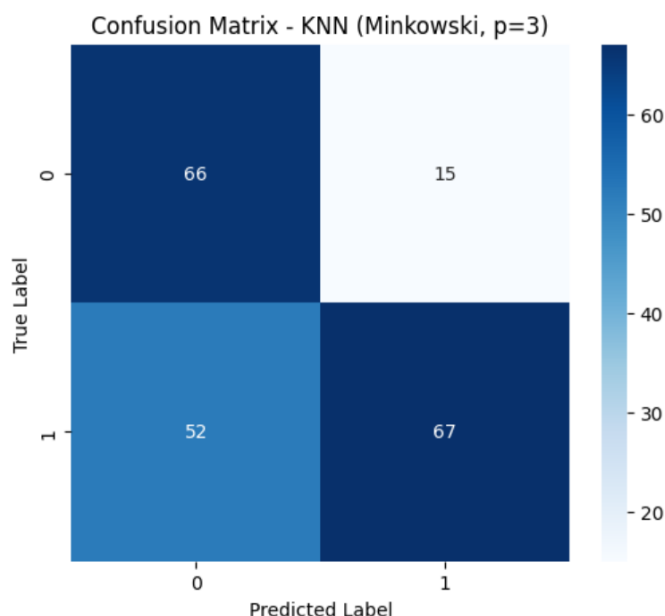


Gambar 4. Confusion matrix dari Model KNN dengan jarak euclidean

Gambar 5 menunjukkan grafik akurasi dan loss dari model KNN dengan jarak Manhattan berdasarkan jumlah tetangga K. Pada grafik akurasi, terlihat bahwa pada $k = 1$, akurasi training sangat tinggi mendekati 100%, sementara akurasi testing rendah, menunjukkan adanya overfitting. Seiring bertambahnya nilai K, akurasi testing meningkat dan stabil di sekitar 80–83%, mengindikasikan kemampuan model yang lebih baik dalam melakukan generalisasi. Pada grafik loss, testing loss awalnya tinggi pada $k = 1$ tetapi menurun tajam hingga stabil di sekitar $k > 10$, sedangkan training loss tetap rendah sepanjang grafik. Nilai K optimal berada pada rentang 10–15, di mana akurasi testing maksimal dan loss testing minimal tercapai, menunjukkan keseimbangan antara bias dan variansi model.



Gambar 5. Grafik Accuracy dan Loss dari Model KNN dengan jarak manhattan



Gambar 6. *Confusion matrix* dari Model KNN dengan jarak Minkowski

Berdasarkan Confusion Matrix untuk model KNN dengan jarak minkowski yang di tunjukkan pada Gambar 6, performa model sesuai dengan evaluasi di tabel. Model berhasil memprediksi 66 data karton dengan benar dan 67 data kertas dengan benar. Namun, terdapat 15 data karton yang salah diprediksi sebagai kertas dan 52 data kertas yang salah diprediksi sebagai karton. Dengan accuracy sebesar 66.50%, model menunjukkan proporsi prediksi yang benar secara keseluruhan. Precision sebesar 71% menunjukkan bahwa sebagian besar prediksi kertas adalah benar, dan recall sebesar 67% menandakan kemampuan model yang cukup baik dalam mendeteksi data kertas meskipun terdapat kesalahan yang signifikan. F1-Score sebesar 67% menunjukkan keseimbangan antara precision dan recall, namun performanya tidak sebaik jarak lainnya.

Kesimpulan

Penelitian ini membahas penerapan algoritma *K-Nearest Neighbors* (KNN) dalam klasifikasi sampah berbasis gambar menggunakan tiga metrik jarak, yaitu Euclidean, Manhattan, dan Minkowski. Hasil evaluasi menunjukkan bahwa jarak Manhattan memberikan performa terbaik dengan akurasi 83%, diikuti oleh Euclidean (75,50%) dan Minkowski (66,50%). Temuan ini mengindikasikan bahwa pemilihan fungsi jarak dalam KNN memiliki dampak signifikan terhadap akurasi klasifikasi.

Selain itu, penelitian ini menunjukkan bahwa dengan *preprocessing* yang tepat, seperti *resize* gambar dan normalisasi, performa model dapat ditingkatkan. Dari hasil analisis grafik akurasi dan *loss*, diketahui bahwa pemilihan nilai *K* yang optimal berperan penting dalam menghindari *overfitting* atau *underfitting*.

Dengan demikian, penelitian ini berkontribusi dalam pengembangan sistem klasifikasi sampah berbasis teknologi, yang dapat membantu meningkatkan efisiensi pengelolaan sampah dan mendukung upaya daur ulang yang lebih baik. Ke depan, integrasi dengan teknik *deep learning* atau pengujian dengan dataset yang lebih besar dapat menjadi langkah selanjutnya untuk meningkatkan akurasi dan kinerja sistem.

Daftar Pustaka

- [1] R. Kurniawan, P. Wintoro, Y. Mulyani, and M. Komarudin, "Implementasi Arsitektur Xception pada Model Machine Learning Klasifikasi Sampah Anorganik," *J. Inform. dan Tek. Elektro Terap.*, vol. 11, Apr. 2023, doi: 10.23960/jitet.v11i2.3034.
- [2] D. Damiri, A. Atika, S. Sodiah, and Y. Ramadhani, "Pemberdayaan Masyarakat melalui Pelatihan Pembuatan Pupuk Kompos dari Limbah Organik dan Rumah Tangga di Desa Perdamaian Kecamatan

- Singkut Kabupaten Sarolangun (Participatory Action Research),” *LOKOMOTIF ABDIMAS J. Pengabd. Kpd. Masy.*, vol. 2, Aug. 2024, doi: 10.30631/lokomotifabdimas.v2i2.2691.
- [3] Nurikah, E. R. Jajuli, H. and E. Furqon, “Waste Management Governance Based On Law Number 18 Of 2008 Of Waste Management Of Waste Based Citizen Participation In The Serang City,” *Gorontalo Law Rev.*, vol. 5, no. 2, pp. 434–442, 2022.
 - [4] S. Stocker, G. Csányi, K. Reuter, and J. Margraf, “Machine learning in chemical reaction space,” *Nat. Commun.*, vol. 11, Oct. 2020, doi: 10.1038/s41467-020-19267-x.
 - [5] H. Darwis, R. Puspitasari, Purnawansyah, W. Astuti, D. Atmajaya, and M. Hasnawi, “A Deep Learning Approach for Improving Waste Classification Accuracy with ResNet50 Feature Extraction,” in *2025 19th International Conference on Ubiquitous Information Management and Communication (IMCOM)*, 2025, pp. 1–8. doi: 10.1109/IMCOM64595.2025.10857536.
 - [6] P. Regina, P. Purnawansyah, T. Hasanuddin, and H. Darwis, “Klasifikasi Daun Herbal Menggunakan K-Nearest Neighbor dan Support Vector Machine dengan Fitur Fourier Descriptor,” *Edumatic J. Pendidik. Inform.*, vol. 7, pp. 160–168, Jun. 2023, doi: 10.29408/edumatic.v7i1.17521.
 - [7] M. Jumalis, M. Mirfan, and A. Manga, “Classification of Coffee Bean Defects Using Gray-Level Co-Occurrence Matrix and K-Nearest Neighbor,” *Ilk. J. Ilm.*, vol. 14, pp. 1–9, Apr. 2022, doi: 10.33096/ilkom.v14i1.910.1-9.
 - [8] L. Informatika, M. D. Taufiqurahman, S. Anraeni, H. Darwis, U. M. Indonesia, and K. Neighbor, “Analisis Sentimen Komentar Konten Kreator Gaming Menggunakan Metode Naive Bayes dan KNN,” vol. 1, no. 4, pp. 317–327, 2024.
 - [9] Herman, H. Darwis, Nurfauliyah, R. Puspitasari, D. Widyawati, and A. Faradibah, “Comparative Analysis of Anxiety Disorder Classification Using Algorithm Naive Bayes, Decision Tree and K-NN,” in *2025 19th International Conference on Ubiquitous Information Management and Communication (IMCOM)*, 2025, pp. 1–6. doi: 10.1109/IMCOM64595.2025.10857485.
 - [10] J. Arfah, Purnawansyah, H. Darwis, and R. Sastra, “Klasifikasi Penyakit Bawang Merah Menggunakan Naive Bayes dan CNN dengan Fitur GLCM,” *Indones. J. Comput. Sci.*, vol. 12, Jun. 2023, doi: 10.33022/ijcs.v12i3.3236.
 - [11] A. Manga, A. Utami, H. Azis, Y. Salim, and A. Faradibah, “Optimizing classification models for medical image diagnosis: a comparative analysis on multi-class datasets,” *Comput. Sci. Inf. Technol.*, vol. 5, pp. 205–214, Nov. 2024, doi: 10.11591/csit.v5i3.p205-214.
 - [12] A. Riska, Purnawansyah, H. Darwis, and W. Astuti, “Studi Perbandingan Kombinasi GMI, HSV, KNN, dan CNN pada Klasifikasi Daun Herbal,” *Indones. J. Comput. Sci.*, vol. 12, Jun. 2023, doi: 10.33022/ijcs.v12i3.3210.
 - [13] H. Azis, P. Purnawansyah, F. Fattah, and I. Putri, “Performa Klasifikasi K-NN dan Cross Validation Pada Data Pasien Pengidap Penyakit Jantung,” *Ilk. J. Ilm.*, vol. 12, pp. 81–86, Aug. 2020, doi: 10.33096/ilkom.v12i2.507.81-86.
 - [14] M. Samir, Purnawansyah, H. Darwis, and F. Umar, “Fourier Descriptor Pada Klasifikasi Daun Herbal Menggunakan Support Vector Machine Dan Naive Bayes,” *J. Teknol. Inf. dan Ilmu Komput.*, vol. 10, pp. 1205–1212, Dec. 2023, doi: 10.25126/jtiik.1067309.
 - [15] H. Darwis, Z. Ali, Y. Salim, and P. Belluano, “Max Feature Map CNN with Support Vector Guided Softmax for Face Recognition,” *JOIV Int. J. Informatics Vis.*, vol. 7, pp. 959–966, Sep. 2023, doi: 10.30630/joiv.7.3.1751.
 - [16] D. Zahirah, N. Kurniati, and H. Darwis, “Digital Image Classification of Herbal Leaves Using Knn and Cnn With Glcm Features,” *J. Tek. Inform.*, vol. 5, no. 1, pp. 61–67, 2024, [Online]. Available: <https://doi.org/10.52436/1.jutif.2024.5.1.1162>
 - [17] I. Putri, “Analisis Performa Metode K- Nearest Neighbor (KNN) dan Crossvalidation pada Data Penyakit Cardiovascular,” *Indones. J. Data Sci.*, vol. 2, pp. 21–28, Mar. 2021, doi: 10.33096/ijodas.v2i1.25.