

Penerapan Metode K-Means Clustering Dalam Pengelompokan Penyakit Pada Ayam Kampung Unggul Balibangtan

Iqbal Afriyadi^{*1}, Sarjon Defit², Sumijan³

^{1,2,3} Fakultas Ilmu Komputer, Universitas Putra Indonesia "YPTK" Padang

Email: ^{*1}iqbalafriyadi@gmail.com, ²sarjond@yahoo.co.uk, ³sumijan@upiyptk.ac.id

(Naskah masuk: 24 Februari 2025, diterima untuk diterbitkan: 20 Oktober 2025)

Abstrak: Penyakit ayam saat ini merupakan salah satu ancaman terbesar pada sebuah peternakan ayam. Penyakit pada ayam bisa disebabkan oleh virus dan bakteri. Ayam KUB merupakan salah satu jenis unggas yang dikembangkan oleh Badan Penelitian dan Pengembangan Pertanian Indonesia, dengan daya tahan tubuh yang baik dan produktivitas tinggi. Kendati demikian ayam KUB ini tetap rentan terhadap berbagai jenis penyakit yang dapat memengaruhi produktivitasnya. Pengelompokan penyakit pada ayam KUB penting untuk dilakukan guna mengidentifikasi pola serangan penyakit serta memberikan langkah preventif yang tepat bagi para peternak. Penelitian ini bertujuan untuk mengelompokkan penyakit yang menyerang Ayam Kampung Unggul Balitbangtan (KUB). Metode yang digunakan pada penelitian ini adalah penerapan machine learning dengan metode K-Means Clustering. Metode ini memiliki beberapa tahapan yaitu penyiapan data, normalisasi data, inisialisasi centroid, mengelompokkan data berdasarkan jarak terdekat, memperbarui centroid, iterasi sampai konvergensi, dan evaluasi hasil. Dataset yang diolah pada penelitian ini bersumber dari pengamatan langsung pada peternakan ayam ASA Farm Padang. Dataset yang digunakan dalam penelitian ini berjumlah 50 dataset yang berasal dari 50 ekor ayam KUB yang masuk kandang karantina pada peternakan tersebut. Pada penelitian ini menghasilkan kelompok penyakit ayam pada 3 kluster yaitu kluster 1 untuk ayam dengan penyakit gejala ringan dengan jumlah sebanyak 12 anggota, kluster 2 dengan penyakit gejala sedang dengan jumlah 14 anggota, dan kluster 3 dengan penyakit gejala tinggi sebanyak 24 anggota. Sehingga penelitian ini diharapkan dapat menjadi acuan bagi peternak, dokter hewan, peneliti selanjutnya atau pihak terkait dalam mengelompokkan penyakit pada ayam kampung atau hewan ternak lainnya..

Kata Kunci - Penyakit ayam, (KUB), pengelompokan penyakit, Machine learning, K-Means Clustering

Application of K-Means Clustering Method in Disease Grouping in Balibangtan Superior Native Chickens

Abstract: Chicken diseases are currently one of the biggest threats to a chicken farm. Diseases in chickens can be caused by viruses and bacteria. KUB chicken is one of the poultry breeds developed by the Indonesian Agricultural Research and Development Agency, with good immune system and high productivity. However, KUB chickens are still susceptible to various diseases that can affect their productivity. Clustering of diseases in KUB chickens is important to identify disease attack patterns and provide appropriate preventive measures for farmers. This study aims to classify diseases that attack Balitbangtan Superior Village Chicken (KUB). The method used in this research is the application of machine learning with the K-Means Clustering method. This method has several stages, namely data preparation, data normalisation, centroid initialisation, clustering data based on the closest distance, updating the centroid, iteration until convergence, and evaluating the results. The dataset processed in this study comes from direct observation at the ASA Farm Padang chicken farm. The dataset used in this study amounted to 50 datasets derived from 50 KUB chickens that entered the quarantine cage on the farm. This study produces a group of chicken diseases in 3 clusters, namely cluster 1 for chickens with mild symptom disease with a total of 12 members, cluster 2 with moderate symptom disease with a total of 14 members, and cluster 3 with high symptom disease with 24 members. So this research is expected to be a reference for farmers, veterinarians, further researchers or related parties in classifying diseases in native chickens or other livestock..

Keywords - Chicken disease, (KUB), disease clustering, Machine learning, K-Means Clustering

1. PENDAHULUAN

KDD adalah bidang penelitian multidisipliner yang berfokus pada upaya mengekstraksi pengetahuan tersembunyi yang sebelumnya tidak diketahui, namun memiliki potensi untuk digunakan, dari kumpulan data besar. Proses ini melibatkan pendekatan komputasi dan analitis untuk mengungkap wawasan baru dari data[1]. Data mining adalah proses menggali pola atau informasi berharga dari kumpulan data besar. Informasi ini biasanya tidak teridentifikasi secara manual dan dapat digunakan untuk menghasilkan pengetahuan baru[2]. Clustering merupakan salah satu metode data mining yang bersifat tanpa arahan (unsupervised), maksudnya metode ini diterapkan tanpa adanya latihan (training) dan tanpa ada guru (teacher) serta tidak memerlukan target output[3]. K-Means Clustering adalah algoritma analisis kluster yang diselesaikan secara iteratif. Langkah-langkahnya dimulai dengan memilih secara acak K objek sebagai pusat kluster awal, kemudian menghitung jarak antara setiap objek dan pusat kluster awal tersebut. Setelah itu, setiap objek akan ditugaskan ke pusat kluster terdekat[4].

Ayam kampung unggul Balitbangtan sering menghadapi berbagai jenis penyakit yang memengaruhi produktivitas dan kelangsungan hidup mereka. Masalah utamanya adalah sulitnya mengidentifikasi jenis penyakit pada ayam, pola penyakit secara cepat dan akurat karena banyaknya faktor yang memengaruhi, seperti gejala klinis, lingkungan, dan pola penyebaran penyakit[5]. Data yang besar dan kompleks ini membuat analisis manual menjadi tidak efisien dan rentan terhadap kesalahan. K-Means Clustering dapat membantu mengelompokkan data penyakit berdasarkan kesamaan gejala, pola penyebaran, atau faktor lingkungan[6]. Menggunakan algoritma ini data penyakit ayam dapat dibagi ke dalam beberapa cluster, di mana setiap cluster mewakili kategori tertentu, seperti jenis penyakit atau tingkat keparahan[7]. Hal ini memungkinkan identifikasi pola penyakit yang lebih cepat dan akurat, sehingga membantu peternak dan peneliti dalam mengambil tindakan pencegahan atau pengobatan yang tepat.

Pada penelitian Prasetyo dkk. berhasil menggunakan kombinasi metode K-Means dan Naïve Bayes untuk menganalisis sentimen di Twitter terkait bantuan bencana alam, mencapai akurasi hingga 87% dalam mengidentifikasi jenis bantuan yang diperlukan[8]. Nugroho dkk. menggunakan algoritma K-Means untuk klasifikasi risiko usaha, dengan validasi silhouette menghasilkan nilai 0,65 dan 0,93 pada variabel terkait[9]. Kurniadi dkk. menggunakan algoritma K-Means untuk mengelompokkan pembangunan jalan, menghasilkan lima kluster dengan indeks Davies-Bouldin terkecil sebesar 0,1617[10]. Marini dan Suhendra menggunakan algoritma K-Means untuk memetakan kluster daerah wisata, menghasilkan dua kluster berdasarkan jumlah kunjungan wisata[11]. Praseptian M. dkk. menggunakan K-Means untuk mengelompokkan tingkat kepuasan pengguna lulusan perguruan tinggi, menunjukkan 94,12% pengguna merasa sangat puas[12].

Ahmad Januar Amriyansah dkk. mengembangkan sistem diagnosis penyakit ayam menggunakan metode Forward Chaining, yang menghasilkan diagnosa akurat untuk 33 gejala dan 10 penyakit[13]. Menurut penelitian yang dilakukan oleh Sefto Pratama dan Agus Alim Muin menerapkan Certainty Factor untuk diagnosis penyakit ayam, mencapai akurasi 85%[14]. Alfonsus Eksi Adi Irawan dkk. mengembangkan sistem diagnosis berbasis Forward Chaining dan Certainty Factor dengan akurasi 92% dan tingkat kepuasan tinggi dari pengguna[15]. Nani Mintarsih dkk. menggunakan CNN DenseNet121 untuk mendiagnosis penyakit ayam melalui citra, mencapai akurasi sebesar 95%[16].

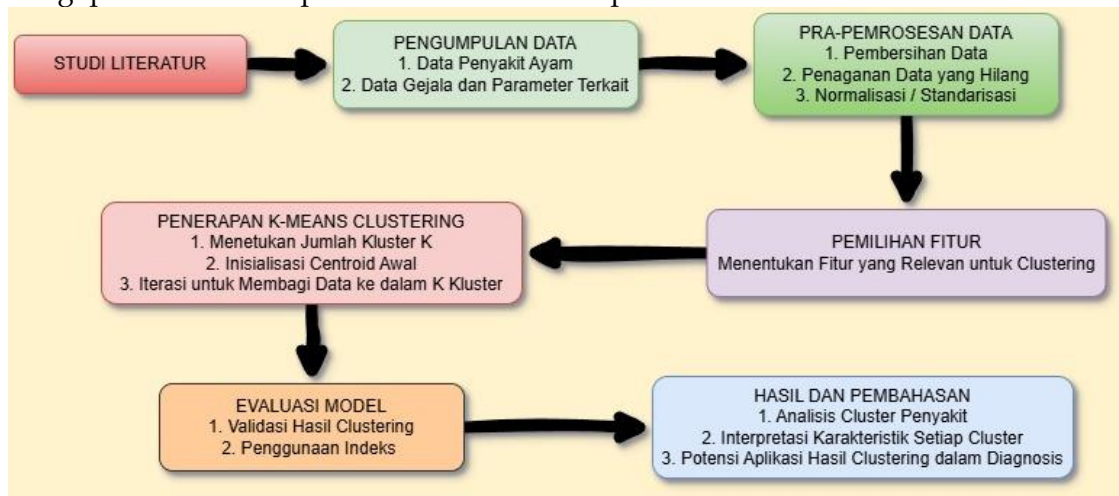
Sebagai salah satu tantangan utama dalam bidang peternakan, pengelompokan penyakit pada hewan berdasarkan gejala dan karakteristiknya merupakan langkah penting dalam meningkatkan upaya pencegahan dan penanganan yang tepat. Penerapan algoritma K-Means Clustering untuk mengelompokkan penyakit pada hewan berbasis machine learning memiliki keunggulan karena fokusnya dalam memberikan identifikasi penyakit yang lebih akurat, sehingga memungkinkan penanganan yang lebih efektif dan terarah. Penelitian ini bertujuan untuk memanfaatkan K-Means Clustering dalam pengelompokan penyakit berdasarkan data hewan yang terinfeksi, yang akan membantu peternak, dokter hewan, dan pihak terkait dalam mengenali pola penyakit yang sering terjadi serta menerapkan langkah penanganan yang optimal. Penerapan algoritma K-Means

Clustering diharapkan mampu meningkatkan efisiensi dalam pengelolaan data penyakit pada hewan ternak. Dengan sistem pengelompokan yang efektif, para peternak dapat mengidentifikasi penyakit lebih cepat dan melakukan tindakan pencegahan yang lebih baik, sehingga produktivitas peternakan dapat meningkat.

Berdasarkan latar belakang yang telah diuraikan, tujuan penelitian ini adalah untuk menerapkan algoritma K-Means Clustering dalam pengelompokan penyakit pada hewan berbasis machine learning. Selain itu, penelitian ini juga bertujuan untuk membantu para tenaga kesehatan hewan dan peternak dalam meningkatkan kualitas pengelolaan kesehatan hewan melalui pengelompokan data yang lebih efektif dan efisien.

2. METODE PENELITIAN

Pendekatan metodologi dalam penelitian ini berfokus pada penggunaan algoritma K-Means Clustering untuk mengelompokkan data penyakit hewan yang telah dikumpulkan. Proses metodologi dimulai dengan tahap pengumpulan dan persiapan data, termasuk pembersihan dan normalisasi data untuk memastikan kualitas data yang digunakan. Setelah data selesai di normalisasi, algoritma K-Means Clustering diterapkan untuk melakukan pengelompokan penyakit hewan berdasarkan karakteristik yang ada, dengan menggunakan tools seperti WEKA. Tahapan akhir penelitian melibatkan pengujian hasil clustering dan analisis untuk menilai efektivitas serta akurasi metode K-Means dalam mengelompokkan penyakit hewan. Metodologi ini digunakan sebagai kerangka kerja untuk memperoleh data yang akurat dan relevan dalam menyelesaikan masalah pengelompokan penyakit hewan secara efektif. Hasil pengelompokan diharapkan memberikan pemahaman yang lebih baik dalam penanganan penyakit pada hewan ternak. Tahapan metodologi penelitian ini dapat dilihat secara detail pada Gambar 1.



Gambar 1. kerangka penelitian

Kerangka penelitian diatas menjelaskan proses kerja penelitian yang menggambarkan tahapan-tahapan utama dalam penelitian ini. Setiap tahapan mulai dari studi literatur hingga analisa hasil dan pembahasan disusun secara sistematis untuk mencapai tujuan penelitian. Hasil dari perhitungan K-Means Clustering ini nanti akan di evaluasi menggunakan alat bantu WEKA untuk mengetahui apakah hasil yang diperoleh sesuai. Kerangka kerja ini memastikan setiap langkah yang dilakukan dalam proses penelitian ini sesuai dan saling berhubungan serta saling mendukung keseluruhan proses. Tahapan pada kerangka kerja sebagai berikut:

2.1. Studi Literatur

Literatur yang digunakan dalam penelitian ini mencakup jurnal-jurnal ilmiah yang membahas penerapan algoritma K-Means Clustering dalam pengelompokan penyakit pada hewan, serta artikel pendukung lainnya sebagai referensi dalam mendukung proses penelitian ini. Studi literatur

dilakukan untuk memahami konsep dasar dan teori yang melandasi penggunaan algoritma K-Means Clustering. Dengan mengkaji berbagai sumber ilmiah, penelitian terdahulu, dan referensi terkait, peneliti dapat menentukan metode dan teknik yang paling relevan untuk diterapkan dalam penelitian ini. Hasil dari studi literatur ini memberikan landasan yang kokoh dalam merancang kerangka kerja penelitian dan menentukan tahapan yang perlu dilakukan dalam proses pengelompokan penyakit pada hewan.

2.2. Pengumpulan Data

Data yang dikumpulkan terkait penyakit-penyakit yang menyerang ayam, seperti gejala, jenis penyakit, serta hasil diagnosis dari para ahli kesehatan hewan. Selain data utama berupa penyakit, data tambahan seperti gejala yang muncul, faktor lingkungan, dan kondisi fisik ayam juga dikumpulkan untuk memperkaya proses analisis.

2.3. Pra-pemrosesan Data dan Pemilihan Fitur

Tahap ini melibatkan pembersihan data dari kesalahan atau inkonsistensi, misalnya kesalahan pencatatan atau duplikasi data. Penanganan data yang hilang (missing values) dilakukan dengan cara mengisi nilai yang hilang atau menghapus data yang tidak lengkap. Data kemudian dinormalisasi atau distandarisasi agar semua variabel memiliki skala yang sama, sehingga tidak ada fitur yang terlalu mendominasi proses clustering. Pada tahap ini, peneliti memilih fitur-fitur yang relevan dan penting untuk diikutsertakan dalam proses clustering. Pemilihan fitur dilakukan untuk memastikan bahwa variabel yang digunakan berkontribusi terhadap akurasi hasil clustering

2.4. Penerapan K-Means Clustering

1. Menentukan Jumlah Kluster K:

Pada tahap ini peneliti menentukan jumlah kluster (K) yang akan digunakan untuk mengelompokkan data. Jumlah kluster ini biasanya ditentukan melalui metode seperti Elbow Method. Rumus yang digunakan tersaji pada persamaan 1.

$$SSE = \sum_{k=1}^K \sum_{x_i \in C_k} \|x_i - \mu_k\|^2 \quad (1)$$

2. Inisialisasi Centroid Awal dan Menghitung jarak antara titik centroid dengan titik tiap objek:

Menentukan centroid awal secara acak adalah langkah pertama dalam proses K-Means clustering yang sangat berpengaruh terhadap hasil akhir dari pengelompokan data. Centroid awal ini dipilih sebagai titik acuan awal untuk memulai pengelompokan data ke dalam kluster yang berbeda. Setelah centroid awal dipilih, algoritma secara iteratif akan memperbarui posisi centroid berdasarkan pola data hingga mencapai titik konvergensi, yaitu saat tidak ada lagi perubahan signifikan dalam pembagian kluster. Hal ini bertujuan agar data dikelompokkan secara optimal sesuai pola distribusi yang ada dalam dataset. Penentuan centroid awal bisa dilakukan dengan memilih titik acak dari sampel data yang digunakan.

Dalam tahap ini, dilakukan perhitungan jarak antara setiap titik data dan pusat kluster terdekat. Jarak tersebut dihitung menggunakan rumus Euclidean Distance, yang mengukur seberapa dekat data dengan centroid kluster yang telah ditetapkan. Dengan menggunakan rumus Euclidean Distance seperti pada Persamaan 2.

$$De = \sqrt{(x_i - s_i)^2 + (y_i - t_i)^2} \quad (2)$$

Data akan dikelompokkan ke kluster dengan centroid yang memiliki jarak terdekat.

3. Iterasi untuk Membagi Data ke dalam K Kluster:

Algoritma K-Means kemudian melakukan iterasi dengan membagi data ke dalam kluster berdasarkan jarak terdekat dari setiap centroid, dan centroid diperbarui hingga tercapai konvergensi.

2.5. Evaluasi Model

Setelah proses clustering selesai, hasil kluster yang terbentuk perlu divalidasi untuk memastikan keakuratan dan relevansi pengelompokan. Metode validasi yang digunakan pada penelitian ini adalah dengan aplikasi WEKA untuk menilai kualitas clustering.

2.6. Hasil dan Pembahasan

Hasil clustering pada penyakit ayam memberikan wawasan mendalam mengenai pola dan karakteristik tiap kelompok berdasarkan gejala yang diamati. Analisis cluster penyakit menunjukkan bahwa setiap cluster memiliki kombinasi gejala yang khas, yang dapat mengarah pada indikasi penyakit tertentu. Dengan memahami pola ini, interpretasi karakteristik tiap cluster dapat dilakukan untuk mengidentifikasi hubungan antara gejala yang dominan, misalnya apakah suatu cluster lebih banyak menunjukkan gejala mata bengkak dan sayap terkulai, yang dapat mengarah pada penyakit tertentu. Lebih lanjut, potensi aplikasi hasil clustering ini sangat berguna dalam membantu proses diagnosis dini secara lebih akurat, memungkinkan peternak atau dokter hewan mengambil keputusan yang lebih tepat dalam penanganan penyakit. Dengan demikian, metode clustering ini dapat diintegrasikan ke dalam sistem pemantauan kesehatan ayam guna meningkatkan efektivitas pengobatan dan langkah-langkah pencegahan penyakit

3. HASIL DAN PEMBAHASAN

3.1. Pengumpulan Data

Berdasarkan studi literatur yang dilakukan di dapat gejala -gejala penyakit pada ayam yang sering di jumpai. Pada penelitian ini dipilih beberapa fitur yang akan digunakan sebagai parameter pengukuran sebagai berikut:

Tabel 1. Data Gejala Penyakit Ayam

No.	Gejala
1	Nafsu makan menurun (G1)
2	Mata bengkak (G2)
3	Jengger dan pial pucat (G3)
4	Keluar cairan dari mata dan hidung (G4)
5	Bintik merah pada kaki (G5)
6	Warna kulit berubah biru keunguan (G6)
7	Mata lesu dan mengantuk (G7)
8	Sayap terkulai (G8)

Berikut dilampirkan dataset dari 50 data yang diperoleh di lapangan setelah dilakukan proses normalisasi data.

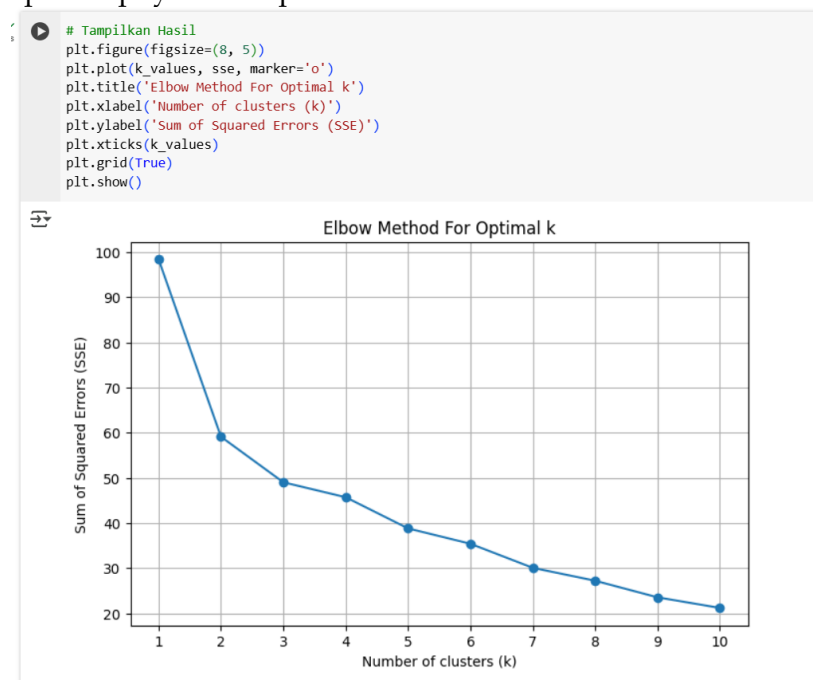
Tabel 2. Dataset Penyakit Ayam

No.	Nama Id	G1	G2	G3	G4	G5	G6	G7	G8
1	Data 1	1	0	1	0	1	0	1	1
2	Data 2	1	1	1	1	0	1	1	1
3	Data 3	1	1	1	0	1	0	1	1
4	Data 4	1	0	1	1	0	1	1	1
5	Data 5	1	1	1	1	1	0	1	1
6	Data 6	1	0	1	0	0	1	1	1
7	Data 7	1	1	1	1	1	1	1	1
8	Data 8	1	1	1	0	1	0	1	1
9	Data 9	1	0	1	1	0	1	1	1
10	Data 10	1	1	1	0	1	0	1	1
.....									
45	Data 45	1	0	1	1	0	1	1	1
46	Data 46	1	1	1	0	1	0	1	1
47	Data 47	1	1	1	1	1	1	1	1
48	Data 48	1	0	1	0	0	0	1	1
49	Data 49	1	0	1	1	0	1	1	1
50	Data 50	1	1	1	0	1	0	1	1

3.2. Penerapan K-Means Clustering

1. Menentukan Jumlah Kluster

Menentukan jumlah kluster dalam algoritma K-Means Clustering adalah langkah kritis yang berpengaruh pada akurasi dan keandalan hasil pengelompokan data. Proses ini memerlukan pertimbangan yang cermat untuk memastikan bahwa jumlah kluster yang dipilih dapat menangkap variasi dan pola penting dalam data. Elbow Method digunakan untuk menentukan jumlah kluster yang optimal, dengan tujuan mencapai keseimbangan antara kompleksitas model dan jumlah variasi data yang dapat dijelaskan. Berdasarkan perhitungan menggunakan Persamaan 1 pada aplikasi phyton didapatkan hasil berikut:



Gambar 2. Hasil Perhitungan Elbow Methode

Berdasarkan grafik diatas jumlah cluster yang akan diterapkan pada penelitian ini berjumlah 3 dengan penjelasan kluster 1 (C1) kelompok ayam dengan gejala ringan, kluster 2 (C2) kelompok ayam dengan gejala sedang dan kluster 3 (C3) kelompok ayam dengan gejala berat.

2. Inisiasi Centroid Awal

Penentuan nilai centroid secara acak merupakan langkah awal dalam algoritma K-Means clustering yang menjadi dasar untuk mengelompokkan data ke dalam cluster yang berbeda. Nilai centroid ini diambil dari sampel data yang telah tersedia dan digunakan sebagai titik referensi dalam proses iterasi untuk memperbarui posisi centroid sampai konvergensi tercapai. Pemilihan nilai centroid secara acak memungkinkan algoritma untuk mengeksplorasi berbagai kemungkinan pengelompokan data dengan lebih efektif. Dalam penelitian ini, nilai centroid awal diambil secara acak dari 50 data penyakit pada ayam KUB yang akan dianalisis. Centroid untuk C1 diambil dari data sampel ke-3 (1,1,1,0,1,0,1,1) C2 diambil dari data sampel ke-14 (1,1,1,0,1,0,1,1) dan untuk C3 diambil dari data sampel ke-26 (1,0,1,0,0,1,1,1).

3. Iterasi Untuk Membagi Data Kedalam K Kluster

Pada tahap ini perhitungan Euclidean Distance digunakan untuk mengukur jarak antara dua titik data dalam ruang multidimensi. Jarak ini dihitung dengan cara mengkuadratkan selisih setiap pasangan koordinat, kemudian menjumlahkan hasilnya, dan mengambil akar kuadrat dari jumlah tersebut. Euclidean Distance sangat berguna dalam K-Means clustering untuk menentukan seberapa dekat data berada dengan centroid yang telah ditentukan, sehingga data dapat dikelompokkan ke dalam cluster yang sesuai. Pada tahap ini, perhitungan Euclidean Distance dilakukan menggunakan persamaan yang telah disediakan, dan hasil perhitungan tersebut ditampilkan dalam Tabel 3.

Tabel 3. Hasil Perhitungan Iterasi 1

No.	Nama Id	C1	C2	C3	Jarak Terdekat	Cluster
1	Data 1	1,000	1,000	1,414	1,000	C1
2	Data 2	1,732	1,732	1,414	1,414	C3
3	Data 3	0,000	0,000	1,732	0,000	C1
4	Data 4	2,000	2,000	1,000	1,000	C3
5	Data 5	1,000	1,000	2,000	1,000	C1
6	Data 6	1,732	1,732	0,000	0,000	C3
7	Data 7	1,414	1,414	1,732	1,414	C1
8	Data 8	0,000	0,000	1,732	0,000	C1
9	Data 9	2,000	2,000	1,000	1,000	C3
10	Data 10	0,000	0,000	1,732	0,000	C1
...	...	...	...	...	...	...
45	Data 45	2,000	2,000	1,000	1,000	C3
46	Data 46	0,000	0,000	1,732	0,000	C1
47	Data 47	1,414	1,414	1,732	1,414	C1
48	Data 48	1,414	1,414	1,000	1,000	C3
49	Data 49	2,000	2,000	1,000	1,000	C3
50	Data 50	0,000	0,000	1,732	0,000	C1

Tabel 3 menampilkan hasil iterasi pertama dari proses K-Means clustering, di mana setiap data objek dikelompokkan ke dalam cluster sesuai dengan jarak terdekatnya dengan centroid yang telah ditentukan. Pada iterasi pertama ini, terdapat tiga cluster, yaitu C1, C2 dan C1. Setiap objek data dianalisis berdasarkan jarak ke centroid, dan di setiap cluster diidentifikasi objek-objek yang memiliki jarak terdekat dengan centroid masing-masing. Hasil pada iterasi pertama menunjukkan bahwa cluster C1 memiliki 26 anggota, cluster C2 memiliki 0 anggota sedangkan cluster C3 memiliki 24 anggota. Setelah iterasi pertama titik centroid baru dihitung dan diperbarui. Hasil perhitungan centroid baru dapat dilihat pada Tabel 4 di bawah ini.

Tabel 4 Hasil Perhitungan Iterasi 1

Cluster	Jumlah Anggota	G1	G2	G3	G4	G5	G6	G7	G8
C1	26	1,000	0,539	1,000	0,039	0,577	0,039	1,000	1,000
C2	0	1,000	1,000	1,000	0,000	1,000	0,000	1,000	1,000
C3	24	1,000	0,500	1,000	1,000	0,458	1,000	1,000	1,000

Berdasarkan hasil perhitungan pada iterasi pertama didapatkan sebaran anggota tiap2 cluster sebagai berikut:

Tabel 5 Sebaran anggota tiap cluster hasil iterasi 1

Cluster	Jumlah Anggota	Anggota Kelompok
C1	26	1,3,5,7,8,10,12,14,15,17,19,22,23,26,27 30,31,34,35,38,39,42,43,46,47,50
C2	0	
C3	24	2,4,6,9,11,13,16,18,20,21,24,25,28,29 32,33,36,37,40,41,44,45,48,49

Proses iterasi K-Means clustering dilanjutkan dengan mengulangi langkah-langkah pengelompokan dan pembaruan centroid hingga posisi centroid tidak mengalami perubahan, dan anggota cluster tetap berada dalam cluster yang sama. Ketika tidak ada lagi perpindahan anggota antara cluster dan nilai centroid telah stabil, proses clustering dianggap mencapai konvergensi. Stabilitas ini menunjukkan bahwa algoritma telah berhasil menemukan pengelompokan data yang optimal sesuai dengan pola yang ada dalam dataset.

Pada penelitian ini, proses iterasi berhenti pada iterasi ke-3. Hal ini disebabkan karena pada iterasi ke-3 tidak terjadi lagi perubahan nilai centroid dengan iterasi sebelumnya, sehingga proses algoritma K-Means dianggap telah selesai. Hasil final dari pengelompokan ini dapat dilihat pada Tabel 6.

Tabel 6 Hasil Perhitungan Iterasi 5

No.	Nama Id	C1	C2	C3	Jarak Terdekat	Cluster
1	Data 1	0,91	1,00	1,57	0,91	C1
2	Data 2	1,73	1,69	0,68	0,68	C3
3	Data 3	1,35	0,07	1,58	0,07	C2
4	Data 4	1,42	1,97	0,65	0,65	C3
5	Data 5	1,68	0,93	1,26	0,93	C2
6	Data 6	1,00	1,73	1,16	1,00	C1
7	Data 7	1,96	1,36	0,77	0,77	C3
8	Data 8	1,35	0,07	1,58	0,07	C2
9	Data 9	1,42	1,97	0,65	0,65	C3
10	Data 10	1,35	0,07	1,58	0,07	C2
...	...	...	...	...	...	...
45	Data 45	1,42	1,97	0,65	0,65	C3
46	Data 46	1,35	0,07	1,58	0,07	C2
47	Data 47	1,96	1,36	0,77	0,77	C3
48	Data 48	0,09	1,42	1,53	0,09	C1
49	Data 49	1,42	1,97	0,65	0,65	C3
50	Data 50	1,35	0,07	1,58	0,07	C2

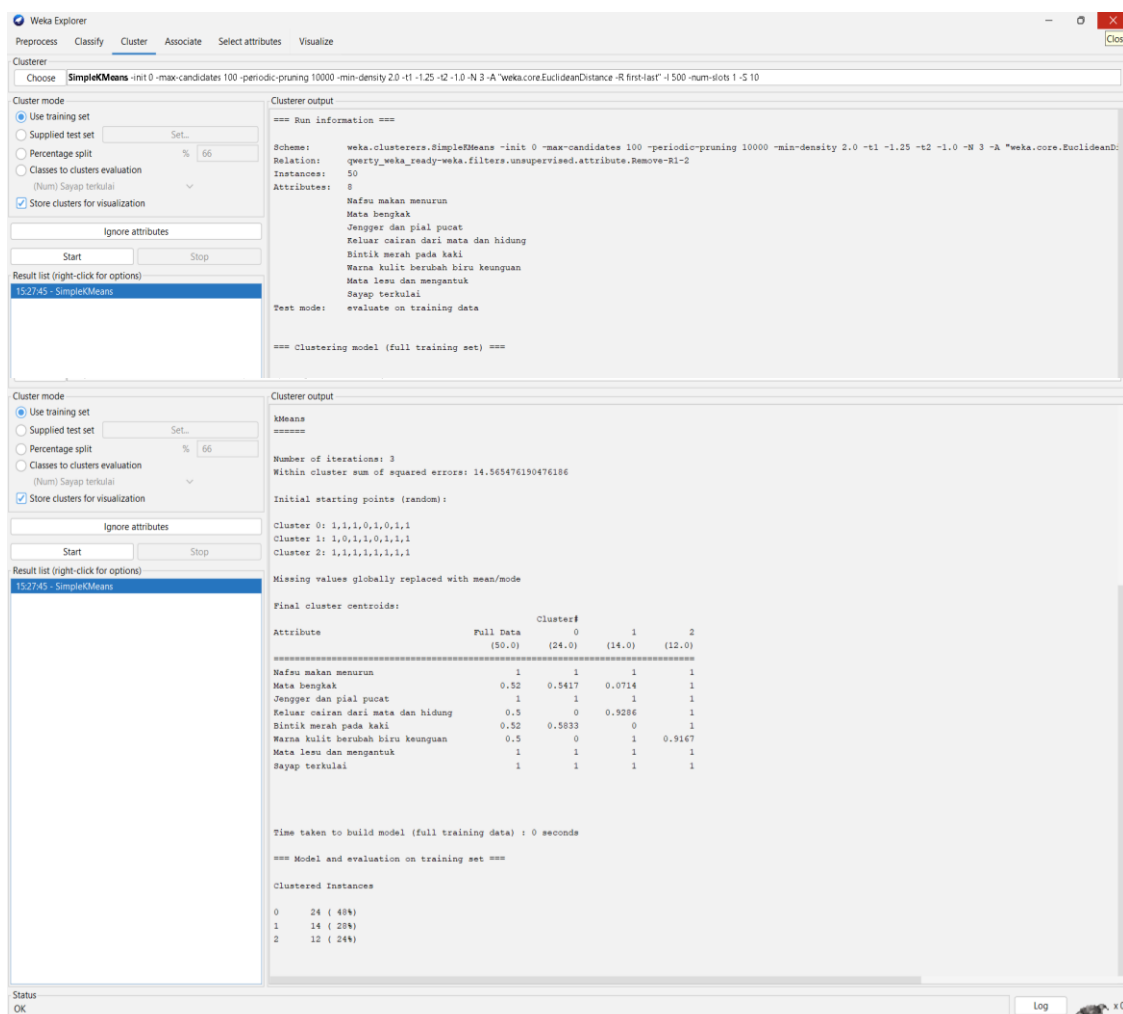
Tabel 7 Sebaran anggota tiap cluster hasil iterasi 3

Cluster	Jumlah Anggota	Anggota Kelompok
C1	12	1,6,16,18,20,24,28,32,36,40,44,48
C2	14	3,5,8,10,12,14,22,26,30,34,38,42,46,50
C3	24	2,4,7,9,11,13,15,17,19,21,23,25,27,29,31 33,35,37,39,41,43,45,47,49

Tabel di atas menunjukkan hasil clustering dengan tiga kelompok. Cluster C1 terdiri dari 12 anggota, yaitu 1, 6, 16, 18, 20, 24, 28, 32, 36, 40, 44, dan 48. Cluster C2 memiliki 14 anggota, yang terdiri dari 3, 5, 8, 10, 12, 14, 22, 26, 30, 34, 38, 42, 46, dan 50. Sementara itu, Cluster C3 mencakup 24 anggota, yakni 2, 4, 7, 9, 11, 13, 15, 17, 19, 21, 23, 25, 27, 29, 31, 33, 35, 37, 39, 41, 43, 45, 47, dan 49. Hasil clustering ini menunjukkan bahwa data telah berhasil dikelompokkan ke dalam tiga kelompok yang berbeda, di mana masing-masing kelompok mencerminkan pola atau hubungan yang ditemukan dalam data.

3.3. Evaluasi Hasil

Pada penelitian ini evaluasi hasil perhitungan pada tahapan sebelumnya di lakukan dengan menggunakan aplikasi WEKA. Hasil dari evaluasi hasil perhitungan dengan penggunaan aplikasi WEKA dapat dilihat dalam Gambar 2 berikut ini:



Gambar 3. Hasil Klasterisasi menggunakan aplikasi WEKA

Berdasarkan hasil evaluasi penggunaan aplikasi weka dapat dilihat bahwasanya jumlah anggota setiap kluster sama dengan hasil perhitungan dengan cara sebelumnya dimana anggota C1 berjumlah 12 anggota, cluster 2 berjumlah 14 anggota dan cluster C3 berjumlah 24 anggota.

4. KESIMPULAN

Hasil klasterisasi menunjukkan tiga pola berbeda dari kondisi kesehatan ayam. Cluster 1 (12 data) ditandai dengan nafsu makan menurun, mata bengkak sebagian (54.17%), jengger pucat, bintik merah pada kaki (58.33%), mata lesu, sayap terkulai, tanpa cairan dari mata/hidung dan kulit tidak berubah menjadi biru keunguan. Cluster 2 (14 data) menunjukkan penurunan nafsu makan, jengger pucat, keluarnya cairan dari mata/hidung (92.86%), perubahan warna kulit menjadi biru keunguan (100%), mata lesu, dan sayap terkulai, tanpa bintik merah pada kaki. Cluster 3 (24 data) melibatkan hampir semua gejala, dengan 100% mata bengkak, 100% cairan dari mata/hidung, 100% bintik merah pada kaki, dan 91.67% perubahan warna kulit menjadi biru keunguan. Pola ini dapat digunakan untuk diagnosis atau pengelompokan ayam berdasarkan gejala yang diamati.

UCAPAN TERIMA KASIH

Saya mengucapkan terima kasih kepada Peternakan Asa Ternak Mandiri atas dukungan dan kerja samanya selama penelitian ini. Bantuan yang diberikan, baik berupa akses informasi maupun fasilitas, sangat membantu kelancaran penelitian. Semoga hasil dari penelitian ini dapat memberikan manfaat bagi peternakan dan mendukung keberlanjutan sektor peternakan di masa depan. Terima kasih atas segala bantuan yang telah diberikan.

DAFTAR PUSTAKA

- [1] Tsui, K. L., Chen, V., Jiang, W., Yang, F., & Kan, C. (2023). Data Mining Methods and Applications. *Springer Handbooks*, 797–816. https://doi.org/10.1007/978-1-4471-7503-2_38
- [2] Apriyani, P., Dikananda, A. R., & Ali, I. (2023). Penerapan Algoritma K-Means dalam Klasterisasi Kasus Stunting Balita Desa Tegalwangi. *Hello World Jurnal Ilmu Komputer*, 2(1), 20–33. <https://doi.org/10.56211/helloworld.v2i1.230>
- [3] Safitri, E. (2024). Implementasi Algoritma K-Means Clustering Dalam Menentukan Strategi Marketing Dalam Penjualan Ikan (Studi Kasus: Grosir Ikan Tani Mas Tanjung Morawa). *JIKTEKS : Jurnal Ilmu Komputer Dan Teknologi Informasi*, 02(02), 23–33.
- [4] Huang, Z., Zheng, H., Li, C., & Che, C. (2024). Application of Machine Learning-Based K-means Clustering for Financial Fraud Detection. *Academic Journal of Science and Technology*, 10(1), 33–39. <https://doi.org/10.54097/74414c90>
- [5] Firmansyah, M., & Susanto, E. S. (2024). Implementasi Sistem Pakar Berbasis Web Untuk Mendiagnosis Penyakit Ayam Broiler Menggunakan Metode Forward Chaining. *Jeis: Jurnal Elektro Dan Informatika Swadharma*, 4(1), 10–16. <https://doi.org/10.56486/jeis.vol4no1.435>
- [6] Ferdy Pangestu, F. P., Nur Yasin, N. Y., Ronald Chistover Hasugian, R. C., & Yunita, Y. (2023). Penerapan Algoritma K-Means Untuk Mengklasifikasi Data Obat. *Jurnal Sisfokom (Sistem Informasi Dan Komputer)*, 12(1), 53–62. <https://doi.org/10.32736/sisfokom.v12i1.1461>
- [7] Pamungkas, L., Dewi, N. A., & Putri, N. A. (2024). Classification of Student Grade Data Using the K-Means Clustering Method. *Jurnal Sisfokom (Sistem Informasi Dan Komputer)*, 13(1), 86–91. <https://doi.org/10.32736/sisfokom.v13i1.1983>
- [8] Prasetyo, V. R., Erlangga, G., & Prima, D. A. (2023). Analisis Sentimen untuk Identifikasi Bantuan Korban Bencana Alam berdasarkan Data di Twitter Menggunakan Metode K-Means dan Naive Bayes. *Jurnal Teknologi Informasi Dan Ilmu Komputer*, 10(5), 1055–1062. <https://doi.org/10.25126/jtiik.20231057077>
- [9] Nugroho, J. R., Suprpto, Y. K., & Setijadi, E. (2022). Clustering Tingkat Risiko Klasifikasi

- Lapangan Usaha (KLU) Menggunakan Metode K-Means. *Jurnal Teknologi Informasi Dan Ilmu Komputer*, 9(3), 533–540. <https://doi.org/10.25126/jtiik.2022915217>
- [10] Kurniadi, D., Agustin, Y. H., Akbar, H. I. N., & Farida, I. (2023). Penerapan Algoritma k-Means Clustering untuk Pengelompokan Pembangunan Jalan pada Dinas Pekerjaan Umum dan Penataan Ruang. *Aiti*, 20(1), 64–77. <https://doi.org/10.24246/aiti.v20i1.64-77>
- [11] Marini, L. F., & Suhendra, C. D. (2023). Penggunaan Algoritma K-Means Pada Aplikasi Pemetaan Klaster Daerah Pariwisata. *Jurnal Media Informatika Budidarma*, 7(2), 707–713. <https://doi.org/10.30865/mib.v7i2.5558>
- [12] Praseptian M, D., Fadlil, A., & Herman, H. (2022). Penerapan Clustering K-Means untuk Pengelompokan Tingkat Kepuasan Pengguna Lulusan Perguruan Tinggi. *Jurnal Media Informatika Budidarma*, 6(3), 1693. <https://doi.org/10.30865/mib.v6i3.4191>
- [13] Amriyansah, A. J., Sulistiani, H., & Amalia, R. (2024). Penerapan Metode Forward Chaining Pada Sistem Pakar Diagnosa Penyakit Ayam Ternak. *Smatika Jurnal*, 14(01), 42–52. <https://doi.org/10.32664/smatika.v14i01.1001>
- [14] Yanto, M., Saputra, D., & Guswandi, D. (2021). Penerapan Metode Certainty Factor Pada Proses Diagnosa Penyakit Aktinomikosis. *Journal of Information System And Informatics Engineering*, 5(1), 37–44
- [15] Fakhriyah, N. N. (2021). Sistem Pakar Diagnosis Penyakit Pada Kambing Dengan Metode Forward Chaining Dan Certainty Factor. *Jurnal Teknologi Informasi, Komputer, Dan Aplikasinya (JTIKA)*, 3(1), 72–84. <https://doi.org/10.29303/jtika.v3i1.138>
- [16] Mintarsih, N., MS, D. D., Ariani, Y. M., & Hilda, A. M. (2024). Implementasi Metode Convolutional Neural Network (Cnn) Densenet121 Pada Diagnosa Penyakit Ayam (Manur). *Infotech: Journal of Technology Information*, 10(1), 85–98. <https://doi.org/10.37365/jti.v10i1.252>