



Implementasi Data Mining Menggunakan Algoritma K-Means untuk Pengelompokan Data Keluhan Pelanggan pada Perumda Tirta Musi Palembang

Fiara Jayantri¹, Ricca Amelia², Andika Maharaztu Putra Wisnu³, Yulistia⁴

^{1,2,3,4}Program Studi Sistem Informasi, Fakultas Ilmu Komputer dan Rekayasa, Universitas Multi Data Palembang, Palembang, Indonesia

¹fiarajayantri19@mhs.mdp.ac.id, ²ricca.amelia023@mhs.mdp.ac.id,

³andika.maharaztu_1@mhs.mdp.ac.id, ⁴yulistia@mdp.ac.id,

Abstract

Penelitian ini membahas penerapan algoritma *K-Means Klustering* untuk membantu Perumda Tirta Musi Palembang dalam mengelompokkan data keluhan pelanggan. Permasalahan utama yang dihadapi adalah tingginya *volume* keluhan yang diterima setiap tahun, namun belum terdapat sistem atau metode yang mampu mengelompokkan dan menganalisis data tersebut secara efisien. Dengan pendekatan *Knowledge Discovery in Databases (KDD)*, proses yang dilakukan meliputi pengumpulan data, pra-pemrosesan, transformasi data kategorikal menjadi bentuk numerik, serta proses klusterisasi menggunakan algoritma *K-Means*. Sistem pendukung berbasis web juga dikembangkan agar proses ini dapat digunakan secara praktis oleh operator atau administrator, bukan hanya sebagai proses satu kali. Hasil akhir membagi pelanggan menjadi dua kluster utama: kluster prioritas (keluhan tinggi) dan kluster stabil (keluhan rendah). Segmentasi ini mendukung pengambilan keputusan yang lebih tepat dalam upaya peningkatan pelayanan.

Keywords: *Klusterisasi K-Means, Kluster Data, Keluhan Pelanggan, Knowledge Discovery in Databases (KDD), Pelayanan Publik.*

1. Pendahuluan

Era digital telah mengubah cara kerja organisasi, termasuk dalam sektor pelayanan publik, yang kini dituntut untuk memberikan layanan cepat, transparan, dan berbasis data. Teknologi informasi memungkinkan pengumpulan dan analisis data dalam skala besar (*big data*), sehingga pemanfaatannya menjadi kebutuhan penting dalam proses pengambilan keputusan. Perumda Tirta Musi Palembang sebagai penyedia layanan air bersih menghadapi tantangan serupa, terutama dalam mengelola keluhan pelanggan yang terus meningkat seiring pertumbuhan penduduk Palembang yang telah mencapai 1,7 juta jiwa [1]. Penelitian ini menjawab ketiadaan sistem klasifikasi otomatis untuk keluhan pelanggan di Perumda Tirta Musi, yang selama ini ditangani secara manual. Hal ini menjadi celah penelitian yang signifikan karena belum ada studi terdahulu yang secara khusus membangun sistem berbasis data mining untuk *klustering* keluhan pelanggan di lingkungan Perumda Tirta Musi Palembang.

Ribuan keluhan tercatat setiap tahun, mencakup masalah seperti gangguan aliran air, kebocoran, keterlambatan teknis, hingga pencatatan tagihan yang tidak akurat. Namun hingga kini, belum tersedia sistem klasifikasi yang efektif; penanganan masih dilakukan

secara manual dan tersebar, sehingga menyulitkan dalam mengenali pola permasalahan secara menyeluruh. Padahal menurut [2], keluhan merupakan sumber informasi strategis yang jika dianalisis dengan benar, dapat digunakan sebagai dasar evaluasi dan perbaikan layanan. Hal ini sejalan dengan [3] yang menekankan pentingnya memahami persepsi pelanggan sebagai bagian dari upaya meningkatkan kualitas layanan publik.

Untuk itu, penelitian ini menawarkan solusi melalui penerapan metode *klustering*, khususnya menggunakan algoritma *K-Means Klustering*. Algoritma ini termasuk dalam *unsupervised learning* yang mampu mengelompokkan data berdasarkan kesamaan karakteristik tanpa memerlukan label. Pendekatan ini mengikuti tahapan *Knowledge Discovery in Databases (KDD)* agar proses pengelompokan lebih sistematis dan terstruktur. Beberapa penelitian terdahulu juga menunjukkan efektivitas K-Means dalam segmentasi layanan publik dan keluhan pelanggan [4], [5], [6], [7]. Oleh karena itu, penelitian ini bertujuan untuk menerapkan *K-Means* dalam pengelompokan keluhan pelanggan Perumda Tirta Musi Palembang guna mendukung peningkatan kualitas layanan secara *data-driven*.

2. Landasan Teori

2.1 Data Mining

Data mining merupakan proses untuk mengekstraksi informasi atau pola tersembunyi dari kumpulan data berukuran besar, yang berguna dalam mendukung pengambilan keputusan di masa depan [8], [6]. Tujuan utamanya adalah menemukan pengetahuan yang bernilai dari data yang tersedia. Metode ini dikenal juga sebagai *pattern recognition* dan banyak digunakan di berbagai bidang seperti teknologi, industri, dan keuangan. Teknik-teknik umum dalam data mining mencakup *klustering*, regresi, asosiasi, klasifikasi, dan peramalan [4].

2.2 Klustering

Klustering merupakan metode dalam unsupervised learning yang bertujuan mengelompokkan data ke dalam beberapa kluster berdasarkan tingkat kemiripan antar objek [6]. Prinsip utamanya adalah memaksimalkan kemiripan objek di dalam kluster yang sama (intra-cluster similarity) sekaligus meminimalkan kemiripan antar kluster yang berbeda (inter-cluster dissimilarity), sehingga setiap kluster merepresentasikan segmen data dengan karakteristik yang relatif homogen. Teknik ini digunakan untuk menemukan struktur atau pola tersembunyi dalam data dan diterapkan di berbagai bidang seperti klasifikasi, pengenalan pola, dan pemrosesan citra [4]. Dalam konteks analisis data, klustering membantu merangkum kumpulan data yang kompleks menjadi beberapa kelompok yang lebih mudah diinterpretasikan, misalnya untuk mengidentifikasi segmen perilaku, tingkat risiko, atau pola kejadian yang berulang. Proses pengelompokan ini dilakukan secara otomatis menggunakan algoritma, sehingga memungkinkan identifikasi kelompok tersembunyi dalam kumpulan data [5]. Artinya, pengelompokan tidak memerlukan label awal dari peneliti, melainkan dibentuk dari karakteristik data itu sendiri melalui ukuran jarak atau ukuran kemiripan tertentu (misalnya Euclidean, Manhattan, atau ukuran lain yang relevan). Hasil klustering kemudian dapat dimanfaatkan sebagai dasar pengambilan keputusan, seperti penentuan prioritas penanganan, pemetaan masalah dominan, hingga penyusunan strategi intervensi yang berbeda untuk masing-masing kluster. Dengan demikian, klustering tidak hanya berfungsi sebagai teknik pengelompokan, tetapi juga sebagai pendekatan analitik untuk menghasilkan wawasan yang lebih terstruktur dari data yang bersifat *heterogeny* [6].

2.3 K-Means

K-Means merupakan salah satu metode dalam *data mining* yang diawali dengan memilih sejumlah komponen dari populasi sebagai pusat awal untuk pembentukan *klaster* [6]. *K-Means* adalah metode analisis pengelompokan yang membagi N objek

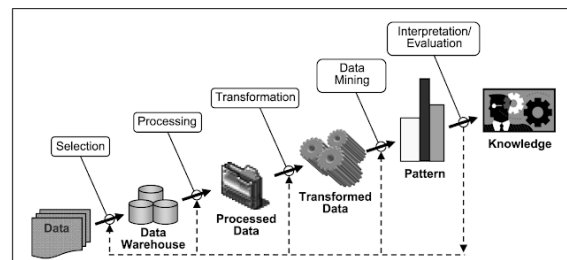
pengamatan ke dalam K *klaster*, di mana setiap objek ditempatkan pada *klaster* dengan rata-rata (mean) terdekat. Metode ini memiliki kemiripan dengan algoritma *Expectation Maximization* pada *Gaussian Mixture Model*, karena keduanya bertujuan menemukan pusat kluster melalui proses iteratif penyempurnaan [5].

2.4 Metode Elbow

Metode Elbow adalah teknik yang digunakan untuk menentukan jumlah kluster optimal dengan mengidentifikasi titik siku pada grafik evaluasi, yang menunjukkan perubahan signifikan dalam nilai evaluasi saat jumlah kluster bertambah [9]. Dengan menerapkan metode *Elbow*, dapat diperoleh informasi yang memudahkan dalam pemilihan jumlah *klaster* yang optimal berdasarkan titik perubahan yang signifikan dalam grafik [10].

2.5 Metode Data Mining

Penerapan data mining dalam penelitian ini menggunakan pendekatan *Knowledge Discovery in Databases* (KDD) yang bertujuan memberikan rekomendasi penyelesaian keluhan pelanggan secara lebih efektif. KDD merupakan proses berjenjang yang mencakup penggunaan data mining untuk menemukan pola dalam data, lalu mengubahnya menjadi informasi yang mudah dipahami. Meskipun istilah data mining dan KDD sering digunakan bergantian, keduanya memiliki perbedaan—data mining merupakan salah satu tahapan penting dalam proses KDD [11].



Gambar 1. Proses KDD Dalam Data Mining
Sumber : [11]

3. Rumusan Masalah

Berdasarkan latar belakang yang telah diuraikan di atas, maka rumusan masalah yang dapat diangkat dalam penelitian ini adalah: Bagaimana penerapan *data mining klustering* menggunakan algoritma *K-Means* dapat meningkatkan efisiensi dan efektivitas pada layanan pengaduan di Perumda Tirta Musi Palembang?

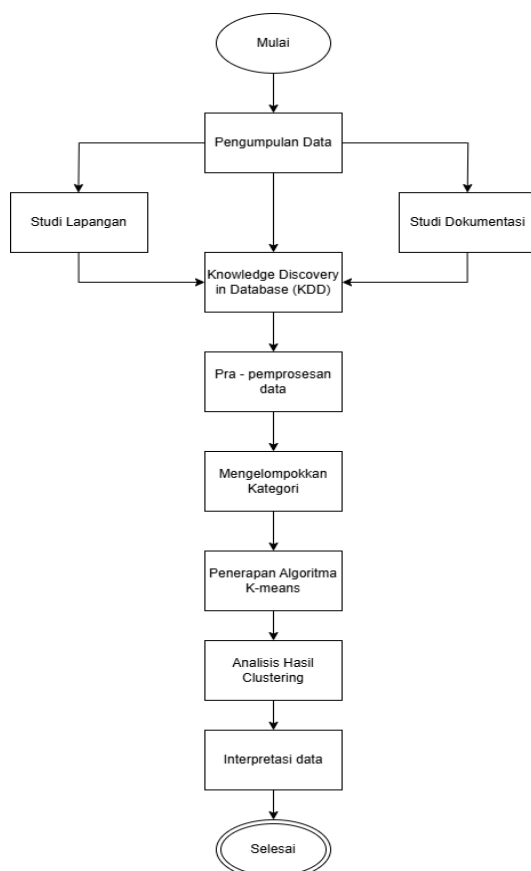
4. Tujuan

Tujuan dalam penelitian ini sebagai berikut :

1. Mengidentifikasi pola keluhan pelanggan Perumda Tirta Musi berdasarkan data pengaduan.
2. Mengelompokkan pelanggan menggunakan algoritma *K-means clustering* untuk menemukan segmen pelanggan dengan karakteristik keluhan yang serupa.
3. Menghasilkan rekomendasi segmentasi yang dapat digunakan oleh Perumda Tirta Musi untuk meningkatkan efisiensi dalam pelayanan pelanggan.

5. Metode Penelitian

Penelitian ini menggunakan pendekatan KDD (*Knowledge Discovery in Databases*) yang mencakup enam tahap berikut:



Gambar 2. Metode Penelitian

a. Pengumpulan Data

Data diperoleh dari arsip keluhan pelanggan selama periode 2023–2024. Setiap data mencakup: nama pelanggan, jenis keluhan, media pengaduan, nama teknisi/tukang yang menangani, dan waktu kejadian.

b. Pra-pemrosesan Data

Dari total 5.056 baris data awal, dilakukan pembersihan terhadap entri kosong, duplikat, dan data tidak relevan. Setelah proses ini, tersisa 2.767 data yang valid dan siap diolah.

c. Transformasi Data

Karena algoritma K-Means bekerja dengan menghitung jarak antar data (umumnya menggunakan Euclidean distance) dan mekanisme tersebut mensyaratkan fitur berbentuk numerik, maka atribut yang bersifat kategorikal seperti “jenis keluhan” dan “media pengaduan” tidak dapat diproses secara langsung dalam bentuk teks. Oleh sebab itu, kategori-kategori tersebut terlebih dahulu ditransformasikan ke representasi numerik menggunakan teknik Label Encoding, yaitu memberikan kode bilangan bulat unik untuk setiap kategori yang berbeda (misalnya WhatsApp = 0, Telepon = 1, Datang langsung = 2, dan seterusnya). Transformasi ini memungkinkan data kategorikal masuk ke ruang fitur numerik sehingga dapat dihitung jaraknya dan dikelompokkan oleh K-Means. Namun demikian, perlu dicatat bahwa Label Encoding dapat menimbulkan kesan adanya urutan (ordinal) antar kategori, padahal kategori seperti jenis keluhan atau media pengaduan pada dasarnya bersifat nominal. Untuk meminimalkan dampak interpretasi urutan semu tersebut, proses analisis selanjutnya difokuskan pada hasil pengelompokan (cluster assignment) dan interpretasi karakteristik tiap kluster, bukan pada makna “besar-kecil” nilai hasil encoding. Dengan langkah ini, data tetap kompatibel dengan K-Means sekaligus menjaga interpretasi agar tetap sesuai dengan sifat asli variabel kategorikal.

d. Penentuan Jumlah Kluster

Metode Elbow digunakan untuk menentukan jumlah kluster yang paling optimal dalam K-Means dengan mengevaluasi perubahan nilai WCSS (within-cluster sum of squares) pada berbagai kandidat jumlah kluster (K). WCSS merepresentasikan total variasi di dalam kluster, yaitu jumlah kuadrat jarak tiap data terhadap centroid klasternya; semakin kecil nilai WCSS, semakin rapat dan homogen anggota kluster. Pada praktiknya, nilai WCSS akan selalu menurun ketika K diperbesar, karena data dipartisi menjadi kelompok yang semakin spesifik. Namun, penurunan tersebut tidak selalu signifikan. Prinsip Elbow adalah memilih K pada titik ketika penurunan WCSS mulai melambat (diminishing returns), sehingga penambahan kluster berikutnya hanya memberikan perbaikan kecil dibanding kompleksitas model yang meningkat.

Berdasarkan grafik Elbow pada penelitian ini, terlihat bahwa pada K = 2 terjadi penurunan WCSS yang paling tajam dibandingkan perubahan pada K berikutnya. Pola tersebut membentuk “siku” (elbow point), yang mengindikasikan bahwa pemisahan data menjadi dua kluster sudah mampu menangkap struktur utama dalam data. Setelah K = 2, penurunan WCSS cenderung lebih landai, sehingga penambahan jumlah kluster tidak memberikan peningkatan kualitas pengelompokan yang sebanding. Oleh karena itu, K = 2 dipilih sebagai jumlah kluster optimal untuk tahap klusterisasi selanjutnya.

e. Klasterisasi dengan K-Means

Data pada penelitian ini diolah menggunakan bahasa pemrograman Python dengan memanfaatkan pustaka Scikit-learn untuk implementasi algoritma K-Means. Pada tahap awal, algoritma menetapkan centroid awal sebagai titik representatif untuk masing-masing klaster. Selanjutnya dilakukan proses iteratif yang terdiri dari dua langkah utama: (1) setiap data pelanggan di-assign ke klaster dengan centroid terdekat berdasarkan ukuran jarak, dan (2) centroid diperbarui dengan menghitung rata-rata (mean) dari seluruh anggota klaster tersebut sehingga pusat klaster merepresentasikan kondisi terbaru dari kelompok data. Siklus ini diulang secara berulang hingga mencapai kondisi konvergen, yaitu ketika posisi centroid tidak lagi berubah secara signifikan atau tidak terjadi perpindahan anggota klaster antar iterasi, sehingga solusi dianggap stabil. Output dari proses ini berupa label klaster untuk tiap data, nilai centroid akhir, serta pembagian data ke dalam kelompok yang lebih homogen. Berdasarkan hasil pengolahan, data pelanggan terbagi menjadi dua klaster, yang selanjutnya diinterpretasikan sebagai dua segmen pelanggan dengan karakteristik pengaduan yang berbeda untuk mendukung analisis dan strategi penanganan.

f. Pembangunan Sistem

Dibuat sistem web berbasis PHP yang memungkinkan admin mengunggah file CSV, menjalankan *klustering*, melihat hasil visualisasi (grafik Elbow, tabel klaster), serta unduh hasil.

6. Hasil dan Pembahasan

Hasil dari penelitian ini berupa sistem pengelompokan keluhan pelanggan. Sistem ini dikembangkan untuk membantu proses *Klustering* data keluhan berdasarkan pola yang terbentuk, sehingga memudahkan pihak Perumda Tirta Musi dalam mengidentifikasi kelompok pelanggan berdasarkan tingkat dan jenis keluhan. Sesuai dengan algoritma *K-Means*.

6.1 Data Set

Data dalam penelitian ini diperoleh dari rekapitulasi pengaduan pelanggan Perumda Tirta Musi Palembang pada rentang waktu Januari 2023 hingga tahun 2024. Dataset awal berjumlah sekitar 5.056 baris yang masih mencakup berbagai entri yang belum layak dianalisis, seperti baris kosong, duplikasi laporan, serta catatan yang tidak relevan dengan konteks pengaduan (misalnya informasi administratif yang tidak memuat substansi keluhan). Oleh karena itu, dilakukan tahap data cleansing untuk meningkatkan kualitas data dan memastikan setiap entri mewakili pengaduan yang valid. Proses ini meliputi penyaringan data yang tidak memiliki isi pengaduan, penghapusan entri yang tidak lengkap atau tidak terbaca, normalisasi format pencatatan, serta seleksi terhadap data yang sesuai dengan tujuan penelitian. Setelah tahapan pembersihan data tersebut, jumlah entri berkurang menjadi 2.767 data yang dinyatakan valid, konsisten, dan siap

digunakan pada tahap analisis selanjutnya. Dengan demikian, dataset akhir diharapkan lebih representatif terhadap pola pengaduan pelanggan dan meminimalkan bias yang dapat muncul akibat keberadaan data kosong maupun tidak relevan.

6.2 Proses Klasterisasi

Proses klasterisasi dilakukan melalui beberapa tahapan teknis sebagai berikut:



Gambar 3. Proses Klasterisasi

a) Pengumpulan Data

Berdasarkan rekapitulasi data pengaduan pelanggan periode Januari 2023 hingga 2024, tercatat ribuan keluhan yang masuk. Keluhan tersebut meliputi berbagai masalah seperti air tidak keluar, kebocoran pipa, masalah meteran, tagihan mahal, dan permintaan teknis lainnya.



Gambar 4. Grafik Keluhan Pelanggan

Sumber: Rekap pengaduan pelanggan Perumda Tirta Musi Palembang (Januari 2023–2024)

Jumlah awal: ±5.056 baris data

Grafik batang berjudul “Jumlah Keluhan Pelanggan Perumda Tirta Musi Palembang per Bulan” menampilkan dinamika keluhan pelanggan secara bulanan pada rentang Januari

2023 hingga Desember 2024. Sumbu-x merepresentasikan periode (bulan), sedangkan sumbu-y menunjukkan jumlah keluhan.

Secara umum, pola data memperlihatkan dua fase yang kontras:

- Fase keluhan tinggi dan relatif stabil (sebagian besar 2023 hingga awal 2024). Pada periode ini, jumlah keluhan berada pada kisaran menengah-tinggi dan berfluktuasi antar-bulan (sekitar ratusan keluhan per bulan). Variasi terlihat wajar sebagai dinamika layanan (misalnya gangguan distribusi, kualitas air, atau lonjakan permintaan musiman), namun tidak tampak tren naik/turun yang benar-benar konsisten sepanjang fase ini.
- Fase penurunan tajam dan bertahan rendah (mulai setelah awal 2024 hingga akhir 2024). Setelah satu titik waktu di awal 2024, grafik menunjukkan penurunan drastis ke level sangat rendah (puluhan kecil hingga mendekati nol) dan bertahan rendah hingga akhir periode, dengan beberapa kenaikan kecil sporadis. Pola ini mencerminkan perubahan struktural (structural break), bukan sekadar fluktuasi normal.

Selain dua fase tersebut, terdapat satu temuan penting, yakni Outlier/lonjakan ekstrem pada satu bulan dengan nilai melampaui 500 keluhan, jauh di atas bulan-bulan lain di fase pertama. Lonjakan sebesar ini biasanya mengindikasikan kejadian luar biasa (misalnya gangguan layanan skala besar, kerusakan infrastruktur utama, atau peristiwa lingkungan), atau artefak pencatatan (misalnya penggabungan laporan, perubahan kanal pelaporan, atau backlog yang baru diinput pada bulan tersebut).

Dari perspektif analisis data, grafik ini menyiratkan deret waktu yang tidak stasioner karena adanya outlier besar dan perubahan level (level shift) yang tajam. Implikasi akademiknya: interpretasi “peningkatan kinerja layanan” tidak boleh langsung disimpulkan hanya dari turunnya angka keluhan, karena penurunan bisa juga disebabkan oleh perubahan mekanisme pelaporan, penurunan akses kanal pengaduan, atau ketidaklengkapan data. Karena itu, pembacaan yang lebih kuat perlu memisahkan analisis menjadi pra-perubahan vs pasca-perubahan, lalu memverifikasi faktor operasional (mis. kebijakan, sistem

pengaduan, atau intervensi layanan) yang terjadi pada titik transisi tersebut.

- Pembersihan Data (*Data Cleansing*)
Menghapus data kosong, duplikat, dan tidak relevan
Hasil akhir: 2.767 data valid
- Transformasi Data Kategorikal
Proses *Label Encoding* untuk mengubah data non-numerik menjadi numerik
- Penentuan Jumlah Kluster (K)
Metode Elbow digunakan untuk menentukan Hasil: $K = 2$
- Proses Klasterisasi dengan K-Means
Parameter yang digunakan dalam algoritma K-Means meliputi jumlah kluster ($n_clusters=2$), iterasi maksimum ($max_iter=300$), dan $random_state=42$ untuk memastikan hasil konsisten. Dilakukan menggunakan Python (*Scikit-learn*) Proses iteratif hingga *centroid konvergen*
- Integrasi ke Sistem Web
Visualisasi hasil dan interpretasi kluster ditampilkan dalam sistem berbasis PHP.

6.3 Hasil Klasterisasi Pelanggan

Algoritma *K-Means* berhasil mengelompokkan pelanggan ke dalam dua kluster utama berdasarkan lima atribut: nama pelanggan, jumlah keluhan, media pengaduan, petugas, dan tukang.

Tabel 1. Kluster Pelanggan

Kluster	Jumlah Data	Karakteristik Umum	Interpretasi
Kluster 1	1.074	- Keluhan tinggi (5 - 8 keluhan) - Media online aktif (1) - Banyak petugas/tukang terlibat (15–30)	Pelanggan kritis / rawan masalah
Kluster 2	1.692	- Keluhan rendah atau nihil (0 - 2 keluhan) - Media penyampaian tidak aktif (0) - Minim keterlibatan petugas (1–8)	Pelanggan stabil / normal

- Kluster 1 (Rawan)
Pelanggan dengan keluhan tinggi, sering menggunakan media online, serta ditangani oleh banyak petugas. Kluster ini menunjukkan

intensitas permasalahan tinggi, sehingga perlu tindakan cepat dan penanganan prioritas.

2. Klaster 2 (Stabil)

Pelanggan dengan keluhan rendah atau tidak ada, media pengaduan pasif, serta hanya ditangani oleh sedikit petugas. Pelanggan ini dinilai lebih stabil, dan cocok untuk pendekatan efisiensi.

6.4 Rekomendasi dan Strategi Penanganan

Hasil Analisis menunjukkan adanya ketidakseimbangan data keluhan, di mana klaster prioritas memiliki jumlah keluhan lebih banyak dibanding klaster stabil. Tidak ditemukan outlier ekstrem, namun distribusi ini menunjukkan perlunya strategi fokus untuk klaster prioritas. Berdasarkan hasil klasterisasi, disusun strategi pelayanan yang berbeda untuk masing-masing klaster:

Tabel 2. Klaster 1 Pelanggan Rawan / PrioritasTinggi

Proaktif Monitoring	Tim Penanganan Khusus	Analisis Lanjutan	Pemberdayaan Pelanggan
Kunjungan berkala atau <i>follow-up</i> otomatis	Alokasikan petugas/tukang senior untuk klaster ini	Telusuri area/wilayah/jenis layanan yang dominan dalam klaster ini	Edukasi pelanggan tentang saluran penyampaian masalah yang efektif
Survey kepuasan setelah setiap penanganan	Waktu <i>respons</i> lebih cepat	Identifikasi kemungkinan akar masalah teknis atau administratif	

Tabel 2 menjelaskan Klaster 1 (pelanggan rawan/prioritas tinggi) yang memerlukan penanganan intensif dan terarah. Strateginya mencakup monitoring proaktif melalui komunikasi berkala atau *follow-up* otomatis, serta penugasan tim khusus dengan alokasi sumber daya yang lebih kuat (misalnya petugas senior) agar respons lebih cepat. Dari sisi analisis, dilakukan penelusuran akar masalah untuk mengidentifikasi jenis layanan yang paling sering memicu keluhan dan menentukan apakah penyebabnya bersifat teknis atau administratif. Selain itu, pelanggan diberdayakan melalui edukasi terkait penanganan awal sehingga pengaduan dapat dicegah atau diminimalkan. Setelah penanganan, dilakukan survei kepuasan sebagai umpan balik untuk evaluasi kualitas layanan.

Tabel 3. Klaster 2 Pelanggan Stabil / Penanganan Umum

Layanan Reguler	Efisiensi Operasional	Peluang Peningkatan Layanan
Tetap diberi akses pengaduan, namun tidak prioritas	Petugas tingkat dasar cukup	Pelanggan ini bisa dijadikan basis loyalitas atau promosi
	Penanganan sesuai SLA standar	Bangun sistem reward/informasi preventif untuk menjaga kepuasan

Tabel 3 menjelaskan bahwa Klaster 2 (pelanggan stabil) termasuk kategori penanganan umum. Layanan tetap diberikan akses dan ditangani petugas, namun tidak menjadi prioritas utama. Dari sisi operasional, penanganan mengikuti SLA standar sehingga efisien dan tidak memerlukan intervensi khusus. Meski stabil, klaster ini tetap memiliki peluang peningkatan layanan, misalnya melalui program loyalitas/promosi serta pembangunan sistem informasi dan komunikasi yang lebih baik untuk menjaga kepuasan pelanggan.

Tabel 4. Strategi Penanganan

Aspek Operasional	Klaster 1 (Kritis)	Klaster 2 (Stabil)
Respons Waktu	Cepat (<24 jam)	Sesuai SLA normal (1–3 hari)
Tim Penanganan	Petugas senior / tim khusus	Petugas umum / lapangan
Monitoring	Proaktif dan rutin (mingguan/bulanan)	Pasif, hanya jika ada pengaduan
Evaluasi Khusus	Ya – evaluasi sumber masalah tiap bulan	Tidak wajib, cukup agregat triwulan
Komunikasi ke Pelanggan	Terjadwal (notifikasi, email, edukasi)	Sesekali (survey berkala atau promosi umum)
Prioritas Penugasan	Tinggi (berbasis analitik klaster)	Menyesuaikan jika tidak ada klaster prioritas lain

Tabel 4. menjelaskan strategi penanganan yang dibedakan berdasarkan hasil klasterisasi menjadi Klaster 1 (Kritis) dan Klaster 2 (Stabil). Pada klaster kritis, penanganan dilakukan lebih agresif: respons dipercepat (<24 jam), ditangani petugas senior/tim khusus, monitoring bersifat proaktif dan rutin, serta dilakukan evaluasi bulanan untuk menelusuri sumber masalah. Komunikasi ke pelanggan juga dibuat terjadwal (mis. notifikasi, email, edukasi), dan kasus dalam klaster ini ditempatkan sebagai prioritas utama berbasis analitik. Sebaliknya, pada klaster stabil, penanganan mengikuti SLA normal (1–3 hari) oleh petugas lapangan umum, monitoring dilakukan jika ada pengaduan, evaluasi tidak wajib dan cukup agregat triwulan, komunikasi hanya sesekali, serta prioritas bersifat menyesuaikan jika tidak ada kasus lain yang lebih mendesak.

7. Kesimpulan

Penelitian ini berhasil mengimplementasikan metode *K-Means* untuk mengelompokkan keluhan pelanggan di Perumda Tirta Musi menjadi beberapa klaster prioritas. Hasil analisis menunjukkan bahwa sebagian besar keluhan terkonsentrasi pada klaster prioritas tinggi, yang mengindikasikan perlunya penanganan khusus dan berkelanjutan pada kelompok ini. Penerapan klasterisasi terbukti membantu dalam memetakan distribusi keluhan dan memberikan gambaran yang jelas mengenai fokus perbaikan pelayanan.

Berdasarkan temuan tersebut, direkomendasikan agar Perumda Tirta Musi menetapkan target operasional yang terukur, yakni mengurangi jumlah keluhan pada klaster prioritas tinggi hingga 20% dalam kurun waktu 12 bulan. Target ini dapat dicapai melalui peningkatan respon layanan, optimalisasi sumber daya teknis, dan pelaksanaan program monitoring berkala terhadap pelanggan yang masuk dalam klaster prioritas.

Selain itu, penerapan sistem klusterisasi ini diharapkan dapat dikembangkan secara berkelanjutan, misalnya dengan menambahkan variabel lain seperti waktu penanganan, lokasi keluhan, dan kategori teknis untuk memperkaya analisis. Langkah ini akan mendukung upaya jangka panjang dalam meningkatkan kepuasan pelanggan dan efektivitas pengelolaan keluhan di sektor pelayanan Perumda Tirta Musi Palembang.

Daftar Pustaka

- [1] BPS, "Statistik Air Bersih 2018-2022," *Bps.Go.Id*, 2023.
- [2] G. G. Wirakanda dan I. S. Putri, "Analisis Penanganan Keluhan Pelanggan (Studi Kasus Di Kantor Pos Bandung 40000)," *J. Bisnis dan Pemasar.*, vol. 10, no. 2, hal. 1–11, 2020, [Daring]. Tersedia pada: <https://ejurnal.poltekpos.ac.id/index.php/promark/article/view/1020/698>
- [3] Sambodo Rio Sasongko, "Faktor-Faktor Kepuasan Pelanggan dan Loyalitas Pelanggan (Literature Review Manajemen Pemasaran)," *J. Ilmu Manaj. Terap.*, vol. 3, no. 1, hal. 104–114, 2021, doi: 10.31933/jimt.v3i1.707.
- [4] E. H. Wisanta dan Y. N. Marlim, "Analisis Algoritma K-Means Untuk Klastering Kepuasan Pelayanan: Mall Pelayanan Publik Pekanbaru," *Semin. Nas. Inform. Pros. Senat*. 2021, hal. 223–228, 2021.
- [5] K. Annisa, B. S. Ginting, dan M. A. Syar, "Penerapan Data Mining Pengelompokan Data Pengguna Air Bersih Berdasarkan Keluhannya Menggunakan Metode Klastering Pada Pdam Langkat," *J. Sist. Inf. Kaputama*, vol. 6, no. 2, hal. 165–179, 2022, doi: 10.59697/jsik.v6i2.167.
- [6] A. Pangestu dan T. Ridwan, "Penerapan Data Mining Menggunakan Algoritma K-Means Pengelompokan Pelanggan Berdasarkan Kubikasi Air Terjual Menggunakan Weka," *JUST IT J. Sist. Informasi, Teknol. Inf. dan Komput.*, vol. 11, no. 3, hal. 67–71, 2022, [Daring]. Tersedia pada: <https://jurnal.umj.ac.id/index.php/just-it/article/view/11591>
- [7] P. Dwi Lestari dan M. Mulyawan, "Datamining Pada Penjualan Air Bersih Di Spam Akidah Menggunakan Algoritma K-Means Klastering Menggunakan Rapidminer," *JATI (Jurnal Mhs. Tek. Inform.*, vol. 7, no. 1, hal. 412–416, 2023, doi: 10.36040/jati.v7i1.6315.
- [8] S. A. Perdana, S. F. Florentin, dan A. Santoso, "Analisis Segmentasi Pelanggan Menggunakan K-Means Klastering Studi Kasus Aplikasi Alfagift," *Sebatik*, vol. 26, no. 2, hal. 420–427, 2022, doi: 10.46984/sebatik.v26i2.2134.
- [9] V. A. Ekasetya dan A. Jananto, "Klusterisasi Optimal Dengan Elbow Method Untuk Pengelompokan Data Kecelakaan Lalu Lintas Di Kota Semarang," *J. Din. Inform.*, vol. 12, no. 1, hal. 20–28, 2020, doi: 10.35315/informatika.v12i1.8159.
- [10] B. G. Sudarsono, M. I. Leo, A. Santoso, dan F. Hendrawan, "Analisis Data Mining Data Netflix Menggunakan Aplikasi Rapid Miner," *JBASE - J. Bus. Audit Inf. Syst.*, vol. 4, no. 1, hal. 13–21, 2021, doi: 10.30813/jbase.v4i1.2729.
- [11] I. Y. Yudo Bismo Utomo, Iin Kurniasari, "Penerapan Knowledge Discovery in Database," *J. Tek. Inform. Kaputama*, vol. 7, no. 1, 2023.