

Analisis Kepuasan Pasien Terhadap Pelayanan Rumah Sakit Umum Seger Waras Menggunakan Metode K-Means Clustering

Ramda Gita Gautama¹, M. Gilang Suryanata², Tugiono³

^{1,2,3} Sistem Informasi, Stmik Triguna Dharma

Email: ¹rgitagautama@gmail.com, ²suryanatagilang@gmail.com, ³tugix.line@gmail.com

Email Penulis Korespondensi: rgitagautama@gmail.com

Article History:

Received Jun 15th, 2026

Revised Jun 30th, 2026

Accepted Jul 30th, 2026

Abstrak

Rumah Sakit Umum Seger Waras merupakan salah satu institusi kesehatan yang memberikan pelayanan rawat inap, rawat jalan, dan rawat darurat bagi masyarakat. Dengan tingginya volume kunjungan pasien setiap harinya, penting bagi rumah sakit untuk mempertahankan kualitas pelayanan yang optimal demi kenyamanan dan kepuasan pasien. Penelitian ini bertujuan untuk mengukur tingkat kepuasan pasien terhadap pelayanan yang diberikan oleh Rumah Sakit Umum Seger Waras menggunakan metode K-Means Clustering. Metode ini memungkinkan pengelompokan data berdasarkan karakteristik yang mirip, sehingga dapat memudahkan analisis kepuasan pasien. Beberapa studi sebelumnya menunjukkan hasil yang signifikan dalam pengukuran kepuasan menggunakan K-Means Clustering, dengan berbagai tingkat kepuasan yang ditemukan dalam layanan publik. Penelitian ini diharapkan dapat memberikan gambaran mengenai elemen pelayanan yang perlu dipertahankan maupun ditingkatkan, serta membantu Rumah Sakit Umum Seger Waras dalam meningkatkan kualitas layanan kesehatan secara keseluruhan.

Kata Kunci: K-Means Clustering, Kepuasan Pasien, Kualitas Pelayanan, Data Mining, Rumah Sakit Umum Seger Waras.

Abstract

Seger Waras General Hospital is a healthcare institution that provides inpatient, outpatient, and emergency services to the community. With the high volume of patient visits every day, it is essential for the hospital to maintain optimal service quality to ensure patient comfort and satisfaction. This study aims to measure patient satisfaction with the services provided by Seger Waras General Hospital using the K-Means Clustering method. This method allows for the grouping of data based on similar characteristics, facilitating the analysis of patient satisfaction. Previous studies have shown significant results in measuring satisfaction using K-Means Clustering, with varying levels of satisfaction identified in public services. This research is expected to provide an overview of the service elements that need to be maintained or improved and assist Seger Waras General Hospital in enhancing the overall quality of healthcare services.

Keywords: K-Means Clustering, Patient Satisfaction, Service Quality, Data Mining, Seger Waras General Hospital.

1. PENDAHULUAN

Rumah sakit merupakan institusi pelayanan kesehatan yang menyelenggarakan pelayanan kesehatan perorangan secara paripurna yang menyediakan pelayanan rawat inap, rawat jalan, dan rawat darurat.

Rumah Sakit Umum Seger Waras saat ini merupakan salah satu unit pelayanan kesehatan masyarakat yang setiap harinya melayani banyak pasien sehingga pihak Rumah Sakit Umum Seger Waras dapat menghasilkan data kunjungan pasien yang cukup banyak. Oleh karena itu rumah sakit harus siap siaga dalam memberikan pelayanan yang berkualitas untuk pasien. Efisiensi dan produktivitas pelayanan kesehatan merupakan aspek penting yang perlu diperhatikan. Pelayanan yang baik terhadap pasien menjadi faktor yang penting dalam memberikan kenyamanan terhadap pasien. Hal ini menjadi pemicu untuk membangun fasilitas kesehatan di kalangan masyarakat dalam memenuhi standarisasi pelayanan kesehatan yang baik.

Clustering merupakan metode yang digunakan dalam data mining yang cara kerjanya mencari dan mengelompokkan data yang mempunyai kemiripan karakteristik antara data satu dengan data lainnya yang telah diperoleh [1]. Dengan menggunakan metode K-Means Clustering dapat membantu untuk mengetahui berapa tingkat kepuasan pasien terhadap pelayanan yang telah diberikan. Metode K-Means Clustering merupakan suatu metode algoritma yang berbeda-beda, sehingga dapat memudahkan untuk menemukan data yang diinginkan. Penelitian yang dilakukan Ulfhi Burelia terkait Mengukur Tingkat Kepuasan Masyarakat Pada Pelayanan Kepolisian Resor (Polres) Dumai Menggunakan Algoritma K-

Means Clustering menghasilkan grafik kepuasan yaitu Step I (68 responden sangat puas dan 11 responden cukup puas), Step II (66 responden sangat puas dan 13 responden puas), Step III (79 responden sangat puas), Step IV (66 responden sangat puas dan 13 responden puas), Step V (67 responden sangat puas dan 12 responden cukup puas). Jadi dari step I-V dapat disimpulkan tingkat kepuasan pada pelayanan Di Polres Dumai Sangat Puas [2]. Sedangkan penelitian dengan judul Implementasi Algoritma K-Means dalam Menentukan Clustering pada Penilaian Kepuasan Pelanggan di Badan Pelatihan Kesehatan Pekanbaru hasilnya sebesar 259 menyatakan puas dan sebesar 44 menyatakan tidak puas dari 303 pelanggan [3]. Selanjutnya Analisis Algoritma K-Means Untuk Clustering Kepuasan Pelayanan: Mall Pelayanan Publik Pekanbaru memiliki cluster C1 mendapatkan nilai rata-rata yang termasuk dalam kategori kurang puas dengan nilai $C1=2.052380$ dan cluster C2 berada pada tingkat puas yaitu $C2=3.064425$, Sementara cluster C3 masuk kedalam kategori sangat puas dengan nilai $C3=4.136904$ [4].

Oleh karena itu, tulisan ini mengulas berupa mengaplikasikan data mining sebagai solusi dalam mengukur dan memahami kepuasan pasien dengan memperhitungkan tingkat kepentingan atribut-atribut layanannya. Kualitas pelayanan yang diberikan kepada pasien merupakan salah satu indikator yang menentukan kepuasan pasien terhadap apa yang diberikan Rumah Sakit Umum Seger Waras. Penelitian yang dilakukan mempunyai tujuan yaitu untuk melihat gambaran kepuasan pasien terhadap pelayanan kesehatan, sehingga dapat diketahui unsur yang dipertahankan dan diperbaiki oleh Rumah Sakit Umum Seger Waras dan dapat lebih meningkatkan kualitas pelayanannya.

2. METODOLOGI PENELITIAN

2.1 Tahapan Penelitian

Metode Penelitian merupakan suatu konsep yang menggunakan metodologi penelitian jenis kuantitatif. Dimana proses penyelidikan yang dimaksud untuk memahami seberapa sering, seberapa banyak, atau sejauh mana suatu fenomena terjadi. Dengan tujuan peneliti akan menganalisis data yang diperoleh. Jika metodologi yang dilakukan baik, maka semakin baik pula hasil yang didapatkan

2.1.1 Pengumpulan Data (*Data Collecting*)

Adapun teknik dalam pengumpulan data ada terdapat beberapa cara yang bisa dilakukan, yaitu :

1. Observasi
Observasi merupakan teknik dalam pengumpulan data dengan melakukan pengamatan secara langsung ke suatu tempat studi kasus yang dimana akan dilakukan suatu penelitian dan menempatkan tempat tersebut sebagai titik utama dalam pengamatan.
2. Kuesioner
Kuesioner adalah teknik pengumpulan data yang dilakukan dengan cara memberi seperangkat pertanyaan atau pernyataan kepada responden. Dalam penelitian ini kuesioner akan diberikan oleh peneliti kepada pasien di rumah sakit umum seger waras. Kuesioner terdiri dari deskripsi dan tujuan survey, pertanyaan berupa identitas responden, dan pertanyaan mengenai pelayanan yang diberikan oleh rumah sakit. Penelitian ini bertujuan mengetahui informasi terkait bidang penelitian serta menjadi bagian dari pengumpulan data. Adapun hasil kuesioner yang telah didapatkan berupa sampel data sebagai berikut

Tabel 1. Data Kuesioner

No	Nama	K1	K2	K3	K4	K5
1	YANNA FITRIYANA	4	6	8	6	6
2	MUSLIMIN	4	6	8	6	4
3	RATIH	6	6	6	6	10
4	IDA SAFITRI	6	6	6	6	6
5	DELLA ALFITA	8	6	6	8	8
6	SINTIA NANDA UTAMI	8	4	6	6	10
7	MUSRIFA	6	8	8	4	4
8	JUNAIDA SARI	8	6	6	6	6
9	AYU LESTARI	6	6	8	8	4
10	YULINDA	8	6	8	6	4
11	ANGGI PERTIWI	8	6	8	8	8
12	RINDI ANDRIANI	8	4	4	4	6
13	ELI AROSA	4	8	8	4	8
14	CAHAYA SARTIKA	8	8	6	4	10
15	SRI WINDARI	4	8	2	6	6
16	ENDORO HANDOKO SIJABAT	6	4	6	4	6
17	FIKA ELIZAH	4	4	8	8	6
18	ROMAULI ARITONANG	4	6	6	4	6
19	FADLY RIFKYANSYAH	4	6	6	6	8

20	ENI SUSANTI	10	6	4	8	6
21	RATI SERVIA	4	4	10	6	6
22	FARIDAH HANNUM	8	6	8	6	8
23	TRIA WANTI	8	4	8	8	6
24	SRI WAHYUNI	8	6	4	8	8
25	IRMAYANI	4	6	4	6	6
26	FAJRIN RAMADINA	8	8	6	6	8
27	ERMINA MARIANA	8	8	8	6	6
28	NIA RAMADHANI	4	6	6	6	4
29	EKA SIAHAAN	6	4	8	6	6
30	YULIANY	8	6	6	8	6
31	MAYA ROMANTI	6	6	6	8	8
32	SUSIYAN	8	8	6	4	8
33	PINKA ERAWATI	4	8	6	4	6
34	RINDI PUSPITA SARI	6	8	6	6	8
35	MAHINDA	8	8	6	8	6
36	DWI ASTUTI	4	6	4	10	8
37	DIANA LESTARI	6	4	8	8	8
38	RISMAWATI	8	6	8	6	6
39	RATNA PUTRI	6	6	6	8	6
40	SITI AMINAH	2	8	6	4	6

2.2 Data Mining

Data Mining adalah suatu proses penambangan atau penemuan informasi baru yang dilakukan dengan cara mencari sebuah pola atau aturan tertentu dari sejumlah data yang menumpuk dan dikatakan data besar. *Data Mining* juga dapat diartikan sebagai serangkaian suatu proses dalam mencari atau menggali nilai tambah suatu data yang berupa pengetahuan yang selama ini tidak diketahui secara manual yang pengetahuannya dapat bermanfaat

Data Mining bukan merupakan suatu bidang yang dapat dikatakan baru. Data Mining adalah sebuah pengembangan dan pencabangan dari ilmu Statistik. Oleh sebab itu Data Mining dan ilmu statistik sangat memiliki keterkaitan satu sama . Data Mining mewarisi sangat banyak bidang, aspek dan teknik dari bidang-bidang ilmu lainnya yang sudah mapan terlebih dahulu. Dimulai dengan beberapa disiplin ilmu terdahulu, hal penting yang terkait dengan Data Mining adalah [5]:

1. Data Mining merupakan suatu proses otomatis terhadap data yang sudah ada.
2. Data yang akan diproses berupa data yang sangat besar.
3. Tujuan *Data Mining* adalah mendapatkan hubungan atau pola yang akan mungkin memberikan indikasi yang bermanfaat.

2.2.1 Teknik Data Mining

Secara umum, teknik data mining dapat dikelompokkan menjadi beberapa kategori utama berdasarkan jenis analisis yang dilakukan dan tujuan pengolahan datanya. Berikut adalah beberapa kelompok utama dalam data mining:

1. Deskripsi
Deskripsi merupakan cara untuk menggambarkan pola dan kecenderungan yang terdapat dalam data yang dimiliki.
2. Estimasi
Estimasi hampir sama dengan klasifikasi, hanya saja nilai peubah atau variable target estimasi lebih ke arah data angka atau numerik.
3. Prediksi
Prediksi adalah suatu cara dalam menerka/menebak sebuah nilai yang belum diketahui sebelumnya.
4. Klasifikasi
Dalam klasifikasi terdapat target variable bertipe kategori, contohnya adalah penggolongan pendapatan yang dapat dipisahkan kedalam tiga kategori, yaitu tinggi, sedang, dan rendah.
5. Pengklasteran
Pengklasteran adalah pengelompokan data record, pengamatan, atau memperhatikan dan membentuk kelas objek-objek atau titik-titik yang memiliki kemiripan satu dengan yang lainnya.
6. Asosiasi
Asosiasi bertugas menemukan atribut yang muncul dalam satu waktu. Dalam dunia bisnis lebih umum disebut analisis keranjang belanja. Suatu sifat yang menjadi sebuah ciri- ciri dari suatu objek disebut dengan karakteristik.

2.2.2 Pengelompokan Data Mining

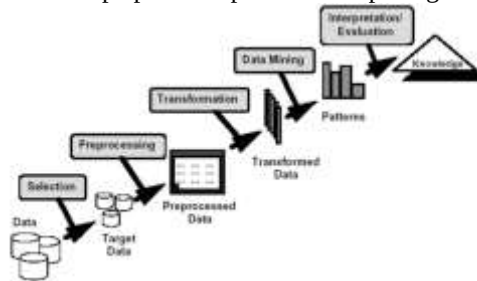
Karakteristik *data mining* mencakup sejumlah atribut dan sifat yang menggambarkan proses dan teknik yang digunakan dalam analisis data. *Data mining* bertujuan untuk mengekstraksi informasi yang berguna dari volume besar data yang tidak terstruktur atau tidak teratur. Ini mencakup pengidentifikasian pola, hubungan, tren, atau pengetahuan baru yang mungkin tidak terlihat secara langsung dalam data mentah. Ada beberapa karakteristik yang dimiliki *Data Mining* yaitu sebagai berikut:

1. Proses dalam menemukan sesuatu objek, informasi atau hal yang belum terlihat dan pola suatu data tertentu yang belum diketahui sebelumnya tanpa menjalankan proses penambangan oleh sipengguna.
2. Data yang menumpuk atau data yang besar sering dipergunakan untuk memperoleh hasil penambangan yang lebih akurat dan bermanfaat karena menggunakan data yang tergolong menumpuk dan sangat besar.

Dapat berguna dalam membuat, merancang ataupun menganalisis sebuah keputusan yang kritis terutama dalam strategi. Dari beberapa penjelasan tersebut dapat ditarik sebuah pernyataan bahwa *Data Mining* bisa dikatakan suatu cara atau teknik dalam menggali sebuah informasi berharga yang diperoleh melalui data yang banyak dan tersembunyi pada suatu koleksi data (*database*) yang sangat besar atau menumpuk sehingga ditemukan suatu pola yang menarik dan bermanfaat yang sebelumnya tidak diketahui pemilik data.

2.2.3 Knowledge Discovery in Databases (KDD)

Pada proses *Data Mining* yang biasanya disebut *knowledge discovery database (KDD)*. *Knowledge Discovery Databases (KDD)* adalah penerapan metode saintifik pada *Data Mining*. Dalam penjelasan ini *Data Mining* merupakan satu langkah dari proses KDD, terdapat beberapa proses seperti terlihat pada gambar dibawah 1. [6].



Gambar 1 Proses *Knowledge Discovery Database*

Penjelasan pada gambar proses *Knowledge Discovery Database (KDD)* terdapat beberapa proses dari tahap menjalankan KDD yaitu sebagai berikut:

1. Seleksi Data (*Selection*)
Selection berarti proses seleksi/pemilihan dari data yang dilakukan sebelum menuju pada tahap penelusuran dalam *Knowledge Discovery Database (KDD)* dimulai dengan ketentuan data dipilih berdasarkan tujuan. Data yang di kumpulkan akan digunakan dalam pemrosesan yang dilakukan oleh *Data Mining*.
2. Pemilihan Data (*Preprocessing/Cleaning*)
Proses *preprocessing* adalah suatu proses yang meliputi antara lain menghapus atau membuang data ganda yang tidak digunakan, memeriksa data yang dianggap tidak konsisten, dan memperbaiki kesalahan pada data, seperti kesalahan penulisan (*tipografi*).
3. Transformasi (*Transformation*)
Pada *fase* ini yang dilakukan adalah mengubah suatu bentuk data yang belum mempunyai beberapa entitas yang jelas ke dalam data yang siap untuk dilakukan proses *Data Mining*.
4. *Data Mining*
Pada proses ini, yang dilakukan adalah melakukan penerapan algoritma atau metode pencarian pengetahuan dari data yang dihasilkan pada proses transformasi.
5. Interpretasi/Evaluasi (*Interpretation/Evaluation*)
Pada *fase* ini yang paling terakhir ini, proses yang dilakukan adalah proses membentuk sebuah output atau hasil yang mudah dimengerti dan berbentuk sebuah informasi yang bermanfaat [7].

2.2.4 Tahapan Proses Data Mining

Tahapan proses dalam cara kerja *Data Mining* yang merupakan suatu pengolahan dalam tahapan yang ada pada tahap *Knowledge Discovery in Databases (KDD)* [8] seperti yang terlihat pada gambar dapat diuraikan sebagai berikut :

1. Paham terhadap sumber aplikasi dalam mengetahui, mencari dan menggali pengetahuan awal sesuai dengan yang diharapkan dan menjadi sasaran.
2. Merancang target data-set yang di butuhkan dalam proses data mining yang meliputi pemilihan sebuah data yang diperlukan dan tetap fokus pada isi –isi sebuah data.
3. Pembersihan dan transformasi data meliputi penghapusan bagian –bagian yang dianggap tidak perlu.
4. Penggunaan algoritma *Data Mining* yang bertujuan mendapatkan hasil berupa evaluasi dan informasi
5. Interpretasi, evaluasi dan visualisasi pola untuk melihat apakah ada sesuatu yang baru dan menarik, serta apa yang menjadi hasil dari penerapan sebuah *algoritma*.

2.3 Pelayanan

Pelayanan merupakan suatu tindakan atau perbuatan seseorang atau organisasi untuk memberikan kepuasan kepada konsumennya. Pelayanan timbul karena adanya kewajiban sebagai suatu proses penyelenggaraan kegiatan organisasi, baik organisasi pemerintah maupun organisasi swasta. Secara umum pelayanan dapat diartikan dengan melakukan kegiatan atau perbuatan yang hasilnya ditujukan untuk kepentingan orang lain, baik perorangan, kelompok atau masyarakat. Pelayanan dikatakan bermutu apabila mampu menimbulkan kepuasan bagi pasien yang dilayaninya.

Kepuasan adalah perasaan senang atau kecewa seseorang yang muncul setelah membandingkan antara kinerja (hasil) yang dipikirkan terhadap kinerja (hasil) yang diharapkan. Kepuasan konsumen dapat mempengaruhi minat untuk kembali ke rumah sakit yang sama. Konsumen yang puas akan menjadi pelanggan yang loyal, seperti promosi dari mulut ke mulut bagi calon konsumen lainnya, yang diharapkan sangat positif bagi rumah sakit [9].

Pelayanan kesehatan adalah pemenuhan kebutuhan dan tuntutan dari para pasien, dimana pasien mengharapkan suatu penyelesaian dari masalah kesehatannya. Tingkat kepuasan pasien menunjukkan tingkat keberhasilan suatu layanan kesehatan dalam meningkatkan mutu pelayanannya [10]. Banyak pasien yang mengeluh perawat kurang ramah dan lambat dalam menangani keluhan pasien. Tingginya beban kerja, banyaknya tugas dari dokter dan banyaknya jumlah pasien sering menjadi alasan mengapa pelayanan menjadi kurang optimal. Kepuasan pasien adalah persepsi bahwa harapannya telah terpenuhi, diperoleh hasil yang optimal bagi pelayanan kesehatan dengan memperhatikan pasien dan keluarganya, perhatian terhadap kebutuhan pasien, serta kondisi lingkungan fisik [11].

2.4 K-Means

Pengertian dari *K-Means Clustering* adalah, K dimaksudkan sebagai konstanta jumlah *cluster* yang diinginkan, Means dalam hal ini berarti nilai suatu rata-rata dari suatu grup data yang dalam hal ini didefinisikan sebagai *cluster*, sehingga *K-Means Clustering* adalah suatu metode penganalisaan data atau metode data mining yang melakukan proses pemodelan tanpa supervisi (unsupervised) dan merupakan salah satu metode yang melakukan pengelompokan data dengan sistem partisi. Metode K-Means berusaha mengelompokkan data yang ada kedalam beberapa kelompok, dimana data dalam satu kelompok mempunyai karakteristik yang sama satu sama lainnya dan mempunyai karakteristik yang berbeda dengan data yang ada didalam kelompok yang lain

K-Means merupakan salah satu algoritma dalam *Data Mining* yang bisa digunakan untuk melakukan pengelompokan/*Clustering* suatu data. Ada banyak pendekatan untuk membuat *cluster*, diantaranya adalah membuat aturan yang mendikte keanggotaan dalam group yang sama berdasarkan tingkat persamaan diantara anggota-anggotanya [12]. Pendekatan lainnya adalah dengan membuat sekumpulan fungsi yang mengukur beberapa properti dari pengelompokan tersebut sebagai fungsi dari beberapa parameter dari sebuah *Clustering*. Metode *K-Means* adalah metode yang termasuk dalam algoritma *Clustering* berbasis jarak yang membagi data ke dalam sejumlah *cluster* dan algoritma ini hanya bekerja pada atribut numerik.

Metode *K-Means* adalah Metode *Clustering* berbasis jarak yang membagi data kedalam *cluster* dan algoritma ini bekerja pada atribut numerik. Metode *K-Means* termasuk dalam *partitioning Clustering* yang memisahkan data ke daerah bagian yang terpisah. Dalam metode *K-Means* setiap data harus termasuk ke *cluster* tertentu pada suatu tahapan proses, pada tahapan proses berikutnya dapat berpindah ke *cluster* yang lain.

Pengelompokan data dengan metode *K-Means* dilakukan dengan algoritma:

1. Tentukan jumlah kelompok.
2. Alokasikan data ke dalam kelompok secara acak
3. Hitung pusat kelompok (*centroid*/rata-rata) dari data yang ada di masing-masing kelompok. Lokasi *centroid* setiap kelompok diambil dari rata-rata (*mean*) semua nilai data pada setiap fiturnya. Jika M menyatakan jumlah data dalam sebuah kelompok, i menyatakan fitur ke-i dalam sebuah kelompok, dan p menyatakan dimensi data, maka persamaan untuk menghitung *centroid* fitur ke-i digunakan persamaan 1.

$$C_i = \frac{1}{M} \sum_{j=1}^M X_j \dots\dots\dots (1)$$

persamaan 1 dilakukan sebanyak p dimensi dari i=1 sampai dengan i=p.

4. Alokasikan masing-masing data ke *centroid*/rata-rata terdekat. Ada beberapa cara yang dapat dilakukan untuk mengukur jarak data ke pusat kelompok, diantaranya adalah *Euclidean*. Pengukuran jarak pada ruang jarak (*distance space*) *Euclidean* dapat dicari menggunakan persamaan 2.

$$d = \sqrt{(X_1 - X_2)^2 + (Y_1 - Y_2)^2} \dots\dots\dots (2)$$

Pengalokasian kembali data ke dalam masing-masing kelompok dalam metode *K-Means* didasarkan pada perbandingan jarak antara data dengan *centroid* setiap kelompok yang ada. Data dialokasikan ulang secara tegas ke kelompok yang mempunyai *centroid* dengan jarak terdekat dari data tersebut. Pengalokasian data ini menurut MacQueen (1967) dapat ditentukan menggunakan persamaan 3.

$$a = \begin{cases} 1 & d = \min\{D(x,c)\} \\ 0 & \text{lainnya} \end{cases} \dots\dots\dots (3)$$

a_{i1} adalah nilai keanggotaan titik x_i ke pusat kelompok c_1 , d adalah jarak terpendek dari data x_i ke K kelompok setelah dibandingkan, dan c_1 adalah *centroid* (pusat kelompok) ke-1. Fungsi objektif yang digunakan untuk metode *K-Means* ditentukan berdasarkan jarak dan nilai keanggotaan data dalam kelompok.

$$J = \sum_{i=1}^n \sum_{c=1}^k A_{ic} D(x, c)^2 \dots\dots\dots (4)$$

n adalah jumlah data, k adalah jumlah kelompok, a_{i1} adalah nilai keanggotaan titik data x_i ke kelompok c_1 yang diikuti a mempunyai nilai 0 atau 1. Apabila data merupakan anggota suatu kelompok, nilai $a_{i1} = 1$. Jika tidak, nilai $a_{i1} = 0$

5. Menghitung nilai rasio antara *Between-Cluster Variation* (BCV) dan *Within-Cluster Variation* (WCV)

$$BCV = d(m_1 - m_2) + d(m_1 - m_3) + d(m_2 - m_3) \dots\dots\dots (5)$$

Kembali ke langkah 3, apabila masih ada data yang berpindah kelompok atau apabila ada perubahan nilai *centroid* di atas nilai ambang yang ditentukan, atau apabila perubahan nilai pada fungsi objektif yang digunakan masih di atas nilai ambang yang ditentukan.

3. HASIL DAN PEMBAHASAN

Dalam konsep penulisan pengembangan sistem merupakan salah satu unsur penting dalam metode penelitian. Dalam metode pengembangan sistem, khususnya *software* atau perangkat lunak dapat diadopsi dari beberapa metode diantaranya algoritma *Waterfall* atau algoritma air terjun. Berikut ini adalah fase yang dilakukan dalam penelitian ini, yaitu:

1. Analisis Masalah dan Kebutuhan

Pada tahapan ini dilakukan dengan wawancara ke Rumah Sakit Umum Seger Waras. Dimana pada tahap ini dilakukan dengan mencari permasalahan dan persoalan-persoalan tentang pengelompokan data penilaian pasien untuk membentuk tingkat kepuasan pasien.

2. Perancangan Sistem dan Pemodelan

Tahap perancangan dan pemodelan berfokus pada struktur data, arsitektur perangkat lunak, *Representasi Interface*, dan *Detail* prosedural. Pada tahapan ini dirancanglah tampilan program dan *database* yang akan digunakan pada sistem. Yang sebelumnya telah dimodelkan dengan menggunakan *Unified Modelling Language* (UML).

3. Pengkodean

Pengkodean dilakukan dengan menerjemahkan hasil dari perancangan dan pemodelan kedalam bahasa pemrograman berbasis *Web* agar dapat dikenali oleh komputer.

4. Uji Coba

Fase ini merupakan fase terpenting untuk pembangunan data mining. Hal ini dikarenakan pada fase ini akan dilakukan *Trial and Error* terhadap keseluruhan aspek aplikasi baik *Coding*, desain sistem dan pemodelan.

5. Implementasi Sistem

Pada tahapan ini memperlihatkan kinerja aplikasi yang di kembangkan dan sejauh mana sistem dapat bekerja dalam pengelompokan data kepuasan pasien menggunakan metode *K-Means*. Dalam penelitian ini pengguna adalah *end user*.

Algoritma sistem adalah langkah-langkah yang dilakukan oleh sebuah sistem dalam perancangan data mining untuk pengelompokan data kepuasan pasien menggunakan metode *K-Means*. Hal ini dilakukan untuk meningkatkan produktifitas dan keberhasilan perusahaan dalam dunia teknologi. Setelah data dikumpulkan, selanjutnya merupakan tahapan analisa dan pembahasan. Data mining digunakan untuk pengelompokan data kepuasan pasien menggunakan metode *K-Means*. Perhitungan menggunakan metode *K-Means* digunakan untuk menganalisa pengelompokan tingkat kepuasan pasien terhadap pelayanan rumah sakit. Adapun kerangka metode Algoritma *K-Means* adalah sebagai berikut :



Gambar 2. Kerangka Metode Algoritma *K-Means*

Berikut ini adalah data kepuasan pasien yang diperoleh dari Rumah Sakit Umum Seger Waras. Pada pengisian kuesioner terdapat 5 pertanyaan yang diisi oleh responden. Adapun atribut yang digunakan ada 6 atribut yaitu nama pasien, sarana prasarana rumah sakit, keamanan rumah sakit, kebersihan rumah sakit, keramahan petugas, dan memahami kebutuhan pasien. Penelitian ini mengambil beberapa atribut yang akan digunakan sebagai variabel dalam perhitungan, berikut ini adalah atribut yang digunakan sebagai variabel:

Tabel 2. Variabel Penilaian

NO	Atribut	Variabel
1.	Sarana Prasarana Rumah Sakit	K1
2.	Keamanan Rumah Sakit	K2
3.	Kebersihan Rumah Sakit	K3
4.	Keramahan Petugas	K4
5.	Memahami Kebutuhan Pasien	K5

Maka dari pengumpulan data yang dilakukan melalui kuesioner perlu dilakukan pembentukan nilai terlebih dahulu yang akan digunakan sebagai acuan dari ukuran kepuasan yang diberikan responden. Pemberian nilai kepuasan diurutkan dari nilai terendah hingga ke nilai tertinggi. Ukuran kepuasan pasien dapat diukur dari tabel berikut :

Tabel 3. Sekala Pengukuran

No	Kepuasan	Score
1.	Tidak Setuju	2
2.	Kurang Setuju	4
3.	Cukup Setuju	6
4.	Setuju	8
5.	Sangat Setuju	10

Setelah data terkumpul, dilakukan perhitungan sebagai berikut:

a. Iterasi ke-1

- Penentuan pusat awal cluster

Untuk menentukan pusat (*centroid*) awal ditentukan dengan mengacak (*random*) dari data nilai yang sudah ada. Pada kasus ini pusat (*centroid*) awalnya adalah:

Tabel 4. Titik Pusat Awal Cluster

NO	Centroid	K1	K2	K3	K4	K5
8	C1	8	6	6	6	6
11	C2	8	6	8	8	8
25	C3	4	6	4	6	6

- Perhitungan jarak pusat cluster

Perhitungan jarak dari data-1 terhadap pusat cluster adalah

$$D(1,C1) = \sqrt{(4-8)^2 + (6-6)^2 + (8-6)^2 + (6-6)^2 + (6-6)^2} = 4,472$$

$$D(1,C2) = \sqrt{(4-8)^2 + (6-6)^2 + (8-8)^2 + (6-8)^2 + (6-8)^2} = 4,898$$

$$D(1,C3) = \sqrt{(4-4)^2 + (6-6)^2 + (8-4)^2 + (6-6)^2 + (6-6)^2} = 4,000$$

Perhitungan jarak dari data-2 terhadap pusat cluster adalah

$$D(2,C1) = \sqrt{(4-8)^2 + (6-6)^2 + (8-6)^2 + (6-6)^2 + (4-6)^2} = 4,898$$

$$D(2,C2) = \sqrt{(4-8)^2 + (6-6)^2 + (8-8)^2 + (6-8)^2 + (4-8)^2} = 6,000$$

$$D(2,C3) = \sqrt{(4-4)^2 + (6-6)^2 + (8-4)^2 + (6-6)^2 + (4-6)^2} = 4,472$$

Dan selanjutnya, perhitungan jarak dilanjutkan untuk setiap data ke-3 dan data seterusnya terhadap pusat cluster. Kemudian akan menghasilkan perhitungan jarak antara setiap data terhadap pusat cluster sebagai berikut:

Tabel 5. Hasil Perhitungan Jarak Setiap Data Terhadap Pusat *Cluster* Iterasi Ke-1

No	C1	C2	C3	Jarak Terdekat	Cluster
1	4,472	4,898	4,000	4,000	C3
2	4,898	6,000	4,472	4,472	C3
3	4,472	4,000	4,898	4,000	C2
4	2,000	4,000	2,828	2,000	C1
5	2,828	2,000	5,291	2,000	C2
6	4,472	4,000	6,324	4,000	C2
7	4,472	6,324	5,656	4,472	C1
8	0,000	3,464	4,472	0,000	C1
9	4,000	4,472	5,291	4,000	C1
10	2,828	4,472	6,000	2,828	C1
11	3,464	0,000	6,324	0,000	C2
12	3,464	6,324	4,898	3,464	C1
13	5,656	6,000	5,291	5,291	C3
14	4,898	5,291	6,633	4,898	C1
15	6,000	8,000	2,828	2,828	C3
16	3,464	5,656	4,000	3,464	C1
17	5,291	4,898	4,898	4,898	C2
18	4,472	6,324	2,828	2,828	C3
19	4,472	4,898	2,828	2,828	C3
20	3,464	4,898	6,324	3,464	C1
21	6,000	5,656	6,324	5,656	C2
22	2,828	2,000	6,000	2,000	C2
23	3,464	2,828	6,324	2,828	C2
24	3,464	4,000	4,898	3,464	C1
25	4,472	6,324	0,000	0,000	C3
26	2,828	3,464	5,291	2,828	C1
27	2,828	3,464	6,000	2,828	C1
28	4,472	6,324	2,828	2,828	C3
29	3,464	4,000	4,898	3,464	C1
30	2,000	2,828	4,898	2,000	C1
31	3,464	2,828	4,000	2,828	C2
32	3,464	4,898	5,656	3,464	C1
33	4,898	6,633	3,464	3,464	C3
34	3,464	4,000	4,000	3,464	C1
35	2,828	3,464	5,291	2,828	C1
36	6,324	6,000	4,472	4,472	C3
37	4,472	2,828	5,656	2,828	C2
38	2,000	2,828	5,656	2,000	C1
39	2,828	3,464	3,464	2,828	C1
40	6,928	7,745	4,472	4,472	C3

Keterangan :

- C1 memiliki 19 anggota, yaitu data ke-4, 7, 8, 9, 10, 12, 14, 16, 20, 24, 26, 27, 29, 30, 32, 34, 35, 38, 39.
- C2 memiliki 10 anggota, yaitu data ke-3, 5, 6, 11, 17, 21, 22, 23, 31, 37.
- C2 memiliki 11 anggota, yaitu data ke-1, 2, 13, 15, 18, 19, 25, 28, 33, 36, 40.

- Penentuan pusat *cluster* baru

Setelah diperoleh anggota dari setiap *cluster*, selanjutnya menghitung pusat *cluster* baru berdasarkan data anggota dari setiap *cluster* yang telah ditemukan. Perhitungan ini menggunakan rumus yang sesuai dengan pusat

anggota *cluster* yaitu : $V = \frac{\sum_{i=1}^n x_i}{n}$; i = 1,2,3,...n.

Sehingga didapatkan perhitungan sebagai berikut :

- C1 memiliki 19 anggota, maka perhitungan pusat *cluster* baru menjadi:

$$C1 = \left(\begin{array}{c} \frac{6+6+8+6+8+8+8+6+6+10+8+8+8+6+8+8+6+8+8+6}{19} ; \\ \frac{6+8+6+6+6+4+8+4+6+6+8+8+4+6+8+8+8+6+6}{19} ; \\ \frac{6+8+6+8+8+4+6+6+4+4+6+8+8+6+6+6+6+8+6}{19} ; \\ \frac{6+4+6+8+6+4+4+4+8+8+6+6+6+8+4+6+8+6+8}{19} ; \\ \frac{6+4+6+4+4+6+10+6+6+8+8+6+6+6+8+8+6+6+6}{19} \end{array} \right)$$

$$C1 = (7,368 ; 6,421 ; 6,315 \ 6,105 ; 6,315)$$

- C2 memiliki 10 anggota, maka perhitungan pusat *cluster* baru menjadi:

$$C2 = \left(\begin{array}{c} \frac{6+8+8+8+4+4+8+8+6+6}{10} ; \\ \frac{6+6+4+6+4+4+6+4+6+4}{10} ; \\ \frac{6+6+6+8+8+10+8+8+6+8}{10} ; \\ \frac{6+8+6+8+8+6+6+8+8+8}{10} ; \\ \frac{10+8+10+8+6+6+8+6+8+8}{10} \end{array} \right)$$

$$C2 = (6,600 ; 5,000 ; 7,400 ; 7,200 ; 7,800)$$

- C3 memiliki 11 anggota, maka perhitungan pusat *cluster* baru menjadi:

$$C3 = \left(\begin{array}{c} \frac{4+4+4+4+4+4+4+4+4+4+2}{11} ; \\ \frac{6+6+8+8+6+6+6+6+8+6+8}{11} ; \\ \frac{8+8+8+2+6+6+4+6+6+4+6}{11} ; \\ \frac{6+6+4+6+4+6+6+6+4+10+4}{11} ; \\ \frac{6+4+8+6+6+8+6+4+6+8+8}{11} \end{array} \right)$$

$$C3 = (3,818 ; 6,727 ; 5,818 ; 5,63 ; 6,363)$$

Tabel 6. Titik Pusat *Cluster*

Centroid	K1	K2	K3	K4	K5
C1	7,368	6,421	6,315	6,105	6,315
C2	6,6	5,0	7,4	7,2	7,8
C3	3,818	6,727	5,818	5,636	6,363

b. Iterasi ke-2

Iterasi berikutnya dilakukan dengan menggunakan langkah yang sama hingga tidak ada perubahan data yang terjadi dalam *cluster*.

- Perhitungan jarak pusat *cluster*

Perhitungan jarak dari data-1 terhadap pusat *cluster* adalah

$$D(1,C1) = \sqrt{(4 - 7,368)^2 + (6 - 6,421)^2 + (8 - 6,315)^2 + (6 - 6,105)^2 + (6 - 6,315)^2}$$

$$= 3,805$$

$$D(1,C2) = \sqrt{(4 - 6,6)^2 + (6 - 5,0)^2 + (8 - 7,4)^2 + (6 - 7,2)^2 + (6 - 7,8)^2}$$

$$= 3,577$$

$$D(1,C3) = \sqrt{(4 - 3,818)^2 + (6 - 6,727)^2 + (8 - 5,818)^2 + (6 - 5,636)^2 + (6 - 6,363)^2}$$

$$= 2,363$$

Perhitungan jarak dari data-2 terhadap pusat *cluster* adalah

$$D(2,C1) = \sqrt{(4 - 7,368)^2 + (6 - 6,421)^2 + (8 - 6,315)^2 + (6 - 6,105)^2 + (4 - 6,315)^2}$$

$$= 4,441$$

$$D(2,C2) = \sqrt{(4 - 6,6)^2 + (6 - 5,0)^2 + (8 - 7,4)^2 + (6 - 7,2)^2 + (4 - 7,8)^2}$$

$$= 4,898$$

$$D(2,C3) = \sqrt{(4 - 3,818)^2 + (6 - 6,727)^2 + (8 - 5,818)^2 + (6 - 5,636)^2 + (4 - 6,363)^2}$$

$$= 3,322$$

Dan selanjutnya, perhitungan jarak dilanjutkan untuk setiap data ke-3 dan data seterusnya terhadap pusat *cluster*. Kemudian akan menghasilkan perhitungan jarak antara setiap data terhadap pusat *cluster* sebagai berikut:

Tabel 7. Hasil Perhitungan Jarak Setiap Data Terhadap Pusat *Cluster* Iterasi Ke-2

No	C1	C2	C3	Jarak Terdekat	Cluster
1	3,805	3,577	2,363	2,363	C3
2	4,441	4,898	3,322	3,322	C3
3	3,966	3,098	4,322	3,098	C2
4	1,502	2,828	2,362	1,502	C1
5	2,665	2,366	5,130	2,366	C2
6	4,466	3,346	6,190	3,346	C2
7	4,122	5,865	4,404	4,122	C1
8	0,887	3,098	4,279	0,887	C1
9	3,719	4,098	4,606	3,719	C1
10	2,963	4,381	5,337	2,963	C1
11	3,137	2,000	5,571	2,000	C2
12	4,018	5,291	5,571	4,018	C1
13	4,893	5,138	3,430	3,430	C3
14	4,582	5,291	5,919	4,582	C1
15	5,706	7,042	4,061	4,061	C3
16	3,515	4,098	3,877	3,515	C1
17	4,871	3,464	4,236	3,464	C2
18	4,018	4,816	1,844	1,844	C3
19	3,803	3,346	1,845	1,845	C3
20	4,019	5,291	6,911	4,019	C1
21	5,558	4,381	5,022	4,381	C2
22	2,502	2,190	5,058	2,190	C2
23	3,576	2,683	5,950	2,683	C2
24	3,516	3,898	5,439	3,516	C1
25	4,121	4,898	2,031	2,031	C3
26	2,416	3,794	4,685	2,416	C1
27	2,416	4,000	4,912	2,416	C1
28	4,123	5,059	2,511	2,511	C3
29	3,268	2,529	4,149	2,529	C2
30	2,089	2,966	4,876	2,089	C1
31	2,928	2,000	3,686	2,000	C2
32	3,203	4,816	4,949	3,203	C1
33	4,297	5,585	2,119	2,119	C3
34	2,704	3,577	3,037	2,704	C1
35	2,584	4,098	4,986	2,584	C1
36	5,906	5,215	5,058	5,058	C3
37	4,123	1,549	5,022	1,549	C2
38	1,877	2,828	4,800	1,877	C1
39	2,415	2,683	3,322	2,415	C1
40	6,218	6,511	3,211	3,211	C3

Keterangan :

- C1 memiliki 18 anggota, yaitu data ke-4, 7, 8, 9, 10, 12, 14, 16, 20, 24, 26, 27, 30, 32, 34, 35, 38, 39.
- C2 memiliki 11 anggota, yaitu data ke-3, 5, 6, 11, 17, 21, 22, 23, 29, 31, 37.
- C3 memiliki 11 anggota, yaitu data ke-1, 2, 13, 15, 18, 19, 25, 28, 33, 36, 40.

- Penentuan pusat *cluster* baru

Setelah diperoleh anggota dari setiap *cluster*, selanjutnya menghitung pusat *cluster* baru berdasarkan data anggota dari setiap *cluster* yang telah ditemukan. Perhitungan ini menggunakan rumus yang sesuai dengan

pusat anggota *cluster* yaitu : $V = \frac{\sum_{i=1}^n x_i}{n}$; i = 1,2,3,...n.

Sehingga didapatkan perhitungan sebagai berikut :

- C1 memiliki 18 anggota, maka perhitungan pusat *cluster* baru menjadi:

$$C1 = \left(\begin{array}{c} \frac{6+6+8+6+8+8+8+6+10+8+8+8+8+6+8+8+6}{18} ; \\ \frac{6+8+6+6+6+4+8+4+6+6+8+8+6+8+8+6+6}{18} ; \\ \frac{6+8+6+8+8+4+6+6+4+4+6+8+6+6+6+6+8+6}{18} ; \\ \frac{6+4+6+8+6+4+4+4+8+8+6+6+8+4+6+8+6+8}{18} ; \\ \frac{6+4+6+4+4+6+10+6+6+8+8+6+6+8+8+6+6+6}{18} \end{array} \right)$$

$$C1 = (7,444 ; 6,555 ; 6,222 ; 6,111 ; 6,333)$$

- C2 memiliki 11 anggota, maka perhitungan pusat *cluster* baru menjadi:

$$C2 = \left(\begin{array}{c} \frac{6+8+8+8+4+4+8+8+6+6+6}{11} ; \\ \frac{6+6+4+6+4+4+6+4+4+6+4}{11} ; \\ \frac{6+6+6+8+8+10+8+8+6+8}{11} ; \\ \frac{6+8+6+8+8+6+6+8+6+8+8}{11} ; \\ \frac{10+8+10+8+6+6+8+6+6+8+8}{11} \end{array} \right)$$

$$C2 = (6,545 ; 4,909 ; 7,454 ; 7,091 ; 7,636)$$

- C3 memiliki 11 anggota, maka perhitungan pusat *cluster* baru menjadi:

$$C3 = \left(\begin{array}{c} \frac{4+4+4+4+4+4+4+4+4+4+2}{11} ; \\ \frac{6+6+8+8+6+6+6+6+8+6+8}{11} ; \\ \frac{8+8+8+2+6+6+4+6+6+4+6}{11} ; \\ \frac{6+6+4+6+4+6+6+6+4+10+4}{11} ; \\ \frac{6+4+8+6+6+8+6+4+6+8+8}{11} \end{array} \right)$$

$$C3 = (3,818 ; 6,727 ; 5,818 ; 5,63 ; 6,363)$$

Tabel 3.9 Titik Pusat *Cluster*

Centroid	K1	K2	K3	K4	K5
C1	7,444	6,555	6,222	6,111	6,333
C2	6,545	4,909	7,454	7,091	7,636
C3	3,818	6,727	5,818	5,636	6,363

c. Iterasi ke-3

Iterasi berikutnya dilakukan dengan menggunakan langkah yang sama hingga tidak ada perubahan data yang terjadi dalam *cluster*.

- Perhitungan jarak pusat *cluster*

Perhitungan jarak dari data-1 terhadap pusat *cluster* adalah

$$D(1,C1) = \sqrt{(4 - 7,444)^2 + (6 - 6,555)^2 + (8 - 6,222)^2 + (6 - 6,111)^2 + (6 - 6,333)^2} = 3,930$$

$$D(1,C2) = \sqrt{(4 - 6,545)^2 + (6 - 4,909)^2 + (8 - 7,454)^2 + (6 - 7,090)^2 + (6 - 7,636)^2} = 3,439$$

$$D(1,C3) = \sqrt{(4 - 3,818)^2 + (6 - 6,727)^2 + (8 - 5,818)^2 + (6 - 5,636)^2 + (6 - 6,363)^2} = 2,363$$

Perhitungan jarak dari data-2 terhadap pusat *cluster* adalah

$$D(2,C1) = \sqrt{(4 - 7,444)^2 + (6 - 6,555)^2 + (8 - 6,222)^2 + (6 - 6,111)^2 + (4 - 6,333)^2} = 4,558$$

$$D(2,C2) = \sqrt{(4 - 6,545)^2 + (6 - 4,909)^2 + (8 - 7,455)^2 + (6 - 7,091)^2 + (4 - 7,636)^2} = 4,730$$

$$D(2,C3) = \sqrt{(4 - 3,818)^2 + (6 - 6,727)^2 + (8 - 5,818)^2 + (6 - 5,636)^2 + (4 - 6,363)^2} = 3,322$$

Dan selanjutnya, perhitungan jarak dilanjutkan untuk setiap data ke-3 dan data seterusnya terhadap pusat *cluster*. Kemudian akan menghasilkan perhitungan jarak antara setiap data terhadap pusat *cluster* sebagai berikut:

Tabel 8. Hasil Perhitungan Jarak Setiap Data Terhadap Pusat *Cluster* Iterasi Ke-3

No	C1	C2	C3	Jarak Terdekat	Cluster
1	3,930	3,439	2,363	2,363	C3
2	4,559	4,730	3,323	3,323	C3
3	3,987	3,222	4,322	3,222	C2
4	1,601	2,733	2,362	1,601	C1
5	2,648	2,526	5,130	2,526	C2
6	4,510	3,440	6,190	3,440	C2
7	4,151	5,737	4,404	4,151	C1
8	0,887	3,047	4,279	0,887	C1
9	3,816	3,978	4,606	3,816	C1
10	3,038	4,244	5,337	3,038	C1
11	3,181	2,135	5,571	2,135	C2
12	4,042	5,206	5,571	4,042	C1
13	4,934	5,100	3,430	3,430	C3
14	4,511	5,378	5,919	4,511	C1
15	5,647	7,046	4,061	4,061	C3
16	3,637	3,933	3,877	3,637	C1
17	5,022	3,331	4,236	3,331	C2
18	4,097	4,691	1,844	1,844	C3
19	3,874	3,332	1,845	1,845	C3
20	3,931	5,344	6,911	3,931	C1
21	5,725	4,200	5,022	4,201	C2
22	2,562	2,219	5,058	2,219	C2
23	3,698	2,595	5,950	2,597	C2
24	3,449	4,025	5,439	3,449	C1
25	4,150	4,844	2,031	2,031	C3
26	2,288	3,887	4,685	2,288	C1
27	2,383	3,978	4,912	2,383	C1
28	4,203	4,918	2,511	2,511	C3
29	3,449	2,298	4,149	2,298	C2
30	2,084	2,987	4,875	2,084	C1
31	2,964	2,135	3,686	2,135	C2
32	3,111	4,844	4,949	3,111	C1
33	4,308	5,511	2,119	2,119	C3
34	2,648	3,645	3,037	2,648	C1
35	2,474	4,158	4,986	2,474	C1
36	5,916	5,310	5,058	5,058	C3
37	4,256	1,542	5,022	1,542	C2
38	1,974	2,733	4,800	1,974	C1
39	2,473	2,665	3,322	2,473	C1
40	6,245	6,481	3,211	3,211	C3

Keterangan :

- C1 memiliki 18 anggota, yaitu data ke-4, 7, 8, 9, 10, 12, 14, 16, 20, 24, 26, 27, 30, 32, 34, 35, 38, 39.
- C2 memiliki 11 anggota, yaitu data ke-3, 5, 6, 11, 17, 21, 22, 23, 29, 31, 37.
- C3 memiliki 11 anggota, yaitu data ke-1, 2, 13, 15, 18, 19, 25, 28, 33, 36, 40.

d. Kesimpulan

Berdasarkan hasil perhitungan pada iterasi ke-3 didapatkan hasil pengelompokan yang sama pada iterasi ke-2. Maka perhitungan dihentikan dan selesai sampai tahap iterasi ke-3. Karena iterasi telah mencapai akhir atau selesai maka dari perhitungan tersebut diperoleh kesimpulan sebagai berikut:

Tabel 9. Data Hasil *Cluster*

No	Nama	Cluster
1	YANNA FITRIYANA	Tidak Puas
2	MUSLIMIN	Tidak Puas
3	RATIH	Cukup
4	IDA SAFITRI	Puas
5	DELLA ALFITA	Cukup
6	SINTIA NANDA UTAMI	Cukup
7	MUSRIFA	Puas
8	JUNAIDA SARI	Puas
9	AYU LESTARI	Puas
10	YULINDA	Puas
11	ANGGI PERTIWI	Cukup
12	RINDI ANDRIANI	Puas
13	ELI AROSA	Tidak Puas
14	CAHAYA SARTIKA	Puas
15	SRI WINDARI	Tidak Puas
16	ENDORO HANDOKO SIJABAT	Puas
17	FIKA ELIZAH	Cukup
18	ROMAULI ARITONANG	Tidak Puas
19	FADLY RIFKYANSYAH	Tidak Puas
20	ENI SUSANTI	Puas
21	RATI SERVIA	Cukup
22	FARIDAH HANNUM	Cukup
23	TRIA WANTI	Cukup
24	SRI WAHYUNI	Puas
25	IRMAYANI	Tidak Puas
26	FAJRIN RAMADINA	Puas
27	ERMINA MARIANA	Puas
28	NIA RAMADHANI	Tidak Puas
29	EKA SIAHAAN	Cukup
30	YULIANY	Puas
31	MAYA ROMANTI	Cukup
32	SUSIYAN	Puas
33	PINKA ERAWATI	Tidak Puas
34	RINDI PUSPITA SARI	Puas
35	MAHINDA	Puas
36	DWI ASTUTI	Tidak Puas
37	DIANA LESTARI	Cukup
38	RISMAWATI	Puas
39	RATNA PUTRI	Puas
40	SITI AMINAH	Tidak Puas

Keterangan:

1. Pada Klasterisasi 1 (C1) terdapat 18 anggota dengan penilaian puas, yaitu 4, 7, 8, 9, 10, 12, 14, 16, 20, 24, 26, 27, 30, 32, 34, 35, 38, 39.
2. Pada Klasterisasi 2 (C2) terdapat 11 anggota dengan penilaian cukup, yaitu 3, 5, 6, 11, 17, 21, 22, 23, 29, 31, 37.
3. Pada Klasterisasi 3 (C3) terdapat 11 anggota dengan penilaian tidak puas, yaitu 1, 2, 13, 15, 18, 19, 25, 28, 33, 36, 40.

4. KESIMPULAN

Berdasarkan analisa pada permasalahan yang terjadi dalam kasus yang diangkat dalam menganalisa kepuasan pasien terhadap pelayanan Rumah Sakit Umum Seger Waras, dalam menganalisis kepuasan pasien terhadap pelayanan di Rumah Sakit Umum Seger Waras, pertama-tama perlu dilakukan pengumpulan data melalui survei atau kuesioner yang mencakup berbagai aspek layanan seperti kualitas perawatan, kecepatan pelayanan, sikap staf medis, dan fasilitas rumah sakit. Data ini kemudian dianalisis untuk mengidentifikasi tingkat kepuasan pasien. Metode analisis bisa meliputi analisis deskriptif untuk mendapatkan gambaran umum dan analisis statistik untuk memahami pola dan hubungan antara variabel. Berdasarkan hasil algoritma k-menas clustering dalam mengelompokkan tingkat kepuasan pasien diperoleh tiga kelompok cluster berbeda, yaitu cluster puas, cluster cukup, dan cluster tidak puas. Analisis lebih lanjut terhadap setiap cluster mengungkapkan bahwa 18 menyatakan puas, 11 menyatakan cukup, dan 11 menyatakan tidak puas dari 40 pasien. Berdasarkan penelitian yang telah dilakukan, untuk melakukan analisis, pertama-tama data tanggapan pasien harus dipersiapkan dan dibersihkan. Kemudian, data tersebut diubah menjadi format numerik jika diperlukan dan dinormalisasi. Metode K-Means digunakan untuk mengelompokkan data ke dalam beberapa cluster berdasarkan kemiripan dalam tanggapan pasien. Analisis ini membantu mengidentifikasi kelompok-kelompok pasien dengan tingkat kepuasan yang serupa atau masalah yang sama, sehingga memudahkan penyesuaian strategi perbaikan. Merancang sistem kepuasan pasien melibatkan pengembangan mekanisme untuk mengumpulkan, menganalisis, dan melaporkan data kepuasan pasien secara efektif. Sistem ini harus mencakup pembuatan formulir survei digital atau fisik, proses pengumpulan data yang efisien, dan alat analisis data untuk mengevaluasi hasil survei. Sistem harus dirancang dengan antarmuka pengguna yang ramah dan mudah diakses oleh pasien serta fitur pelaporan untuk manajemen rumah sakit. Dengan adanya sistem ini, Rumah Sakit Umum Seger Waras dapat terus memantau dan meningkatkan kualitas pelayanan berdasarkan umpan balik pasien.

UCAPAN TERIMA KASIH

Terima kasih disampaikan kepada pihak-pihak yang telah mendukung terlaksananya penelitian ini. Yaitu Bapak M. Gilang Suryanata, S.Kom., M.Kom dan Bapak Tugiono, S.Kom., M.Kom.

Referensi

- [1] A. H. R. A. G. Castaka Agus Sugianto, "algoritme k means untuk pengelompokan penyakit pasien pada puskesmas cigugur tengah," *Journal of Information Technology*, p. 39, 2020.
- [2] G. U. A. S. Ulfhi Burelia, "MENGUKUR TINGKAT KEPUASAN MASYARAKAT PADA PELAYANAN KEPOLISIAN RESOR(POLRES) DUMAI MENGGUNAKAN ALGORITMA K-MEANS CLUSTERING," *JUTEKINF*, p. 17, 2022.
- [3] F. I. E. B. I. A. Aqshol Al Fahrozi, "Implementasi Algoritma K-Means dalam Menentukan Clustering pada Penelitian Kepuasan Pelanggan Kesehatan Pekanbaru," *Indonesian Journal of Innovati on Multidisipliner Research*, p. 490, 2023.
- [4] Y. N. M. Eben Haezer Wisanta, "Analisis Algoritma K-Means Untuk Clustering Kepuasan Pelayanan: Mall Pelayanan Publik Pekanbaru," *Seminar Nasional Informatika*, p. 227, 2021.
- [5] Yane Laheroi Nainel, Efori Buulolo, Ikwan Lubis, "Penerapan Data Mining Untuk Estimasi Penjualan Obat Berdasarkan Pengaruh Brand Image Dengan Algoritma Expectation Maximization (Studi Kasus: PT. Pyridam Farma Tbk)," *JURIKOM (Jurnal Riset Komputer)*, vol. 7, no. 2, 2020.
- [6] F. Alghifari, "Penerapan Data Mining Pada Penjualan Makanan Dan Minuman Menggunakan Metode Algoritma Naïve Bayes," *Jurnal Media Infotama*, vol. 9, no. 2, 2021.
- [7] P. M. S. Tarigan, "IMPLEMENTASI DATA MINING MENGGUNAKAN ALGORITMA APRIORI DALAM MENENTUKAN PERSEDIAAN BARANG (STUDI KASUS : TOKO SINAR HARAHAHAP)," *Semesta Teknika*, vol. 12, no. 2, 2022.
- [8] D. N. Sari, "Penerapan Data Mining Untuk Klasifikasi Gaya Belajar Siswa Menggunakan Algoritma C4.5," *Jurnal Smart Teknologi*, vol. 3, no. 4, p. 413 – 422, 26 9 2022.
- [9] Y. T. Hayaza, "Analisis kepuasan pasien terhadap kualitas pelayanan kamar obat di puskesmas surabaya utara," *Jurnal Ilmiah Mahasiswa Universitas Surabaya*, p. 3, 2013.
- [10] W. I. F. V. L. Harminto, "Analisis Kepuasan Pasien Terhadap Kualitas Pelayanan Umum di Klinik Cipinang Jakarta dengan Metode Servqual," *JURNAL MANAJEMEN FE-UB*, p. 103, 2021.
- [11] R. M. F. Y. Vera Sesrianty, "nalisa Kepuasan Pasien Terhadap Mutu Pelayanan Keperawatan," *Jurnal Kesehatan Perintis (Perintis's Health Journal)*, p. 117, 2019.
- [12] Haviluddin, "Implementasi Metode K-Means untuk Pengelompokkan Rekomendasi Tugas Akhir," *Informatika Mulawarman : Jurnal Ilmiah Ilmu Komputer*, vol. 16, no. 1, 2021.