

PENGELOMPOKAN DATA OBAT MENGGUNAKAN METODE *K-MEANS* CLUSTERING PADA UPT PUSKESMAS KONDORAN KEC.SANGALLA'

Efraim Harpendi Bara, Yosep Agus Pranoto, F.X Ariwibisono
Program Studi Teknik Informatika S1, Fakultas Teknologi Industri
Institut Teknologi Nasional Malang, Jalan Raya Karanglo km 2 Malang, Indonesia
eharpendi012@gmail.com

ABSTRAK

Dalam hal pengadaan dan pendistribusian obat di tempat pelayanan kesehatan dipengaruhi oleh perencanaan yang dilakukan. Agar obat dapat tersedia dengan jenis dan jumlah yang cukup sesuai dengan kebutuhan pelayanan kesehatan. Pada saat ini sebagian besar Puskesmas di Indonesia khususnya di daerah Tana Toraja belum melakukan kegiatan pelayanan farmasi seperti yang diharapkan, mengingat beberapa kendala antara lain kemampuan tenaga farmasi, terbatasnya pengetahuan manajemen Puskesmas, kebijakan manajemen institusi yang membawahi suatu Puskesmas, terbatasnya akses sarana dan prasarana internet. Akibat dari kondisi ini maka pelayanan farmasi Puskesmas masih bersifat konvensional dimana data yang diolah manual dinilai belum efektif dari segi waktu dan belum memberikan informasi yang jelas mengenai data obat yang harus diprioritaskan.

Penelitian ini akan mengembangkan sebuah sistem data mining yang dapat membantu proses pengelompokan data obat kedalam beberapa *cluster*, sehingga mengurangi lama waktu proses pengolahan data untuk rencana pembelian obat pada periode berikutnya, serta mengurangi penumpukan obat yang tidak cepat pakai yang mengakibatkan obat sering kadaluarsa pada gudang obat suatu puskesmas. Penelitian ini menggunakan metode *K-Means clustering* yaitu sebuah metode yang dapat digunakan untuk mengelompokkan data berdasarkan tingkat kemiripan dan karakteristik data.

Penerapan metode *K-Means* pada penelitian ini digunakan untuk mengolah data yang akan dimasukkan kedalam masing-masing 2 *cluster*, *cluster* yang pertama yaitu data obat dengan pemakaian lambat, dan yang kedua data obat dengan pemakaian cepat. Dataset yang digunakan adalah dataset yang didapatkan dari hasil penelitian dilapangan dimana terdapat 204 daset yang akan diuji pada sistem. Berdasarkan hasil pengujian pada sistem menggunakan metode *K-Means Clustering* didapatkan hasil bahwa 183 data obat termasuk kedalam *cluster* pertama dan 21 data obat termasuk kedalam *cluster* kedua.

Kata Kunci : Data Mining, *K-Means*, Clustering, Dekstop

1. PENDAHULUAN

Perkembangan pembangunan sebuah bangsa dapat ditentukan melalui banyak indikator salah satunya yaitu di bidang kesehatan. Kesehatan merupakan hak asasi dari setiap warga negara Indonesia. Perkembangan dunia dan teknologi juga mendorong perkembangan dalam dunia medis. Obat menjadi salah satu bagian yang tidak terpisahkan dari dunia medis itu sendiri. Perlunya pembangunan kesehatan secara menyeluruh dan berkesinambungan guna meningkatkan kesadaran, kemampuan dan kemauan hidup sehat agar setiap warga masyarakat mendapatkan akses serta pelayanan kesehatan yang sama rata. Kebutuhan obat tiap tahun mengalami peningkatan seiring dengan perkembangan jaman, terlebih dengan penemuan – penemuan penyakit jenis baru. Untuk itu obat perlu dikelola dengan efektif dan efisien.

Dalam pengelolaan obat, perencanaan kebutuhan merupakan salah satu aspek penting dan menentukan pengelolaan obat tersebut. Untuk menjamin ketersediaan dan pemerataan obat diperlukan perencanaan kebutuhan obat agar dapat diperoleh dengan cepat pada tempat dan waktu yang tepat. Dalam hal pengadaan, pendistribusian dan

pengadaan obat di tempat pelayanan kesehatan dipengaruhi oleh perencanaan yang dilakukan. Agar obat dapat tersedia dengan jenis dan jumlah yang cukup sesuai dengan kebutuhan pelayanan kesehatan.

Kemajuan pada bidang IT mendorong berkembangnya berbagai jenis *software* yang dapat membantu kehidupan manusia. Salah satunya adalah *software* yang menerapkan sistem data mining. Dalam penerapan data mining untuk klustering data terdapat beberapa algoritma yang dapat digunakan yaitu *K-Nearest Neighbors*, *Fuzzy C-Means*, dan *K-Means* dimana algoritma *K-Means* yang paling populer dan banyak digunakan. Karena *K-means* mempunyai kemampuan mengelompokkan data dalam jumlah yang cukup besar dengan waktu komputasi yang relatif cepat. Berdasarkan permasalahan diatas maka dalam skripsi ini dibuat sistem yang dapat membantu proses pengelompokan data obat kedalam beberapa *cluster*, sehingga mengurangi lama waktu proses pengolahan data untuk rencana pembelian obat pada periode berikutnya, serta mengurangi penumpukan obat yang tidak cepat pakai yang mengakibatkan obat sering kadaluarsa pada gudang obat suatu puskesmas.

2. TINJAUAN PUSTAKA

2.1 Penelitian Terdahulu

Penelitian yang dilakukan oleh Sulastris dan Gufroni pada tahun 2017 yang diberi judul Penerapan Data Mining dalam pengelompokan penderita Thalassemia memiliki tujuan untuk mengelompokkan penderita Thalassemia kedalam beberapa kelompok, agar persiapan sebelum memberikan jumlah obat terapi kelasi dan kesiapan dari pendonor darah dapat terjadwal dengan baik. Data set yang digunakan dalam penelitian tersebut berjumlah 2068 data. Dari data set tersebut peneliti melakukan data *cleaning* menghasilkan 78 data prematur, sehingga data baru yang dihasilkan adalah 1990 data. Setelah melakukan data *cleaning* kemudian peneliti melakukan seleksi data ulang sehingga memperoleh 374 data yang akan diproses. Metode yang digunakan oleh peneliti adalah K-Means *clustering*. Dari penelitian tersebut dihasilkan bahwa dari 374 data, 214 data termasuk kedalam cluster 1, pada cluster 2 berjumlah 137 dan cluster 3 sebanyak 23 data. [1]

Ditahun yang sama Astuti melakukan penelitian yang berjudul Penerapan data mining untuk clustering data penduduk miskin menggunakan algoritma Hard C-Means. Data yang digunakan peneliti yaitu data penduduk miskin kecamatan Bantul yang berjumlah 1313. Tujuan dari penelitian tersebut adalah agar penyaluran bantuan oleh BKKBN dapat tersalur dengan tepat. Dari data tersebut kemudian peneliti melakukan clustering data menggunakan metode Hard C-Means yang menghasilkan 253 keluarga termasuk kedalam cluster pertama, 763 pada cluster kedua dan 297 keluarga pada cluster ketiga. [2]

Berbeda dengan Astuti, ditahun sebelumnya yaitu 2016 Darmi dan Setiawan sudah melakukan penelitian yang berjudul Penerapan Metode Clustering K-Means dalam pengelompokan penjualan produk. Tujuan peneliti dalam melakukan penelitian ini adalah untuk memberikan informasi kepada pihak swalayan MM TIKKA bahwa produk mana saja yang laku dan mana yang kurang laku, sehingga pemesanan barang yang kurang laku dapat dikurangi. Hasil dari penelitian ini yaitu produk yang laku pada swalayan tersebut adalah makanan, dan minuman sedangkan yang tidak laku yaitu produk kosmetik. [3]

Pada tahun 2018 Fitra R.K Afiani, melakukan penelitian dengan mengelompokkan varietas padi unggul. Penelitian yang dilakukan menggunakan metode K-Means. Data yang digunakan peneliti dalam penelitian tersebut berasal dari Balai Pengkajian Teknologi Pertanian Jawa Timur. Peneliti menggunakan 3 klaster dalam penelitiannya yaitu kategori unggul, sedang, dan kurang. Pada penelitian tersebut iterasi penghitungan metode K-Means berlangsung sebanyak 3x. Dari penelitian tersebut didapatkan hasil dari 67 data didapatkan hasil 13 varietas padi berkategori Unggul, 33 Sedang dan 21

Kurang yang didapatkan dari hasil iterasi ketiga yaitu iterasi terakhir dalam perhitungannya. [4]

Kemudian pada tahun berikutnya yaitu 2019 Chusyairi et al, melakukan penelitian untuk mengelompokkan data puskesmas Banyuwangi dalam pemberian imunisasi. Tujuan penelitian tersebut adalah memberikan informasi bagi dinas terkait untuk memberikan tugas tambahan bagi kelompok Puskesmas yang memiliki target IDL dengan status kurang untuk mengurangi angka penyakit-penyakit yang dapat dicegah dengan imunisasi (PD2I). Data yang diolah oleh peneliti akan dimasukkan kedalam 3 *cluster* yaitu psukesmas yang mencapai target IDL dengan status cukup, kurang, dan status sangat baik. Hasil dari penelitian ini menunjukkan dari 45 puskesmas 19 data puskesmas dengan target imunisasi cukup, 24 data puskesmas dengan imunisasi kurang, dan 2 data puskesmas dengan target imunisasi sangat baik. [5]

2.2 Dasar Teori

2.2.1 Puskesmas

Puskesmas adalah sebuah unit organisasi yang bergerak di bidang pelayanan kesehatan di garda terdepan dan mempunyai misi sebagai pusat pengembangan pelayanan kesehatan, yang melaksanakan pembinaan dan pelayanan kesehatan secara menyeluruh dan terpadu untuk masyarakat di suatu wilayah kerja tertentu yang telah ditentukan secara mandiri dalam menentukan kegiatan pelayanan namun tidak mencakup aspek pembiayaan. [6]

2.2.2 Algoritma K-Means

K-Means merupakan salah satu dari beberapa metode *clustering* dengan cara kerja mempartisi data yang ada ke dalam bentuk satu atau lebih cluster/kelompok. Pembagian data ke dalam cluster/kelompok pada metode ini menggunakan data dengan karakteristik yang sama yang dikelompokkan ke dalam satu cluster yang sama. [7]

Langkah-langkah dalam algoritma K-Means :

- Identifikasikan data yang akan dikelompokkan dan tentukan *cluster* data, X_{ij} ($i=1,2,...,n$; $j=1,2,...,m$) dimana n adalah jumlah data yang akan di *cluster* dan m adalah jumlah variabel data
- Awal iterasi pusat setiap *cluster* ditentukan dengan cara acak.
- Tentukan jarak setiap data dengan pusat *cluster* menggunakan rumus *Euclidean* berikut :
- Kelompokkan data berdasarkan jarak terdekat data terhadap pusat *centroid* awal
- Selanjutnya cari pusat *cluster* baru dengan menghitung pusat *cluster* berdasarkan nilai rata-rata dengan rumus :

$$C_k = \frac{1}{n_k} \sum d_i$$

Persamaan 2

Dimana n_k adalah jumlah data dalam cluster D_i jumlah dari nilai jarak yang masuk dalam masing-masing *cluster*

- Cek kondisi apakah pusat *cluster* berubah jika ya maka iterasi dilanjutkan dan hitung ulang masing-masing data terhadap nilai *cluster* baru
- Jika tidak terdapat perbedaan pada pusat *cluster* baru dan lama maka iterasi dihentikan.

2.2.3 Clustering

Clustering merupakan pengelompokan data berdasarkan kemiripan. Data dengan tingkat kemiripan tinggi akan dikelompokkan dalam sebuah *cluster*.

Clustering membagi data ke dalam grup dengan objek yang memiliki karakteristik sama. Setiap data dimasukkan kedalam cluster dengan karakteristik yang paling mirip lalu kemudian berusaha menjauhkannya dengan cluster yang lain. [9]. *Clustering* banyak digunakan untuk mengelompokkan data dengan tujuan dan manfaat tertentu tergantung dengan kebutuhan

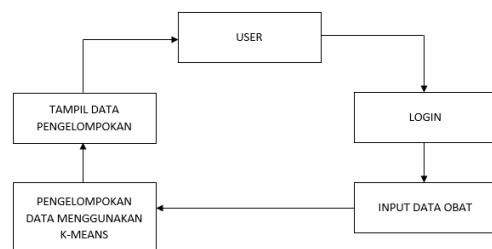
3. METODE PENELITIAN

3.1 Analisis Sistem

Sistem yang akan dibangun pada sistem pengelompokan data obat menggunakan metode K-Means *Clustering* merupakan sistem untuk mengolah data obat pada UPT Puskesmas Kondoran menggunakan sebuah aplikasi berbasis desktop. Sehingga untuk dibangunnya sistem ini diperlukan data-data obat serta informasi dari UPT Puskesmas Kondoran Kec. Sangalla'. Dari dibangunnya sistem ini diharapkan dapat memberikan informasi mengenai obat mana yang pemakaiannya cepat dan mana yang tidak.

3.2 Blok Diagram Sistem

Berikut merupakan diagram blok sistem yang akan dibangun pada sistem pengelompokan data obat menggunakan metode K-Means



Gambar 1 Diagram Blok Sistem

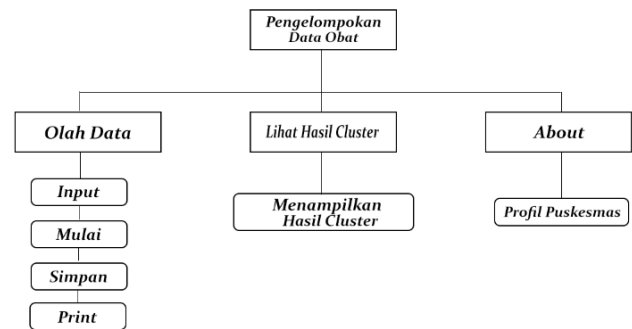
Penjelasan :

Pada tahap pertama user akan diminta untuk melakukan login terlebih dahulu, setelah login berhasil maka user harus menginputkan data obat pada sistem. Selanjutnya sistem akan mengolah dan memulai mengelompokkan data obat yang telah diinputkan kedalam 2 *cluster* yang telah ditentukan menggunakan metode K-Means *Clustering* yaitu

obat dengan pemakaian cepat dan yang tidak, setelah pengelompokan data selesai maka sistem akan menampilkan data yang telah dikelompokkan dan user data melihat data tersebut.

3.3 Struktur Menu

Gambar 2 merupakan rancangan struktur menu sistem yang akan dibangun. Sistem yang akan dibangun akan memiliki 3 menu utama yaitu Olah Data, Lihat Hasil Cluster, dan About.



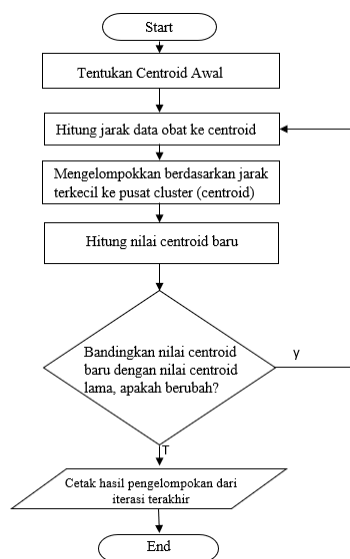
Gambar 2 Rancangan Struktur Menu

Penjelasan :

Menu Olah Data berfungsi untuk mengolah data obat didalamnya terdapat fitur yang pertama yaitu input untuk menginputkan data pada sistem sebelum diolah, kedua mulai merupakan sebuah tombol yang berfungsi untuk memulai proses pengelompokan data menggunakan metode K-Means *Clustering*, setelah itu didalam Olah Data juga terdapat fitur Simpan yang berfungsi untuk menyimpan data yang telah diolah serta Print untuk mencetak hasil clustering. Menu Lihat Hasil Cluster adalah menu yang berfungsi untuk menampilkan data hasil pengelompokan yang telah di proses sebelumnya. Menu About berfungsi untuk menampilkan informasi mengenai profil dari Puskesmas tersebut.

3.4 Flowchart Metode K-Means

Berikut merupakan flowchart untuk metode K-Means yang digunakan dalam proses pengelompokan data obat pada UPT Puskesmas Kondoran ditunjukkan pada gambar 3

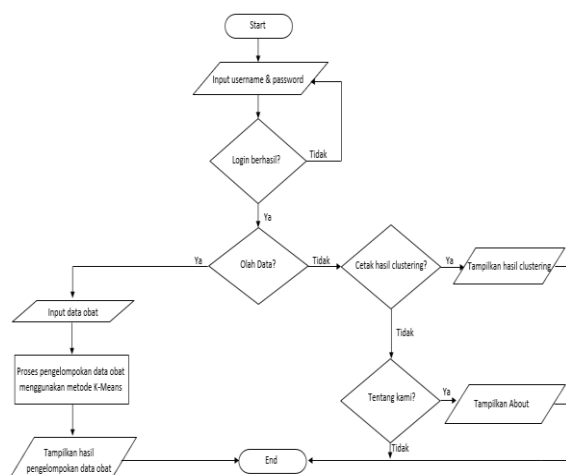


Gambar 3 Flowchart Metode K-Means

Penjelasan :

Pada gambar 3.3 diatas merupakan flowchart metode dari k-means, tahap pertama adalah menginputkan jumlah *cluster* yang diinginkan atau ditentukan, kemudian menentukan titik pusat atau *centroid* awal secara acak, kemudian melakukan perhitungan jarak masing-masing data ke pusat *cluster* menggunakan rumus *Euclidean Distance*, setelah selesai menghitung maka tentukan *cluster* dari masing-masing data dengan melihat nilai paling mendekati pada *centroid* atau titik pusat *cluster*, setelah masing-masing data mendapatkan kelompoknya maka selanjutnya menghitung nilai *centroid* baru, kemudian membandingkan apakah ada perubahan pada nilai *centroid* baru dengan nilai *centroid* lama jika terdapat perbedaan maka iterasi tetap berlanjut namun jika tidak terdapat perubahan maka proses algoritma selesai, hasil pengelompokan adalah hasil dari proses iterasi terakhir.

3.5 Flowchart Sistem



Gambar 4 Flowchart Sistem

Penjelasan :

Pada gambar 4 merupakan flowchart sistem yang akan dibuat, ketika sistem pertama kali dijalankan akan menampilkan halaman *loading page* pada aplikasi setelah itu user akan diminta untuk memasukkan *username* dan *password*, setelah memasukkan sistem akan mengecek apakah *username* & *password* tersebut cocok atau tidak, jika ya maka sistem akan melanjutkan ke tahap selanjutnya jika tidak maka sistem akan memunculkan form login lagi, setelah login berhasil sistem akan menampilkan halaman awal dari kemudian user akan dituntun untuk memilih apakah akan mengolah data? Jika ya, maka proses selanjutnya adalah menginputkan data pada sistem setelah itu sistem akan memproses data tersebut dan melakukan proses pengelompokan data obat menggunakan metode K-Means, setelah proses pengelompokan data selesai maka sistem akan menampilkan data hasil pengelompokan setelah itu proses akan selesai, jika tidak maka user akan dituntun untuk memilih apakah akan mencetak hasil clustering jika ya maka sistem akan menampilkan hasil dari clustering yang telah dilakukan sebelumnya jika tidak maka user akan dituntun untuk memilih tentang kami jika ya maka sistem akan menampilkan halaman *about* jika tidak maka proses sistem akan berakhir.

4. HASIL DAN PEMBAHASAN

4.1 Data Cleaning

Data Cleaning merupakan proses untuk dapat mengatasi nilai yang hilang, noise dan data yang tidak konsisten. Dataset yang didapatkan dari Puskesmas sebanyak 243 data dengan 14 kolom dilakukan data *cleaning* sehingga menghasilkan data *inskonsisten* sebanyak 34 data sehingga didapatkan data baru sebanyak 204 dengan 5 kolom yang akan digunakan dalam proses pengelompokan.

4.2 Pengujian Metode

Pengujian dilakukan pada sistem dengan menggunakan 204 data set yang didapatkan dari penelitian di lapangan.

4.2.1 Penentuan Cluster

Penentuan *cluster* pada penelitian ini dilakukan diawal yakni terdapat 2 cluster yang akan dijadikan tempat untuk mengelompokkan data obat masing-masing adalah :

1. *cluster* pertama obat dengan penggunaan lambat
2. *cluster* kedua obat dengan penggunaan lambat

4.2.2 Penentuan Centroid awal

Penentuan *centroid* awal pada penelitian ini dilakukan secara acak dengan memilih 2 data dari 204 data yang akan diuji

Masing-masing adalah :

Tabel 1 Tabel Pusat Awal Cluster

Cluster	a	b	c	d
1	3720	9000	8640	4080
2	7700	40000	19300	28400

4.2.3 Menghitung Masing-Masing Jarak Ke Pusat Cluster

Setelah didapat titik pusat awal cluster, kemudian dilakukan perhitungan jarak Euclidian, dan mengelompokkan berdasarkan jarak terkecil selanjutnya akan di dapat nilai centroid baru untuk acuan perhitungan berikutnya sampai nilai centroid sebelum dan sesudah bernilai sama.

Perhitungan jarak Euclidean pada iterasi 1

- 1) Berikut ini beberapa hasil perhitungan jarak Euclidean pada titik pusat 1 (3720,9000,8640,4080) :

$$D(1,1)=\sqrt{(61500-3720)^2+(10000-9000)^2+(43100-8640)^2+(28400-4080)^2}=71554$$

$$D(2,1)=\sqrt{(0-3720)^2+(20000-9000)^2+(18750-8640)^2+(1250-4080)^2}=16654$$

$$\dots$$

$$D(204,1)=\sqrt{(2350-3720)^2+(0-9000)^2+(872-8640)^2+(1478-4080)^2}=12247$$

- 2) Berikut ini beberapa hasil perhitungan jarak Euclidean pada titik pusat 2 (7700,40000,19300,28400) :

$$D(1,2)=\sqrt{(61500-7700)^2+(10000-40000)^2+(43100-19300)^2+(28400-28400)^2}=66037$$

$$D(2,2)=\sqrt{(0-7700)^2+(20000-40000)^2+(18750-19300)^2+(1250-28400)^2}=34594$$

$$\dots$$

$$D(204,2)=\sqrt{(2350-7700)^2+(0-40000)^2+(872-19300)^2+(1478-28400)^2}=51894$$

Setelah semua data telah diolah menggunakan rumus *Euclidean* maka kelompokkan data kedalam *cluster* berdasarkan nilai terkecil/jarak terkecil ke masing-masing *cluster*

Setelah itu hitung nilai *centroid* baru dengan menggunakan persamaan (2)

Tabel 2 Tabel Centroid baru

Cluster	a	b	c	d
1	9,424	3,029	9,148	3,304
2	134,405	138,307	180,918	91,795

Setelah nilai *centroid* baru didapatkan maka hitung nilai masing-masing data terhadap nilai *centroid* baru Perhitungan jarak Euclidean pada iterasi 2

- 1) Berikut ini beberapa hasil perhitungan jarak Euclidean pada titik pusat 1

$$D(1,1)=\sqrt{(61500-9424)^2+(10000-3029)^2+(43100-9148)^2+(28400-3304)^2}=67402$$

$$D(2,1)=\sqrt{(0-9424)^2+(20000-3029)^2+(18750-9148)^2+(1250-3304)^2}=21754$$

$$\dots$$

$$D(204,1)=\sqrt{(2350-9424)^2+(0-3029)^2+(872-9148)^2+(1478-3304)^2}=6344$$

- 2) Berikut ini beberapa hasil perhitungan jarak Euclidean pada titik pusat 2

$$D(1,2)=\sqrt{(61500-134405)^2+(10000-138307)^2+(43100-180918)^2+(28400-91795)^2}=211638$$

$$D(2,2)=\sqrt{(0-134405)^2+(20000-138307)^2+(18750-180918)^2+(1250-91795)^2}=257988$$

$$\dots$$

$$D(204,2)=\sqrt{(2350-134405)^2+(0-138307)^2+(872-180918)^2+(1478-91795)^2}=277743$$

Lakukan perhitungan *centroid* baru sampai dengan nilai *centroid* baru tidak berubah dari *centroid* lama/sebelumnya.

Berikut merupakan tampilan pada pengolahan data yang dilakukan pada sistem ditampilkan pada gambar 5

The screenshot displays a software interface for data clustering. It includes sections for 'DATA' (input data table), 'HASIL CLUSTERING DATA' (output clustering results), and 'PROSES TRAINING DATA' (training process details). The 'DATA' section shows a table with columns for 'No', 'Nama Obat', 'Sisa Obat', 'Persebaran', and 'Penggunaan'. The 'HASIL CLUSTERING DATA' section shows a table with columns for 'No', 'Nama Obat', 'Sisa Obat', 'Persebaran', 'Penggunaan', and 'Cluster'. The 'PROSES TRAINING DATA' section shows a table with columns for 'No', 'Nama Obat', 'Sisa Obat', 'Persebaran', 'Penggunaan', and 'Cluster'.

4.3 Table pengujian fungsional

Dari hasil pengujian diatas kita sudah mendapatkan hasil dari data uji masing-masing halaman monitoring dan tombol button nya, sebagaimana hasil tertera pada table 3.

Tabel 3. Tabel Hasil Uji Monitoring

No.	Item Uji	Hasil Uji	
		Berhasil	Gagal
1.	Halama Login	√	-
2.	Halaman Utama	√	-
3.	Halaman Olah Data	√	-
4.	Halaman Lihat Data	√	-
5.	Halaman About	√	-
6.	Button Proses	√	-
7.	Exit	√	-
8.	Button Home	√	-
9.	Button About	√	-
10.	Button Olah Data	√	-
11.	Button show entries data	√	-
12.	CRUD	√	-
13.	Button Print	√	-

5. KESIMPULAN DAN SARAN

5.1 Kesimpulan

Berdasarkan hasil perancangan sistem pengelompokan Data Obat Pada UPT Puskesmas Kondoran Kec. Sangalla' maka dapat diambil kesimpulan :

1. Berdasarkan hasil *cleaning* data didapatkan 39 data yang *premature* atau *inkonsisten* dari 243 data sehingga data yang digunakan dalam penelitian ini adalah 204 data
2. Berdasarkan hasil pengujian metode pada sistem dari 204 data yang diuji sebanyak 183 data masuk kedalam *cluster* 1 obat dengan pemakaian lambat dan 21 data masuk kedalam *cluster* 2 obat dengan pemakaian cepat

3. Berdasarkan hasil pengujian metode pada sistem dan penghitungan secara manual menggunakan excell menunjukkan kesamaan 100% pada hasil akhir yang didapatkan
4. Berdasarkan hasil pengujian fungsional menunjukkan 100% sistem sudah dapat berjalan dengan baik dan sesuai dengan fungsinya

5.2 Saran

Agar sistem dapat berjalan dengan baik dapat diterapkan beberapa saran :

1. Dataset yang digunakan bisa mengambil dari beberapa tahun agar hasil yang didapatkan lebih akurat
2. Agar lebih muda dalam *maintenance* sistem dibangun berbasis web
3. Agar mendapatkan hasil yang lebih akurat perbandingan 2 metode dapat dilakukan yaitu membandingkan hasil dari pengelompokan menggunakan metode *K-Means* dan *K-Nearest Neighbors*

DAFTAR PUSTAKA

- [1] Sulastri, Gufroni, 2017. Penerapan Data Mining dalam Pengelompokan Penderita Thalassaemia, *Jurnal Nasional Teknologi dan Sistem Informasi*, Vol 03, No 02
- [2] Astuti, 2017. Penerapan Data Mining untuk Clustering Data Penduduk Miskin Menggunakan Algoritma Hard C-Means, *Jurnal Ilmiah DASI*, Vol 18, No 01
- [3] Darmi, Setiawan, 2016. Penerapan Metode Clustering K-Means dalam pengelompokan penjualan produk, *Jurnal Media Infotama*, Vol 12, No 2.
- [4] Afiani R.K Fitra, 2018. Penerapan K-Means Clustering Untuk Mengetahui Varietas Padi Unggul Produksi Balai Pengkajian Teknologi Pertanian Jawa, *JATI (Jurnal Mahasiswa Teknik Informatika)*, Vol 02, No.01
- [5] Chusyairi A, Saputra R.N. Pelsri, 2019. Pengelompokan data Puskesmas Banyuwangi Dalam Pemberian Imunisasi Menggunakan Metode K-Means Clustering, *Telematika*, Vol 12, No 02
- [6] Sanah Nor, 2017. Pelaksanaan Fungsi Puskesmas (Pusat Kesehatan Masyarakat) Dalam Meningkatkan Kualitas Pelayanan Kesehatan di Kecamatan Long Kali Kabupaten Paser, *eJurnal Ilmu Pemerintahan*, Vol 05, No 01.
- [7] Purnamaningsih Chandra, S., & Saptono Ristu, 2014. Pemanfaatan Metode K-Means Clustering dalam Penentuan Penjurusan Siswa SMA, *Jurnal ITSMART*, Vol 03, No. 01.