

PERBANDINGAN METODE FUZZY C-CMEANS DAN K-MEANS CLUSTERING PADA DATA PENGGUNAAN OBAT DI R.S NATIONAL HOSPITAL SURABAYA

Teddy Gideon Manik, Woro Isti Rahayu, Rd. Nuraini Siti Fathonah

Program Studi Sarjana Terapan Teknik Informatika

Universitas Logistik Bisnis Internasional

Jl. Sari Asih No.54, Kota Bandung, Jawa Barat, Indonesia

teddygideonmanik@gmail.com, woroisti@ulbi.ac.id, nurainisf@ulbi.ac.id

ABSTRAK

Rumah sakit merupakan institusi dibidang pelayanan kesehatan masyarakat, berfungsi untuk melayani masyarakat secara luas dalam bentuk jasa. Namun di Rumah sakit National Hospital Surabaya saat ini mencatat total penggunaan obat tanpa menampilkannya dalam bentuk visual sebaran data. Penelitian ini melakukan perbandingan klustering untuk penggunaan obat dengan metode K-Means dan Fuzzy C-Means menggunakan data tahun 2021. Untuk Fuzzy C-Means menggunakan fungsi objektif dari data dan pada K-Means hanya dengan menentukan titik centroid yang akan digunakan dalam *Clustering*. Dari perbandingan yang dibuat untuk menentukan banyak, sedang dan sedikitnya penggunaan obat. K-Means hanya memerlukan 2 iterasi saja untuk mendapatkan rasio yang tepat untuk pengelompokan. Sedangkan untuk Fuzzy C-means perlu melakukan iterasi sebanyak 6 kali untuk mendapatkan error terkecilnya untuk pengelompokan.

Kata kunci: Rumah Sakit, K-Means, Fuzzy C-Means, Clustering

1. PENDAHULUAN

Sebagai institusi di bidang pelayanan kesehatan masyarakat, rumah sakit National Hospital Surabaya beroperasi untuk melayani masyarakat secara luas dalam bentuk jasa. Oleh karenanya rumah sakit menyelenggarakan pelayanan kesehatan secara paripurna yang memiliki pelayanan seperti rawat inap, rawat jalan dan gawat darurat. Kesimpulannya, rumah sakit melayani pengobatan.

Melihat secara khusus bidang pengobatan, obat menjadi komponen tak tergantikan dalam pelayanan kesehatan di rumah sakit mengingat obat adalah segala bahan atau ramuan yang digunakan untuk mencegah, meringankan, dan menyembuhkan penyakit pada makhluk hidup. Jadi, obat menjadi hal yang paling dibutuhkan oleh makhluk hidup saat sakit. Adapun kebutuhan obat tidak lepas dari bagaimana obat itu diolah dan bagaimana kebutuhan obat direncanakan. [1].

Pengelolaan persediaan obat yang baik dapat menjadi salah satu pendukung dalam peningkatan mutu pelayanan persediaan obat. Namun kenyataan masih terdapat kendala dalam pengelompokan jenis penyakit untuk jenis obat yang digunakan. Metode clustering yang dapat digunakan adalah metode *Fuzzy C-Means* dan *K-Means* adalah metode pembelajaran mesin tanpa pengawasan yang paling banyak digunakan dan relatif berhasil di antara banyak algoritma pengelompokan metode *Fuzzy C-Means* dan *K-Means*. Konsep dasar fcm adalah melakukan perulangan untuk memperbaiki pusat cluster dan derajat keanggotaan yang pada awalnya masih belum akurat, pengulangan ini terus dilakukan sampai didapatkan pusat cluster yang tepat sedangkan untuk metode *K-Means* merupakan algoritma yang dibutuhkan parameter input sebanyak k dan membagi

sekumpulan n objek kedalam k cluster. Clustering adalah proses pengelompokan objek data ke dalam kelompok berdasarkan kemiripannya dalam satu algoritma yang sering digunakan dalam clustering menggunakan algoritma *K-Means* dan *Fuzzy C-Means*. Perbandingan dalam metode *Fuzzy C-Means* dan *K-Means*, pada metode *Fuzzy C-Means* adalah pengelompokan data ditentukan oleh derajat keanggotaan, sedangkan metode *K-Means* adalah pengelompokan data ditentukan dari centroid penggunaan obat.[2]

Diperlukan sebuah metode yang tepat dan efektif dalam penentuan jenis obat apa saja yang sudah diambil atau digunakan untuk resep obat oleh seorang dokter di Rumah Sakit, yang dapat mempermudah untuk mengetahui nya lebih mudah dengan menghitung menggunakan metode *K-Means* dan metode *Fuzzy K-Means*. Kepentingannya dan dalam mengatasi permasalahan tersebut menggunakan teknik Data mining yaitu metode *Fuzzy C-Means* dan *K-Means* merupakan salah satu metode clustering non hirarki yang berusaha mempartisi data yang ada ke dalam bentuk satu *cluster* atau lebih. Pada proses pengelompokan (*clustering*), dalam metode *K-Means* suatu objek hanya akan menjadi anggota satu *cluster* dan sulit untuk mencapai kekonvergenan. Oleh karena itu, digunakan metode yang lain yaitu metode *Fuzzy C-Means (FCM)*. *FCM* merupakan metode untuk melakukan pengelompokan data yang diberikan satu atau lebih data cluster. Algoritma *FCM* digunakan karena pada metode *FCM* kemungkinan kegagalan untuk kesamaan lebih kecil dibandingkan metode *K-Means* [3].

Berdasarkan keterangan di atas penulis membandingkan secara teoritis dan aplikasi antara metode *K-Means* dan metode *FCM* untuk

mengelompokkan pada data obat yang digunakan dokter di Rumah Sakit National Surabaya, dengan judul “Perbandingan Metode *K-Means* dan Metode *Fuzzy C-Means* (FCM) untuk Clustering Data obat digunakan dokter di Rumah Sakit National Surabaya. Adapun tujuan pada penelitian ini adalah untuk mengetahui hasil dari jumlah obat yang digunakan setiap dokter spesialis di Rumah Sakit National Hospital Surabaya. dan untuk mengetahui hasil perbandingan dalam penerapan dari metode *K-Means* dan *Fuzzy C-Means* pada sistem [4].

2. TINJAUAN PUSTAKA

2.1. DATA MINING

Data mining merupakan disiplin ilmu yang bertujuan untuk menemukan, mengekstraksi, menggali atau menambang pengetahuan berdasarkan dari suatu data atau informasi. Data mining dapat disebut pula sebagai Knowledge Discovery in Database (KDD) yang merupakan proses dalam mengumpulkan dan menggunakan data dengan tujuan mencari hubungan antar setiap data atau pola dalam suatu Big Data (data yang besar). Clustering merupakan salah bagian dari data mining yang digunakan dalam pengelompokan data[5].

2.2. *K-Means* Clustering

K-Means adalah pengelompokkan data dengan memaksimalkan kemiripan data dalam satu kluster dan meminimalkan kemiripan data antara kluster [6]. *K-Means* memiliki fungsi untuk mengelompokkan data kedalam data *cluster*. Algoritma ini dapat menampung data tanpa menggunakan label kategori. *K-Means Clustering* Algoritma juga merupakan metode non-hierarchy. Metode dengan menghitung ukuran jarak ini terdiri dari metode cluster berhierarchy dengan penggabungan (*agglomerative*) dan pemisahan (*decisive*). Contohnya adalah metode pautan lengkap complete linkage, metode pautan rata-rata average linkage, metode Ward serta metode cluster tak berhierarchy yaitu metode *K-Means* [7]. *Cluster Sampling* adalah cara pengambilan contoh di mana bagian-bagian populasi dipilih secara random dari kelompok yang sudah tersedia yang disebut dengan ‘cluster’, *Clustering* atau klusterisasi adalah salah satu masalah yang menggunakan teknik *unsupervised learning*.

K Means Clustering memiliki *objective* yaitu meminimalis dari *object* yang telah di susun pada proses klusterisasi. Dengan cara minimalisasi varian antara satu kluster dengan maksimalis variasi pada data di kluster lainnya. *K-Means clustering* merupakan bentuk metode algoritma yang mendasar[8] penerapannya adalah:

- Menentukan jumlah kluster
- Secara acak mendistribusikan data kluster
- Menghitung rata rata dari data yang ada di kluster.
- Menggunakan langkah baris 3 kembali sesuai nilai threshold
- Menghitung jarak antara data dan nilai

centroid(*K-Means clustering*)

- Distance space* dapat diimplementasikan untuk menghitung jarak data dan *centroid*. Tujuan dari *K-Means clustering*, seperti metode kluster lainnya, adalah untuk mendapatkan kelompok data dengan memaksimalkan kesamaan karakter dalam kluster dan memaksimalkan perbedaan antara setiap kluster.

Algoritma *K-Means clustering* mengelompokkan data dari jarak antara data terhadap titik sumbu centroid kluster yang didapatkan melalui perhitungan proses secara berulang. Analisis yang diperlukan menentukan jumlah K sebagai inputan untuk algoritma. Dalam machine learning, algoritma *K-Means clustering* termasuk dalam bagian dari jenis *unsupervised learning*.

Tahap pertama yang perlu dilakukanyaitu dengan menghitung jarak pada setiap data terhadap setiap *centroid* dengan menggunakan persamaan *Euclidean Distance*:

$$d(x_i, \mu_i) = \sqrt{(x_i - \mu_i)^2}$$

Pada tahap selanjutnya akan dilakukan pengelompokkan setiap data terhadap jarak pada titik pusat centroid terdekat. Dari data tersebut, Ubah nilai centroid yang diperoleh dari rata-rata cluster yang bersangkutan dengan persamaan berikut:

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

Dimana n_k merupakan jumlah data dalam cluster, sedangkan d_i merupakan jumlah dari nilai jarak yang masuk dalam masing-masing cluster. Jika anggota tiap cluster berubah, maka iterasi akan dilanjutkan dengan menggunakan langkah kedua hingga kelima. Jika anggota tiap cluster tidak ada yang berubah, maka iterasi berhenti dan nilai rata-rata pusat cluster (μ_i) akan digunakan sebagai parameter dalam penentuan pembagian data. Proses perhitungan *K-Means Clustering* telah selesai. [9].

2.3. *Euclidean Distance*

Perhitungan jarak dari 2 buah titik adalah, ruang *Euclidean*. Ruang *Euclidean* diperkenalkan oleh Euclid, seorang matematikawan yang berasal dari Yunani tahun 300 B.C.E. Untuk mempelajari hubungan antara sudut dan jarak. *Euclidean* ini berhubungan dengan Teorema Pythagoras dan biasanya diterapkan pada 1 sampai 3 dimensi. Tapi juga dapat lebih disederhanakan jika diterapkan pada dimensi yang lebih tinggi. [10]

2.4. *Fuzzy C-Means Clustering*

Fuzzy C-Means (FCM) adalah teknik mengelompokkan data untuk setiap titik memiliki tingkat anggota cluster berdasarkan logika *Fuzzy*. *Fuzzy C-Means clustering* (FCM) menggunakan

matrix partisi U Fuzzy sehingga suatu titik data dapat dimiliki oleh lebih dari satu kelompok dengan memiliki nilai keanggotaan (bobot) yang berbeda yang dimulai dari 0 hingga 1. Semakin dekat titik data ke pusat cluster, maka keanggotaan titik datanya akan menuju pusat cluster yang tepat. Jumlah nilai keanggotaan dari setiap objek harus sama dengan satu [11].

Berikut rumus untuk mengkluster:

$$Q_i = \sum_{k=1}^c \mu_{ik}$$

$$\mu_{ik} = \frac{\mu_{ik}}{Q_i}$$

Kemudian menghitung pusat kluster V_{kj} dengan menggunakan:

$$V_{kj} = \frac{\sum_{i=1}^n ((\mu_{ik})) w_* x_{ij}}{\sum_{i=1}^n (\mu_{ik})^w}$$

Menghitung fungsi objektif yang didapat dari t iterasi menggunakan:

$$p_t = \sum_{i=1}^n \sum_{k=1}^c \left(\left[\sum_{j=1}^m (X_{ij} - V_{kj})^2 \right] \mu_{ik}^w \right)$$

Tahap lanjutan menghitung hasil perubahan dari partisi matriks :

$$\mu_{ik} = \frac{\left[\sum_{j=1}^m (X_{ij} - V_{kj})^2 \right]^{-\frac{1}{w-1}}}{\sum_{k=1}^c \left[\sum_{j=1}^m (X_{ij} - V_{kj})^2 \right]^{-\frac{1}{w-1}}}$$

2.5. Machine Learning

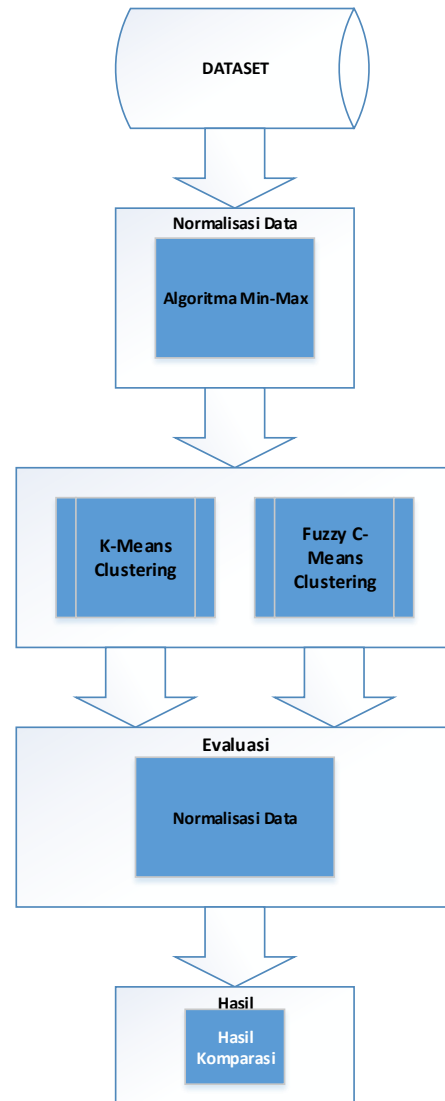
Machine learning mengarah pada sebuah metode yang membuat komputer memiliki kemampuan dalam mengerjakan sebuah pekerjaan secara otomatis dan terstruktur. Proses *machine learning* dilakukan melalui algoritma tertentu, sehingga pekerjaan yang dilakukan komputer dapat dilakukan secara otomatis [12].

Machine Learning merupakan salah satu ilmu Kecerdasan Buatan (*Artificial Intelligence*) yang membahas mengenai pembangunan sistem yang berdasarkan pada jenis data yang digunakan. Dalam pemetaan *Machine Learning*, terdapat beberapa scenario, seperti:

1. Supervised Learning merupakan pembelajaran menggunakan masukan data pembelajaran yang telah diberi label. Melakukan prediksi pada data yang memiliki label.
2. menggunakan masukan data yang tidak diberikan label. Setelah itu mencoba untuk mengelompokkan data berdasarkan karakteristik jenis data yang ditemui.
3. Reinforcement Learning fase pembelajaran dan tes saling dicampur. Untuk mengumpulkan informasi yang digunakan dengan berinteraksi ke lingkungan sehingga mendapatkan hasil yang akan digunakan untuk pembelajaran [13]

3. METODE PENELITIAN

Metode penelitian adalah ilmu yang mempelajari cara-cara untuk menyusun rencana penelitian, pelaksanaan, dan penulisan laporan dengan metode ilmiah secara efisien dan sistematis yang hasilnya berguna untuk memecahkan masalah dan pengembangan ilmu pengetahuan guna menyusun keputusan.



Gambar 1. Diagram alur penelitian

3.1. Dataset

Data penggunaan obat yang telah dilaksanakan oleh dokter spesialis bagi pasien di Rumah Sakit Nation Hospital Surabaya yang diambil dari bulan maret-agustus 2021 sebagai data percobaan dalam penelitian ini. Data ditampilkan di tabel I. total data dokter yang bekerja di Rumah Sakit National Hospital Surabaya sebanyak 51. Berdasarkan data ini, data normalisasi digunakan untuk data obat yang digunakan pada tabel I, sehingga lebih mudah untuk menghitung menggunakan *K-Means* dan *Fuzzy C-Mean*. [14]

3.2. Algoritma Min-Max

Algoritma Min-Max merupakan algoritma yang dipergunakan dalam penormalisasian data dengan cara menyamakan dan memperkecil rentang data. Algoritma Min-Max merupakan teknik sederhana di mana teknik tersebut dapat secara khusus sesuai dengan data dalam batas yang ditentukan sebelumnya dengan batas yang telah ditentukan sebelumnya [15].

3.3. K-Means Clustering

Algoritma K-Means merupakan metode pengelompokan data yang melakukan pembagian data ke dalam kelompok tertentu yang memisahkan data ke dalam kelompok yang berbeda. Secara iteratif, K-Means mampu meminimalkan rata-rata jarak setiap data ke klusternya. Pada dasarnya penggunaan algoritma K-Means dalam melakukan proses clustering tergantung dari data yang ada dan konklusi yang ingin dicapai. Untuk itu digunakan algoritma [16].

3.4. Fuzzy C-Means (FCM)

Fuzzy C-Means (FCM) adalah suatu cara pengelompokan data untuk keberadaan setiap titik data dalam suatu kluster ditentukan oleh derajat keanggotaan dari fungsi objektif. FCM bukan merupakan keanggotaan *Fuzzy inference system*, namun tepatnya urutan baris pusat cluster dan beberapa derajat keanggotaan untuk setiap titik data. Informasi ini dapat digunakan untuk membangun suatu *Fuzzy inference system* [17].

3.5. Normalisasi Data

Pada tahap analisis data yaitu untuk mengetahui apakah data tersebut memiliki missing value dan sebagainya. Data yang digunakan berupa jumlah data jenis obat di Rumah Sakit. Analisa data yang dilakukan dalam penelitian ini ialah sebagai berikut: Dengan melakukan pengumpulan data dan konversi nilai bentuk object dalam data excel [18].

Tabel 1. Dataset Obat

No.	OBAT	CARDIOLOGIST	GENERAL PRACTITIONER	INTER NIST	NEURO SURGEON	NEURO LOGIST	SURGEON	THT
1	ABBOTIC 125 MG/5 ML 30 ML SYR	1	21	30	35	1	7	14
2	ABBOTIC 125 MG/5 ML 60 ML SYR	0	3	30	0	0	0	0
3	ABBOTIC 250 MG/5 ML 50 ML SYR	0	2	10	0	0	0	0
4	ABBOTIC XL 500 MG TAB	0	1	3	0	0	0	0
5	ABIXA 10 MG TAB	0	30	5	0	0	0	0
6	ACPULSIF 5 MG TAB	1	2	2	1	21	2	180
7	ACTIFED PLUS COUGH SYR	5	1	1	15	5	2	75
8	ACTILYSE 50 MG INJ	0	5	20	0	0	0	0
9	ACTOS 30 MG TAB	0	155	120	0	0	0	0
10	ADALAT OROS 30 MG TAB	0	70	30	0	0	0	0
11	ADONA AC17 10 MG TAB	0	4	45	0	0	0	0
...								
1124	ZITHROMAX 500 MG TAB	0	10	1	0	0	0	0
1125	ZITHROMAX 500 MG VIAL	0	11	30	0	0	0	0
1126	ZITHROMAX 600 MG 15 ML DRY SYR	0	166	5	0	0	0	0
1127	ZOLMIA 10 MG TAB	0	1	15	0	0	0	0
1128	ZOVIRAX 200 MG TAB	0	15	120	0	0	0	0
1129	ZOVIRAX 5% CR 2 GR	8	1	1	3	50	5	3
1130	ZOVIRAX 5% CR 5 GR	0	3	30	0	0	0	0
1131	ZYPRAZ 0,25 MG TAB	0	25	20	0	0	0	0

3.6. Hasil Komparasi

Menurut Winarno Surakhmad dalam bukunya Pengantar Pengetahuan Ilmiah (1986: 84), komparasi adalah penyelidikan deskriptif yang berusaha mencari pemecahan melalui analisis tentang hubungan sebab akibat, yakni memilih faktor-faktor tertentu yang berhubungan dengan situasi atau fenomena yang diselidiki dan membandingkan satu faktor dengan faktor lain [19].

4. HASIL DAN PEMBAHASAN

Jumlah data yang digunakan dalam penelitian ini adalah sebanyak 1,131 data, dengan membuat attribute sebanyak 3, banyak, sedang, sedikit.

4.1. Analisis dengan K-Means Clustering

Untuk menguji metode K-Means, penulis melakukan pengujian metode K-Means dengan menggunakan google colab. Langkah-langkah melakukan perhitungan dengan menggunakan metode K-Means Clustering sebagai berikut:

```
kmean = KMeans(n_clusters=3)
kmean
```

Dengan menentukan 3 pusat *cluster* yang kemudian untuk dicari pusat dari *cluster* menggunakan:

```
y_cluster = kmean.fit_predict(x_train)
y_cluster
```

Selanjutnya dengan menggunakan:

```
kmean.cluster_centers_
```

Untuk mendapatkan pusat centroid yang digunakan dan mendapatkan hasil seperti:

```
array([[ 1.      , 24.      , 6.      , 6.5      ,
        11.      , 1.      , 276.     ],
       [13.2     , 43.2     , 128.6    , 0.8      ,
        10.2     , 0.4      , 38.4     ],
       [ 3.69230769, 11.65384615, 18.5     , 2.15384615,
        0.65384615, 0.34615385, 3.57692308]])
```

Gambar 2. Pusat *centroid* K-Means Clustering

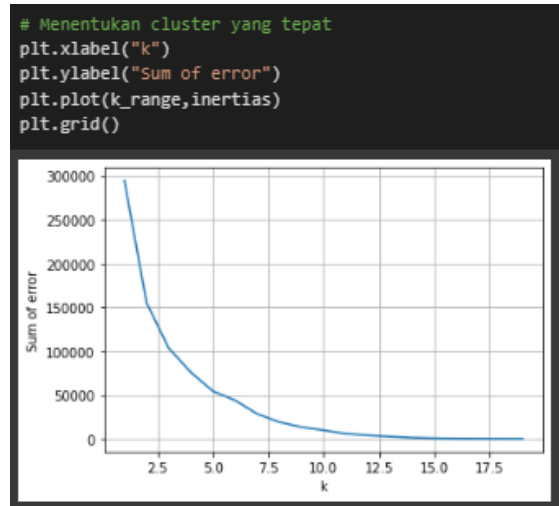
Kemudian menentukan pusat *cluster* menggunakan bantuan metode *Elbow*:

```
# Mencari jumlah cluster
inertias = []
k_range = range(1,20)
for k in k_range:
    km = KMeans(n_clusters=k).fit(x_train)
    inertias.append(km.inertia_)

inertias

[294601.5151515152,
155535.19999999995,
103753.10344827587,
76114.75,
54504.55128205128,
44074.54545454546,
28614.519999999997,
19278.661904761902,
13566.449999999999,
9932.95,
6094.788888888889,
4158.288888888889,
2761.3888888888887,
1523.4384615384615,
837.0818181818181,
591.7083333333334,
419.0333333333333]
```

Gambar 3. Menentukan Kluster dengan *Elbow*



Gambar 4. Hasil plot scatter elbow

Untuk menentukan *cluster* menggunakan:

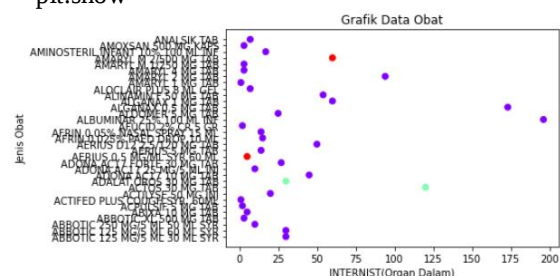
```
y_cluster = kmean.fit_predict(x_train)
y_cluster
```

```
array([2, 0, 0, 0, 0, 1, 0, 0, 0, 0, 0, 1, 0, 0, 0, 0, 0, 0, 0, 0, 1,
       0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0], dtype=int32)
```

Gambar 5. Hasil *cluster* K-Means

Dengan hasil yang didapatkan untuk sebaran data menggunakan K-Means seperti dibawah ini :

```
plt.scatter(dataset.iloc[:, 3], dataset.iloc[:, 0], c =
kmeans.labels_, cmap='rainbow')
plt.xlabel("INTERNIST(Organ Dalam)")
plt.ylabel("Jenis Obat")
plt.title("Grafik Data Obat")
plt.show
```

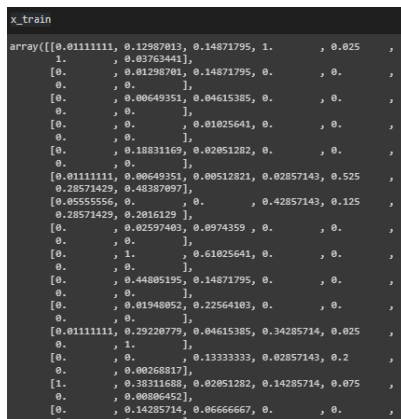


Gambar 6. Visual sebaran data K-Means Clustering

4.2. Analisis dengan Fuzzy C-Means Clustering

Dengan menentukan dan mengubah data ke bentuk *array* untuk mengubah data berada di rentang 0 sampai 1 dengan perintah:

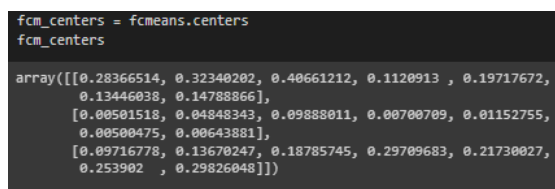
```
scaler = MinMaxScaler()
x_train=scaler.fit_transform(x_train)
```



Gambar 7. Data dalam bentuk *Array*

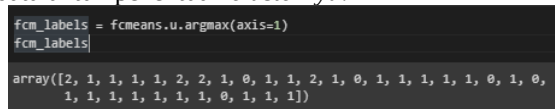
Kemudian menentukan titik *cluster* yang diambil dari data diatas dengan perintah dibawah berikut:

```
fcmeans = FCM(n_clusters=3)
fcmeans.fit(x_train)
```



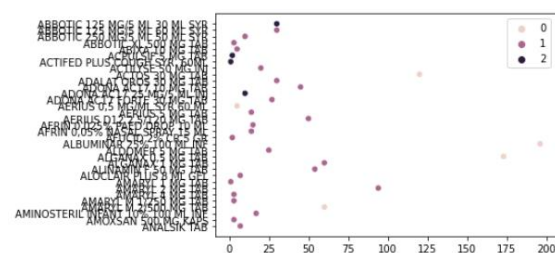
Gambar 8. Pusat *cluster*

Kemudian untuk membuat pengelompokan dari data untuk penentuan *clusternya*:



Gambar 9. Hasil *Cluster Fuzzy C-Means*

```
scatter(x,y, hue=fcm_labels)
```



Gambar 10. Visual sebaran data *Fuzzy C-Means*

5. KESIMPULAN DAN SARAN

Berdasarkan hasil analisa terdapat perbedaan dalam menentukan banyak, sedang, dan sedikit penggunaan obat. Dari kedua metode tersebut didapatkan hasil bahwa *K-Means* diperlukan 2 iterasi untuk mendapatkan rasio yang tepat untuk mengelompokkan. Untuk metode *Fuzzy C-means* didapati hasil bahwa iterasi yang diperlukan sebanyak

6 kali untuk mendapatkan error terkecilnya untuk pengelompokan.

Adapun saran yang dapat diberikan untuk penelitian selanjutnya dapat dikembangkan untuk membandingkan metode-metode lain seperti *purity* dan *average linkage* dan membuat perancangan system untuk digunakan.

DAFTAR PUSTAKA

- [1] H. S. Firdaus, A. L. Nugraha, B. Sasmitho, and M. Awaluddin, "PERBANDINGAN METODE FUZZY C-MEANS DAN K-MEANS UNTUK PEMETAAN DAERAH RAWAN KRIMINALITAS DI KOTA SEMARANG," *Elipsoida : Jurnal Geodesi dan Geomatika*, vol. 4, no. 01, pp. 58-64, Nov. 2021. <https://doi.org/10.14710/elipsoida.2021.9219>
- [2] Rahman Syarif, Muhammad Tanzil Furqon dan Sigit Adinugroho, Perbandingan Algoritme K-Means Dengan Algoritme Fuzzy C Means (FCM) Dalam Clustering Moda Transportasi Berbasis GPS, Vol. 2, No. 10, Oktober 2018, hlm. 4107-4115
- [3] Agusta, Yudhi. 2007. 'K-Means penerapan permasalahan dan metode terkait'. *Jurnal Sistem dan Informatika*, Vol 3.
- [4] Yohannes, Analisis Perbandingan Algoritma Fuzzy C-Means dan K-Means, Vol 2 No. 1, <https://ars.ilkom.unsri.ac.od/>
- [5] Aditya Ramadhan, Zuliar Efendi, Mustakim Mustakim, PERBANDINGAN K-MEANS DAN FUZZY C-MEANS UNTUK PENGELOMPOKAN DATA USER KNOWLEDGE MODELING, ISSN (Online) : 2579-5406, <https://ejournal.uin-suska.ac.id/index.php/SNTIKI/article/view/3268>
- [6] Nyoman Mahayasa Adiputra, CLUSTERING PENYAKIT DBD PADA RUMAH SAKIT DHARMA KERTI MENGGUNAKAN ALGORITMA K-MEANS, *Information System and Emerging Technology Journal*. Vol. 2, No. 2, Desember 2021, <https://ejournal.undiksha.ac.id/index.php/insert/article/view/41673/20969>
- [7] Andesberg, 1973, Clustering Algoritma (K-Means), <https://sis.binus.ac.id/2022/01/31/clustering-algoritma-k-means/>
- [8] Ariyanti, Dini Puspita Sukma, Perbandingan fuzzy c-means dan k-means untuk text clustering menggunakan LSI, 2021, <http://hdl.handle.net/123456789/12047>
- [9] Prihandoko, Perbandingan Algoritma K-Means dengan Fuzzy C-Means Untuk Clustering Tingkat Kedisiplinan Kinerja Karyawan, Vol 2 No 3 (2018): Desember 2018, <https://doi.org/10.29207/resti.v2i3.492>
- [10] Fitria Febrianti, Ahmad Hanif Asyhar, Moh. Hafiyusholeh, PERBANDINGAN

- PENGLUSTERAN DATA IRIS MENGGUNAKAN METODE K-MEANS DAN FUZZY C-MEANS, 10.15642/mantik.2016.2.1.7-13
- [11] Jimmy Pujoseno, IMPLEMENTASI DEEP LEARNING MENGGUNAKAN CONVOLUTIONAL NEURAL NETWORK UNTUK KLASIFIKASI ALAT TULIS, <https://www.scribd.com/document/503450031/DeepLearning-CNN-Kasus2#>
- [12] Al Faruqi, Muhammad (2021) Sistem Pemetaan Posisi Objek Kendaraan Menggunakan Pengolahan Citra Pada Area 360°. Other thesis, Universitas Komputer Indonesia.
- [13] Safwandi Safwandi, Novia Hasdyna, Nur Azizah, Analisis K-Means Clustering pada Data Sepeda Motor, VOL 5 NO 1 (2020), DOI: <https://doi.org/10.19184/isj.v5i1.17071>
- [14] Muhammad Kurniawan, Afib Pamungkas, Salman Hadi, ALGORITMA MINIMAX SEBAGAI PENGAMBIL KEPUTUSAN DALAM GAME TIC-TAC-TOE, Vol 4, No 1 (2016), <https://ojs.amikom.ac.id/index.php/semnastekno/media/article/view/1243>
- [15] BUSTAMI YUSUF, ANALISIS DAN PERBANDINGAN KUALITAS PENGELOMPOKAN DOKUMEN (DOCUMENT CLUSTERING) DENGAN MENGGUNAKAN METODE K-MEANS DAN K-MEDIANS, <https://jurnal.ar-raniry.ac.id/index.php/elkawnie/article/download/536/2699>
- [16] Herlina Latipa Sari, Perbandingan Algoritma Fuzzy C-Means (FCM) Dan Algoritma Mixture Dalam Pencusteran Data Curah Hujan Kota Bengkulu, <https://123dok.com/document/yjr11p5z-perbandingan-algoritma-fuzzy-c-means-fcm-dan-algoritma-mixture-dalam-pencusteran-data-curah-hujan-kota-bengkulu.html/>
- [17] Shavira Siti Rachmasari, Abdul Kudus, PERBANDINGAN PENERAPAN ALGORITME K-MEANS DAN FUZZY C-MEANS UNTUK MENGELOMPOKKAN DATA KINERJA DOSEN UNIVERSITAS ISLAM BANDUNG, Vol 7, No 2, Prosiding Statistika (Agustus, 2021), <https://karyailmiah.unisba.ac.id/index.php/statistika/article/view/28917>
- [18] Parmaji, Komparasi Pendapatan Usahatani Cabai Merah Dan Padi Sawah Di Lahan Irigasi Pada MT I Di Desa Triyoso Belitang OKU Timur, <http://meaningaccordingtoexperts.blogspot.com/2017/04/pengertian-dan-macam-macam-penelitian.html?m=1>
- [19] Suparmi Suparmi, Soeheri Soeheri, SISTEM INFORMASI GEOGRAFIS PEMETAAN TEMPAT KOST BERBASIS WEB MENGGUNAKAN METODE EUCLIDEAN DISTANCE, Vol 5, No 1, DOI: <http://dx.doi.org/10.22303/infosys.5.1.2020.105-113>