



Pengelompokan Data Kasus Keracunan Makanan Biologis Berdasarkan Faktor Penyebab Menggunakan Metode Clustering

Dwi Astuti¹, Relita Buaton², Magdalena Simanjuntak³

*dwi406606@gmail.com, *bbcbuaton@gmail.com, *magdalena.simanjuntak84@gmail.com

Program studi Sistem Informasi, STMIK Kaputama, Indonesia

Alamat: Jl. Veteran No.4A, Tangsi, Kec. Binjai Kota, Kota Binjai, Sumatera Utara 20714

Abstract: Cases of biological food poisoning can be caused by several causative factors, one of which is a food processing site that does not meet health requirements. According to the BPOM report (2016) cases of food poisoning in Indonesia in 2016 reached 1,068 cases. In 2016, 60 extraordinary events (KLB) of food poisoning were reported by 31 BB/BPOM throughout Indonesia. From the many cases of food poisoning that occur, it is necessary to take action in prevention by processing data on existing cases of poisoning to follow up on existing problems to reduce the number of cases of food poisoning by using a system on a computer so that the managed data can be processed quickly to obtain further information. Therefore the author wants to use a system with the clustering method to assist in processing data on biological poisoning cases grouping objects based on the characteristics of each object. Based on the research conducted, it can be seen that in cluster 2 in the data group of biological poisoning cases there are 11 data with centroid point age (x) 2, namely 12-16 years, centroid point on the type of poisoning (y) 6.36, namely sandwiches, and centroid point on the causative factor (z) 2.9, namely Gram-negative rod-shaped bacteria which are usually found in the intestines of humans and warm-blooded animals.

Keywords: Data Mining, K-Means Algorithm, Food Poisoning

Abstrak: Kasus keracunan makanan biologis dapat disebabkan oleh beberapa faktor penyebab salah satunya yaitu tempat pengolahan makanan yang tidak memenuhi syarat kesehatan. Menurut laporan BPOM (2016) kasus keracunan akibat makanan di Indonesia pada tahun 2016 mencapai 1.068 kasus. Pada tahun 2016 sebanyak 60 kejadian luar biasa (KLB) keracunan pangan dilaporkan oleh 31 BB/BPOM di seluruh Indonesia. Dari banyaknya kasus keracunan makanan yang terjadi maka perlu adanya tindakan dalam melakukan pencegahan dengan mengolah data kasus keracunan yang ada untuk menindaklanjuti masalah yang ada untuk mengurangi angka kasus keracunan makanan dengan menggunakan sistem pada komputer agar data yang dikelola dapat diproses dengan cepat untuk mendapatkan informasi lebih lanjut. Maka dari itu penulis ingin menggunakan suatu sistem dengan metode clustering untuk membantu dalam mengolah data kasus keracunan biologis mengelompokkan obyek-obyek berdasarkan karakteristik yang dimiliki masing-masing obyek. Berdasarkan penelitian yang dilakukan, dapat diketahui bahwa pada cluster 2 pada kelompok data kasus keracunan biologis terdapat 11 data dengan titik centroid usia (x) 2 yaitu 12-16 tahun, titik centroid pada jenis keracunan (y) 6.36 yaitu sandwich, dan titik centroid pada faktor penyebab (z) 2.9 yaitu bakteri Gram negatif berbentuk batang yang biasanya terdapat pada usus manusia dan hewan berdarah panas.

Kata kunci: Data Mining, Algoritma K-Means, Keracunan Makanan

1. PENDAHULUAN

Makanan merupakan kebutuhan dasar manusia untuk keberlangsungan hidup dan sebagai sumber energi untuk menjalankan aktivitas fisik maupun biologis didalam kehidupan sehari-hari. Makanan yang dibutuhkan harus sehat yakni memiliki nilai gizi yang optimal dan lengkap seperti vitamin, mineral, karbohidrat, protein lemak dan lainnya. Kasus keracunan makanan biologis dapat disebabkan oleh beberapa faktor penyebab salah satunya yaitu tempat pengolahan makanan yang tidak memenuhi syarat kesehatan. Menurut laporan BPOM (2016) kasus keracunan akibat makanan di Indonesia pada tahun 2016 mencapai 1.068 kasus.

Pada tahun 2016 sebanyak 60 kejadian luar biasa (KLB) keracunan pangan dilaporkan oleh 31 BB/BPOM di seluruh Indonesia. Dari kejadian luar biasa tersebut sebanyak 5.673 orang terpapar, 3.351 orang mengalami sakit dan 7 orang meninggal dunia. (Santoso et al., 2019)

Dari banyaknya kasus keracunan makanan yang terjadi maka perlu adanya tindakan dalam melakukan pencegahan dengan mengolah data kasus keracunan yang ada untuk menindaklanjuti masalah yang ada untuk mengurangi angka kasus keracunan makanan dengan menggunakan sistem pada komputer agar data yang dikelola dapat diproses dengan cepat untuk mendapatkan informasi lebih lanjut. Maka dari itu penulis ingin menggunakan suatu sistem dengan metode *clustering* untuk membantu dalam mengolah data kasus keracunan biologis mengelompokkan obyek-obyek berdasarkan karakteristik yang dimiliki masing- masing obyek.

2. KAJIAN PUSTAKA

Data Mining

Data Mining Merupakan serangkaian proses untuk menggali nilai tambah dari suatu kumpulan data berupa pengetahuan yang selama ini tidak diketahui secara manual dari suatu kumpulan data. (Relita Buaton et al., 2019)

Data mining merupakan proses penggalian informasi dan pola yang bermanfaat dari data yang sangat besar. *Data mining* mencakup pengumpulan data, ekstraksi data, analisis data, dan statistik data. *Data mining* juga dikenal sebagai *Knowledge discovery*, *Knowledge extraction*, *data/pattern analysis*, *information harvesting*, dan lain-lain. (Arhami et al., 2020).

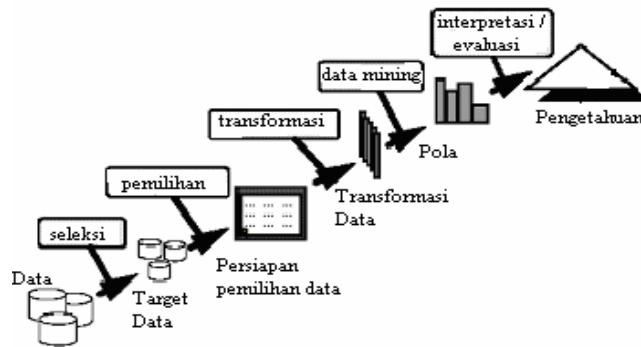
Kelebihan dengan menggunakan *data mining* ialah sebagai berikut:

1. Manajemen hubungan pelanggan yang lebih baik
2. Perkiraan trend pasar
3. Membantu dalam Persaingan
4. Menarik dan mempertahankan pelanggan

Sedangkan kekurangan dari *data mining* adalah sebagai berikut:

1. Keamanan data berisiko untuk diretas oleh pihak yang tidak diinginkan.
2. Data mining melanggar privasi pengguna. Hal tersebut wajib kita waspadai mengingatkan profil pengguna dapat dijual, di perdagangan, atau dapat digunakan sebagai alat untuk mencari keuntungan yang merugikan orang. dan lain-lain. (Prasetyo, 2012)

Ada beberapa tahapan proses dalam *data mining*. Diagram di bawah menggambarkan beberapa tahap / proses yang berlangsung dalam *data mining*. Fase awal dimulai dari data sumber dan berakhir dengan adanya informasi yang dihasilkan dari beberapa tahapan, yaitu :



Gambar 1 Fase-Fase dalam Data Mining

Sumber : (Prasetyo, 2012)

Tahapan proses dalam *Data Mining* dapat dijelaskan sebagai berikut :

1. Seleksi Data

Pemilihan (seleksi) data dari sekumpulan data operasional perlu dilakukan sebelum tahap penggalian informasi dalam KDD dimulai. Data hasil seleksi yang akan digunakan untuk proses *data mining*, disimpan dalam suatu berkas, terpisah dari basis data operasional.

2. *Pre-processing/ Cleaning* (pemilihan data)

Sebelum proses data mining dapat dilaksanakan, perlu dilakukan proses *cleaning* pada data yang menjadi fokus KDD. Proses *cleaning* mencakup antara lain membuang duplikasi data, memeriksa data yang inkonsisten, dan memperbaiki kesalahan pada data, seperti kesalahan cetak (*typografi*). Juga dilakukan proses *enrichment*, yaitu proses “memperkaya” data yang sudah ada dengan data atau informasi lain yang relevan dan diperlukan untuk KDD, seperti data atau informasi eksternal.

3. Transformasi

Coding adalah proses transformasi pada data yang telah dipilih, sehingga data tersebut sesuai untuk proses data mining. Proses coding dalam KDD merupakan proses kreatif dan sangat tergantung pada jenis atau pola informasi yang akan dicari dalam basis data

4. *Data mining*

Data mining adalah proses mencari pola atau informasi menarik dalam data terpilih dengan menggunakan teknik atau metode tertentu. Teknik, metode, atau algoritma dalam data mining sangat bervariasi. Pemilihan metode atau algoritma yang tepat sangat bergantung pada tujuan dan proses KDD secara keseluruhan.

5. *Interpretasi / Evaluasi*

Pola informasi yang dihasilkan dari proses data mining perlu ditampilkan dalam bentuk yang mudah dimengerti oleh pihak yang berkepentingan. Tahap ini merupakan bagian dari proses KDD yang disebut dengan *interpretation*. Tahap ini mencakup pemeriksaan apakah pola atau informasi yang ditemukan bertentangan dengan fakta atau hipotesa yang ada sebelumnya. (Relita Buaton et al., 2019)

Metode Algoritma K-Means (Clustering)

Algoritma K-Means ditemukan oleh beberapa orang yaitu *Lloyd* (1957), *Forgey* (1965), *Friedman* dan *Rubin* (1967), dan *McQueen* (1967). Ide dari pengelompokan (*Clustering*) pertama kali ditemukan oleh *Lloyd* pada tahun 1957, namun hal tersebut baru dipublikasi pada tahun 1982. Pada tahun 1965 *Forgey* juga mempublikasikan teknik yang sama sehingga terkadang dikenal sebagai *Lloyd-Forgey*. (Wanto et al., 2020)

K-Means merupakan salah satu *algoritma clustering* yang masuk dalam kelompok *Unsupervised learning* yang digunakan untuk membagi data menjadi beberapa kelompok dengan sistem partisi. *Algoritma* ini menerima masukan berupa data tanpa label kelas. Hal ini berbeda dengan *K-Nearest Neighbor* (KNN) dan *algoritma supervised learning* lainnya yang menerima masukan berupa vektor. Pada *algoritma K-Means*, komputer mengelompokkan sendiri data-data yang menjadi masukannya tanpa mengetahui terlebih dahulu target kelasnya. Masukan yang diterima adalah data atau objek dan k buah kelompok (*cluster*) yang diinginkan. *Algoritma* ini akan mengelompokkan data atau objek kedalam sebuah kelompok tersebut. (Wanto et al., 2020)

Pada setiap *cluster* terdapat titik pusat (*Centroid*) yang mempresentasikan *cluster* tersebut. Secara sederhana *algoritma K-Means* dapat dijelaskan sebagai *algoritma data mining* yang digunakan untuk menyelesaikan masalah pengelompokan (*Clustering*). Pada pemrosesan data *algoritma K-Means Clustering*, akan diawali dengan pengelompokan *Centroid* pertama yang dipilih secara acak sebagai titik awal untuk setiap *cluster*, kemudian menghitung secara berulang agar posisi *Centroid* optimal. (Wanto et al., 2020)

Adapun langkah-langkah dalam pengelompokan data dengan *Algoritma K-Means* adalah sebagai berikut: (Relita Buaton et al., 2019)

1. Menentukan Jumlah cluster (k) pada data set.
2. Menentukan nilai Pusat (centroid)
3. Hitung jarak dekat dengan centroid
4. Jarak centroid yang digunakan adalah *dEuclidean Distance*, dan *dManhattan* dengan rumus seperti dibawah ini:

Rumus *dEuclidean*;

$$d_{ij} = \sqrt{(x_{1i} - x_{1j})^2 + (x_{2i} - x_{2j})^2 + \dots + (x_{ki} - k_j)^2} \dots\dots\dots (1)$$

Keterangan:

d_{ij} = jarak da data ke i ke pusat cluster j

x_{kj} = data dari ke-i pada *attribute* data ke-k

x_{kj} = data dari ke-j pada *attribute* data ke-k

Rumus *dManhattan*;

$$X,Y = \sqrt{\sum i |X_i - Y_i|} \dots\dots\dots (2)$$

Clustering merupakan suatu metode untuk mencari dan mengelompokkan data yang memiliki kemiripan karakteritik (*similarity*) antara satu data dengan data yang lain. *Clustering* merupakan salah satu metode data mining yang bersifat tanpa arahan (*unsupervised*), maksudnya metode ini diterapkan tanpa adanya latihan (*taining*) dan tanpa ada guru (*teacher*) serta tidak memerlukan target *output*. Dalam data mining ada dua jenis metode clustering yang digunakan dalam pengelompokan data, yaitu *hierarchical clustering* dan *non-hierarchical clustering*.

Metode *hierarchical clustering* adalah suatu metode pengelompokan data yang dimulai dengan mengelompokkan dua atau lebih objek yang memiliki kesamaan paling dekat. Kemudian proses diteruskan ke objek lain yang memiliki kedekatan kedua. Demikian seterusnya sehingga cluster akan membentuk semacam pohon dimana ada hierarki (tingkatan) yang jelas antar objek, dari yang paling mirip sampai yang paling tidak mirip.

Selanjutnya, berbeda dengan metode *hierarchical clustering*, metode *nonhierarchical clustering* justru dimulai dengan menentukan terlebih dahulu jumlah cluster yang diinginkan (dua cluster, tiga cluster, atau lain sebagainya). Setelah jumlah cluster diketahui, baru proses cluster dilakukan tanpa mengikuti proses hierarki. (Relita Buaton et al., 2019)

3. METODE PENELITIAN

Dalam menerapkan metode *clustering*, proses awal yang dilakukan dalam pembentukan cluster adalah melakukan transformasi data kedalam bentuk numerik dengan kode-kode yang telah di tentukan jumlah *group* (K), hitung *centroid*, hitung jarak ke *centroid* dan kemudian groupkan berdasarkan jarak terdekat, jika tidak ada objek yang pindah *group* maka *iterasi* selesai. Kemudian akan dilakukan inisialisasi pada data agar mudah dalam melakukan tranformasi data usia, jenis keracunan, dan faktor penyebab yaitu sebagai berikut:

Berikut ini merupakan data-data yang diperoleh selama proses pengumpulan data dan diambil 20 data secara acak sebagai sampel. Adapun data yang digunakan yaitu sebagai berikut:

Tabel 1 Data Pendukung Penelitian

No.	Usia	Jenis Keracunan	Penyebab
1	12 Tahun	kesalahan dalam pengolahan makanan	jamur beracun
2	6 Tahun	susu / daging yang tidak dimasak	bakteri Penyebab gastroenteritis atau radang lambung dan usus
3	8 Tahun	makanan siap santap, keju, sosis dan yogurt	bakteri Gram-positif berbentuk batang yang dapat ditemukan di lingkungan seperti tanah, air, dan pada hewan
4	10 tahun	makanan kadaluarsa	bakteri Anaerob gram-positif yang dapat menghasilkan racun botulinum
5	21 Tahun	daging mentah	bakteri Gram-negatif berbentuk batang yang biasanya ditemukan di usus manusia dan hewan berdarah panas
6	22 Tahun	susu / daging yang tidak dimasak	bakteri Gram-negatif berbentuk batang yang dapat menyebabkan infeksi pada manusia dan hewan
7	23 Tahun	salad buah	bakteri Gram-positif berbentuk bulat yang biasanya membentuk kelompok seperti anggur
8	25 tahun	kesalahan dalam pengolahan makanan	jamur beracun
9	15 Tahun	kue	bakteri Gram-positif berbentuk bulat yang biasanya membentuk kelompok seperti anggur
10	31 Tahun	susu / daging yang tidak dimasak	bakteri Gram-negatif berbentuk batang yang dapat menyebabkan infeksi pada manusia dan hewan
11	32 Tahun	roti isi	bakteri Gram-positif berbentuk bulat yang biasanya membentuk kelompok seperti anggur

No.	Usia	Jenis Keracunan	Penyebab
12	8 tahun	makanan kadaluarsa	bakteri Anaerob gram-positif yang dapat menghasilkan racun botulinum
13	35 tahun	susu / daging yang tidak dimasak	bakteri Penyebab gastroenteritis atau radang lambung dan usus
14	12 tahun	kesalahan dalam pengolahan makanan	jamur beracun
15	6 tahun	susu / daging yang tidak dimasak	bakteri Penyebab gastroenteritis atau radang lambung dan usus
16	8 tahun	makanan siap santap, keju, sosis dan yogurt	bakteri Gram-positif berbentuk batang yang dapat ditemukan di lingkungan seperti tanah, air, dan pada hewan
17	12 tahun	makanan kadaluarsa	bakteri Anaerob gram-positif yang dapat menghasilkan racun botulinum
18	28 tahun	daging mentah	bakteri Gram-negatif berbentuk batang yang biasanya ditemukan di usus manusia dan hewan berdarah panas
19	22 tahun	susu / daging yang tidak dimasak	bakteri Gram-negatif berbentuk batang yang dapat menyebabkan infeksi pada manusia dan hewan
20	21 tahun	salad buah	bakteri Gram-positif berbentuk bulat yang biasanya membentuk kelompok seperti anggur

Tabel 2 Kriteria Usia

Kode	Usia	
1	< 11 Tahun	Anak
2	12 - 20 Tahun	Remaja
3	21 - 29 Tahun	Dewasa awal
4	30 - 38 Tahun	Dewasa Akhir
5	39 – 47 Tahun	Lansia Awal
6	> 48 Tahun	Lansia Akhir

Tabel 3 Jenis Keracunan

Kode	Jenis Keracunan
1	Daging Mentah
2	Kesalahan Dalam Pengolahan Makanan
3	Kue Tart
4	Makanan Kadaluarsa
5	Makanan Siap Santap, Keju, Sosis, Yogurt
6	Roti Isi
7	Salad Buah
8	Susu / Daging Yang Tidak Dimasak

Tabel 4 Faktor Penyebab

Kode	Faktor Penyebab
1	Bakteri gastroenteritis atau radang lambung dan usus
2	Bakteri Anaerob gram-positif yang dapat menghasilkan racun botulinum
3	Bakteri Gram-negatif berbentuk batang yang biasanya ditemukan di usus manusia dan hewan berdarah panas
4	Bakteri Gram-positif berbentuk batang yang dapat ditemukan di lingkungan seperti tanah, air, dan pada hewan
5	Bakteri Gram-negatif berbentuk batang yang dapat menyebabkan infeksi pada manusia dan hewan
6	Bakteri Gram-positif berbentuk bulat yang biasanya membentuk kelompok seperti anggur
7	Jamur beracun

Untuk menentukan group dari satu objek, pertama yang harus dilakukan adalah mengukur jarak *Deuclidean* antara dua titik atau objek atau X dan Y yang didefinisikan dalam suatu variabel-variabel yang independen yaitu usia (X), Jenis Keracunan (Y), dan faktor penyebab (Z).

$$Deuclidean (X,Y) = \sqrt{(sX1 - Y2)^2}$$

Tabel 5 Transformasi Data

No	Nama	Usia (X)	Jenis Keracunan (Y)	faktor penyebab (Z)
1	A	2	2	7
2	B	1	8	1
3	C	1	5	4
4	D	1	4	2
5	E	3	1	3
6	F	3	8	5
7	G	3	7	6
8	H	3	2	7
9	I	2	3	6
10	J	4	8	5
11	K	4	6	6
12	L	1	4	2
13	M	4	8	1
14	N	2	2	7
15	O	1	8	1
16	P	1	5	4
17	Q	2	4	2
18	R	3	1	3
19	S	3	8	5
20	T	3	7	6

Selanjutnya langkah yang dilakukan adalah perhitungan data berdasarkan algoritma *k-means clustering*.

Iterasi 1

Centroid 1 = (2, 2, 7) diambil dari secara acak dari data 1

Centroid 2 = (1, 4, 2) diambil dari secara acak dari data 4

Centroid 3 = (3, 1, 3) diambil dari secara acak dari data 5

Tabel 6 Hasil Iterasi 1

No	Nama	Usia (X)	jenis keracunan (Y)	faktor penyebab (Z)	Jarak Dari C1	Jarak Dari C2	Jarak Dari C3	Group
1	A	2	2	7	0	5.48	4.24	1
2	B	1	8	1	8.54	4.12	7.55	2
3	C	1	5	4	4.36	2.24	4.58	2
4	D	1	4	2	5	0	3.74	2
5	E	3	1	3	4.24	3.74	0	3
6	F	3	8	5	6	5.39	7.28	2
7	G	3	7	6	5.20	5.39	6.71	1
8	H	3	2	7	1	5.74	4.12	1
9	I	2	3	6	1.41	4.24	3.74	1
10	J	4	8	5	6.63	5.83	7.35	2
11	K	4	6	6	4.58	5.39	5.92	1
12	L	1	4	2	5.48	0	3.74	2
13	M	4	8	1	8.72	5.10	7.35	2
14	N	2	2	7	0	5.48	4.24	1
15	O	1	8	1	8.54	4.12	7.55	2
16	P	1	5	4	4.36	2.24	4.58	2
17	Q	2	4	2	5.39	1	3.32	2
18	R	3	1	3	4.24	3.74	0	3
19	S	3	8	5	6.40	5.39	7.28	2
20	T	3	7	6	5.20	5.39	6.71	1

Keterangan :

1. Jika pada centroid 1 lebih kecil maka hasil *cluster* masuk pada grup 1.
2. Jika pada centroid 2 lebih kecil maka hasil *cluster* masuk pada grup 2.
3. Jika pada centroid 3 lebih kecil maka hasil *cluster* masuk pada grup 3.

Group lama : {0,0,0,0,0,0,0,0,0,0,0,0,0,0,0,0,0}

Group baru : {1,2,2,2,3,2,1,1,1,2,1,2,2,1,2,2,3,2,1}

Terjadi perubahan *group*, maka dilanjutkan ke iterasi berikutnya:

Untuk group 1 ada 7 data :

$$\text{Grup 1 (X)} = (2+3+3+2+4+2+3)/7 = 2.71$$

$$\text{Grup 1 (Y)} = (2+7+2+3+6+2+7)/7 = 4.14$$

$$\text{Grup 1 (Z)} = (7+6+7+6+6+7+6)/7 = 6.43$$

Untuk group 2 ada 9 data :

$$\text{Grup 2 (X)} = (1+1+1+3+4+1+4+1+1)/9 = 2$$

$$\text{Grup 2 (Y)} = (8+5+4+8+8+4+8+8+5)/9 = 6.36$$

$$\text{Grup 2 (Z)} = (8+5+4+8+8+4+8+8+5)/9 = 2.91$$

Untuk group 3 ada 2 data :

$$\text{Grup 3 (X)} = (3+3)/2 = 3$$

$$\text{Grup 3 (Y)} = (1+1)/2 = 1$$

$$\text{Grup 3 (Z)} = (3+3)/2 = 3$$

Iterasi 2

$$\text{Centroid 1} = (2.71, 4.14, 6.43)$$

$$\text{Centroid 2} = (2, 6.36, 2.91)$$

$$\text{Centroid 3} = (3, 1, 3)$$

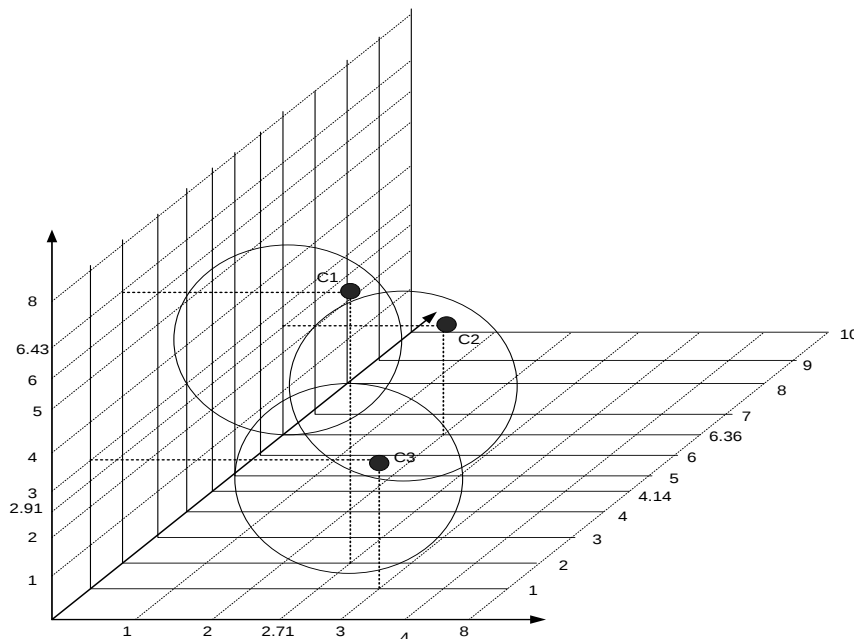
Tabel 7 Hasil Iterasi II

No	Nama	Usia (X)	Jenis Keracunan (Y)	faktor penyebab (Z)	Jarak Dari C1	Jarak Dari C2	Jarak Dari C3	Group
1	A	2	2	7	2.33	5.98	4.24	1
2	B	1	8	1	6.87	2.70	7.55	2
3	C	1	5	4	3.09	2.01	4.58	2
4	D	1	4	2	5	2.72	3.74	2
5	E	3	1	3	4.65	5.45	0	3
6	F	3	8	5	4	2.85	7.28	2
7	G	3	7	6	2.91	3.32	6.71	1
8	H	3	2	7	2.24	6.07	4.12	1
9	I	2	3	6	1.41	4.57	3.74	1
10	J	4	8	5	4.31	3.33	7.35	2
11	K	4	6	6	2.30	3.71	5.92	1
12	L	1	4	2	4.74	2.72	3.74	2
13	M	4	8	1	6.78	3.21	7.35	2
14	N	2	2	7	2.33	5.98	4.24	1
15	O	1	8	1	6.87	2.70	7.55	2
16	P	1	5	4	3.09	2.01	4.58	2
17	Q	2	4	2	4.48	2.53	3.32	2
18	R	3	1	3	4.65	5.45	0	3
19	S	3	8	5	4.12	2.85	7.28	2
20	T	3	7	6	2.91	3.32	6.71	1

Group lama : {1,2,2,2,3,2,1,1,1,2,1,2,2,1,2,2,2,3,2,1}

Group baru : {1,2,2,2,3,2,1,1,1,2,1,2,2,1,2,2,2,3,2,1}

Karena pada iterasi ke-1 dan ke-2 posisi cluster tidak berubah atau terdapat persamaan, maka perhitungan iterasi dihentikan dan mendapatkan hasil akhir yang dapat disimpulkan adalah sebagai berikut:



Gambar 2 Grafik Cluster

PUSAT

	X	Y	Z
Centroid 1 =	(2.71,	4.14,	6.43)
Centroid 2 =	(2,	6.36,	2.91)
Centroid 3 =	(3,	1,	3)

dan mendapatkan hasil akhir yang dapat disimpulkan dari 20 data terdapat 3 grup yaitu grup 1 terdapat 7 data dan grup 2 terdapat 11 data, dan grup 3 terdapat 2 data. Adapun penjelasan dari hasil 3 grup adalah sebagai berikut:

1. Berdasarkan perhitungan diatas dapat diketahui bahwasannya pada cluster 1 pada kelompok data kasus keracunan biologis terdapat 6 data dengan titik centroid usia (x) 2.71 yaitu 17 - 25 Tahun, titik centroid pada jenis keracunan (y) 4.14 yaitu makanan kadaluarsa, dan titik centroid pada factor penyebab (z) 6.42 yaitu Bakteri Gram-positif berbentuk bulat yang biasanya membentuk kelompok seperti anggur.
2. Berdasarkan perhitungan diatas dapat diketahui bahwasannya pada cluster 2 pada kelompok data kasus keracunan biologis terdapat 11 data dengan titik centroid usia (x) 2 yaitu 12-16 Tahun, titik centroid pada jenis keracunan (y) 6.36 yaitu roti isi, dan

titik centroid pada factor penyebab (z) 2.9 yaitu Bakteri Gram-negatif berbentuk batang yang biasanya ditemukan di usus manusia dan hewan berdarah panas.

3. Berdasarkan perhitungan diatas dapat diketahui bahwasannya pada cluster 2 pada kelompok data kasus keracunan biologis terdapat 11 data dengan titik centroid usia (x) 3 yaitu 17-25 Tahun, titik centroid pada jenis keracunan (y) 1 yaitu daging mentah, dan titik centroid pada factor penyebab (z) 3 yaitu Bakteri Gram-negatif berbentuk batang yang biasanya ditemukan di usus manusia dan hewan berdarah panas.

KESIMPULAN

Berdasarkan penelitian yang dilakukan dengan menggunakan metode algoritma k-means maka dapat diketahui bahwasannya hasil terbanyak terdapat pada cluster 2 dengan data kasus keracunan biologis terdapat 11 data dengan titik centroid usia (x) 2 yaitu 12-16 Tahun, titik centroid pada jenis keracunan (y) 6.36 yaitu roti isi, dan titik centroid pada factor penyebab (z) 2.9 yaitu Bakteri Gram-negatif berbentuk batang yang biasanya ditemukan di usus manusia dan hewan berdarah panas.

REFERENSI

- Adawiah, N., Suliatiyowati, N., & Jajuli, M. (2021). Klasterisasi kasus kekerasan terhadap anak dan perempuan berdasarkan algoritma K-Means. *Generation Journal*, 5(2), 69–80.
- Alfikri, T. R., Fakhriza, M., & Ikhwan, A. (2024). Penerapan algoritma K-Means dalam mengelompokkan penyebaran penyakit di Kecamatan Sei Balai. *Jurnal Teknik Informatika Kaputama (JTIK)*, 8(1), 1–7.
- Arhami, M., & Nasir, M. (2020). *Data mining* (R. Indah Utami, Ed.; 1st ed.). CV Andi Offset.
- Dewi, R. (2023). *Aplikasi Matlab untuk simulasi pengolahan sinyal* (A. Prijono, Ed.; 1st ed.). Zahir Publishing.
- Efori Buulolo, S. Kom., & M. Kom. (2020). *Data mining* (1st ed.). CV Budi Utama.
- Fitriana, N. F. (2021). Gambaran pengetahuan pertolongan pertama keracunan makanan. *Jurnal Kesehatan Tambusai*, 2(3), 173–178.
- Laksana, F., Hidayat, R., & Dewi, Y. (2023). Pengelompokan jumlah kekerasan terhadap anak berdasarkan kecamatan di Kabupaten Banyumas menggunakan penerapan algoritma K-Means. *INDEXIA: Informatic and Computational Intelligent Journal*, 5(2), 123–135.

- Matheos Sariole, F., & Hakim, L. (2024). Klasifikasi barang menggunakan metode clustering K-Means dalam penentuan prediksi stok barang. *Jurnal Sains dan Teknologi*, 5(3), 846–854. <https://doi.org/10.55338/saintek.v5i1.2709>
- Nurhakim, B., Septiani, I., Anam, K., & Pratama, D. (2024). Penerapan algoritma K-Means clustering dalam menganalisis risiko penyakit stroke. *Jurnal Mahasiswa Teknik Informatika*, 8(1), 318–322.
- Otong Kadang, M. (2021). *Algoritma dan pemrograman* (A. K. Muzakir, Ed.; Pertama). Humanities Genius.
- Prasetyo, D. (2012). *Data mining: Konsep dan aplikasi menggunakan MATLAB* (N. WK, Ed.; 1st ed., Vol. 1). ANDI OFFSET.
- Relita Buaton, Zarlis, M., Efendi, S., & Yasin, V. (2019). *Data mining time series* (1st ed., Vol. 1). Wade Group.
- Santoso, T. B., & Hadi, A. S. (2019). Implementasi data mining untuk clustering daerah penyebaran penyakit diare di DKI Jakarta menggunakan algoritma K-Means. *Jurnal Ilmiah FIFO*, 1(2), 1–12.
- Wanto, A., Noor Hasan Siregar, M., Perdana Windarto, A., Hartama, D., Sri Rahayu Ginantra, N. L., Napitupulu, D., Surya Negara, E., Ridwan Lubis, M., Vita Dewi, S., & Prianto Cahyo. (2020). *Data mining: Algoritma & implementasi* (T. Limbong, Ed.; 1st ed.). Yayasan Kita Menulis.
- Yuniar, D. P. (2022). *Kesehatan dan gizi anak usia dini* (B. Adi Laksono, Ed.; 1st ed., Vol. 1). Bayfa Cendekia Indonesia.