

PEMBANGKITAN ATURAN FUZZY MENGGUNAKAN ALGORITMA FUZZY C-MEANS UNTUK PEMILIHAN KONSUMSI PANGAN HARIAN

Rizky Salsabella¹, Candra Dewi², Randy Cahya W³

^{1, 2, 3}Fakultas Ilmu Komputer Universitas Brawijaya

Email : salsabella13@gmail.com¹, dewi_candra@ub.ac.id², rendicahya@ub.ac.id³

Abstrak

Kesehatan merupakan suatu hal yang penting dan merupakan kebutuhan dasar manusia. Namun masih banyak masyarakat yang kurang peduli dengan kesehatannya. Salah satu cara untuk menjaga kesehatan yaitu dengan menjaga pola makan dan juga jenis makanan yang dikonsumsi. Pola makan yang baik yaitu merupakan pola makan yang memenuhi kebutuhan gizi tubuh. Agar gizi tubuh terpenuhi, perlu dilakukan suatu perencanaan konsumsi harian untuk seseorang. Namun untuk melakukan perencanaan konsumsi harian tidaklah mudah, terbilang rumit dan cukup kompleks karena harus benar-benar menyesuaikan dengan jumlah kebutuhan kalori harian setiap orang. Nilai kalori yang dihasilkan turut melibatkan parameter jenis kelamin, usia, berat badan, tinggi badan, dan aktivitas harian. Sistem ini akan membantu masyarakat untuk melakukan perencanaan konsumsi harian mereka berdasarkan nilai kalori dengan menggunakan algoritma Fuzzy C-means dan metode Fuzzy Sugeno. Metode ini memiliki dua proses utama yaitu pelatihan dan pengujian. Proses pelatihan merupakan proses pembangkitan aturan, yang hasilnya digunakan untuk proses pengujian, sedangkan proses pengujian menghasilkan nilai kalori yang digunakan untuk perencanaan menu konsumsi harian. Hasil dari penelitian ini menunjukkan bahwa nilai *error* terendah sistem untuk jenis kelamin perempuan adalah 23,472% dan untuk jenis kelamin laki-laki adalah 17,238%.

Kata kunci: gizi, kalori, pembangkitan aturan, clustering, fuzzy c-means, fuzzy sugeno

Abstract

Health is an important thing and is a basic human need. But there are many people who are still less concerned about their health condition. There is a way to maintain their health condition, by controlling their dietary habit and the type of food that they consumed. A good dietary habit is a diet that can fulfill nutrition needs of each human body. In order to fulfill that, there must be a plan for their daily food consumption. But to planning the daily food consumption is not easy, a bit complicated, and complex because we need to measure daily calories value for each person. Calorie value can be measured by involving some of parameters such as gender, age, weight, height, and daily activity. This system will help people to planning their daily food consumption based on the calorie value of each people by using Fuzzy C-Means Clustering Algorithm and Fuzzy Sugeno. This method has two main processes, which are Training and Testing. Training process is a generating rules process that the result can be used for the testing process, meanwhile the testing process will generates calorie value that can be used for planning the daily food consumption. The result of this study indicates the system's lowest error value for the female category is 23,473 %, while the male category is 17,238 %.

Keywords: nutrition, calorie, generating rules, clustering, fuzzy c-means, fuzzy sugeno

1. PENDAHULUAN

Cara mudah untuk menjaga kesehatan, salah satunya yaitu dengan mengonsumsi makanan yang sehat. Namun banyak diantara kita yang belum mengetahui tentang pola makan yang sehat, baik itu berupa jenis makanan maupun waktu makan, sehingga banyak sekali orang yang terserang penyakit gara-gara salah dalam mengonsumsi makanan dan pola makan yang tidak benar karena banyak penyakit yang kita derita berawal dari kesalahan pola makan. Pola makan yang seimbang merupakan pola makan yang dapat memenuhi kebutuhan gizi tubuh. Asupan makanan yang berlebih dapat mengakibatkan kelebihan berat badan

dan penyakit lain yang disebabkan oleh kelebihan zat gizi. Sebaliknya, asupan makanan yang kurang dari kebutuhan dapat menyebabkan tubuh menjadi kurus dan rentan terhadap penyakit. Kedua kondisi tersebut sama tidak baiknya bagi tubuh (Sulistyoningsih, 2011). Oleh karena itu perlu dilakukan suatu perencanaan konsumsi pangan harian untuk menyeimbangkan zat gizi seseorang.

Untuk membuat suatu perencanaan konsumsi pangan harian tidaklah mudah. Perencanaan konsumsi harian merupakan hal yang kompleks dan terbilang rumit karena harus benar-benar menyesuaikan dengan jumlah kebutuhan kalori harian setiap orang. Seorang ahli gizi memerlukan waktu yang lama dalam merencanakan

konsumsi pangan seseorang karena jumlah kebutuhan kalori seseorang pasti berbeda-beda. Melihat kenyataan yang demikian, maka hal tersebut menjadi dasar bagi penulis untuk membuat sistem yang dapat digunakan untuk menggantikan pakar, yaitu sistem *fuzzy*.

Logika *fuzzy* merupakan suatu bentuk logika yang mempunyai nilai kabur, tidak pasti, dan samar, yang dapat berfungsi untuk menyelesaikan kasus penalaran (Kusumadewi & Purnomo, 2010). Dalam penalaran tersebut terjadi proses evaluasi aturan *fuzzy* untuk menghasilkan *output* setiap aturan. Aturan *fuzzy* yang dievaluasi tidak langsung tersedia, melainkan harus ditentukan terlebih dahulu. Digunakannya teknik pembentukan aturan secara otomatis pada permasalahan kali ini karena pada umumnya aturan milik logika *fuzzy* didefinisikan terlebih dahulu oleh pakar, dan hal ini memerlukan banyak waktu, pengalaman, dan keahlian pakar.

Aturan yang dibangun dengan berguna untuk proses inferensi pada sistem *fuzzy*. Pada kasus pembangkitan aturan ini, digunakan metode inferensi TSK karena memiliki proses perhitungan yang cepat dan efisien. TSK sering digunakan pada sistem yang kompleks karena tidak terlalu membebani sistem dalam melakukan komputasi (Priyono et al, 2007). Berkenaan dengan hal tersebut, inferensi menjadi kekuatan dari sistem *fuzzy* dalam pengambilan keputusan terhadap suatu permasalahan. Sehingga sistem *fuzzy* dapat digunakan untuk menyelesaikan permasalahan yang rumit, termasuk menentukan perencanaan konsumsi pangan harian seseorang.

Beberapa penelitian menggunakan teknik *clustering* dalam pembentukan aturan pada sistem *fuzzy*. Teknik ini menggunakan gagasan pembagian *input* ke dalam suatu daerah untuk membentuk antiseden suatu aturan *fuzzy*. Teknik ekstraksi aturan berbasis *clustering* mempunyai keuntungan dalam hal efisiensi saat melakukan komputasi (Koczy dan Botzheim, 2002). Salah satu algoritma *clustering* yang terkenal adalah *Fuzzy C-Means* (FCM). Algoritma FCM adalah adaptasi dari algoritma *k-means* dengan fungsi keanggotaan halus. Tidak seperti algoritma *k-means*, dimana penentuan setiap titik data didasarkan pada pusat yang terdekat, FCM memungkinkan suatu titik data menjadi bagian untuk semua pusat (Abas & Martinez, 2003).

Pada penelitian yang telah dilakukan sebelumnya, digunakan metode *Fuzzy C-Means* sebagai pembelajaran dalam pembangkitan aturan dan menggunakan Sugeno orde-satu sebagai metode inferensi *Fuzzy* dengan objek penelitiannya berupa data deteksi dini risiko penyakit stroke. *Output* dari sistem ini merupakan nilai deteksi dini risiko penyakit stroke dengan akurasi terbesar yang dihasilkan adalah sebesar 93,33% (Valensia, 2015). Lalu pada penelitian selanjutnya, dilakukan penelitian menggunakan metode yang sama yaitu, *Fuzzy C-Means* dan menggunakan FIS Sugeno orde-

sat. Objek penelitiannya merupakan data mammografi yang didapat dari repositori UCI. Data yang ada kemudian digunakan untuk pembangkitan aturan *Fuzzy* menggunakan *Fuzzy C-Means Clustering*. Nilai akurasi tertinggi yang dihasilkan sebesar 87% (Sayekti, 2014).

Berdasarkan kelebihan pembangkitan aturan serta logika *fuzzy* yang telah dijelaskan, pada penelitian ini dilakukan pemilihan konsumsi pangan berdasarkan nilai kalori yang didapatkan dari perhitungan menggunakan algoritma FCM dan logika *fuzzy*. Selain itu dilakukan pengujian MAPE untuk mengetahui nilai *error* terendah sistem.

2. PEMBANGKITAN ATURAN FUZZY MENGGUNAKAN ALGORITMA FUZZY C-MEANS

Dalam bukunya yang berjudul *An Introduction to Data Mining*, Dr Daniel T. Larose (2005) mendefinisikan *Clustering* sebagai upaya mengelompokkan *record*, observasi, atau mengelompokkan kedalam kelas yang memiliki kesamaan objek (Larose, 2005). Pengklasteran berbeda dengan klasifikasi yaitu tidak adanya variabel target dalam pengklasteran. Pengklasteran tidak mencoba untuk melakukan klasifikasi, mengestimasi, atau memprediksi nilai dari variabel target. Akan tetapi, algoritma pengklasteran mencoba untuk melakukan pembagian terhadap keseluruhan data menjadi kelompok-kelompok yang memiliki kemiripan (*homogeny*), yang mana kemiripan *record* dalam suatu kelompok akan bernilai maksimal, sedangkan kemiripan dengan *record* dalam kelompok lain akan bernilai minimal.

2.1. Fuzzy C-Means Clustering

Fuzzy C-Means merupakan suatu teknik pengklasteran data yang mana keberadaan tiap titik data dalam suatu *cluster* ditentukan oleh derajat keanggotaan. Konsep dasar dari *Fuzzy C-Means* adalah dengan menentukan pusat *cluster* sebagai penanda lokasi rata-rata untuk tiap *cluster*. Dengan memperbaiki pusat *cluster* dan derajat keanggotaan setiap titik data secara berulang, maka akan dapat dilihat bahwa pusat *cluster* akan bergerak menuju lokasi yang tepat. Perulangan ini didasarkan pada minimasi fungsi obyektif yang menggambarkan jarak dari titik data yang diberikan menuju pusat *cluster* yang terbobot oleh derajat keanggotaan titik data tersebut. Algoritma *Fuzzy C-Means* (FCM) adalah sebagai berikut (Kusumadewi & Purnomo, 2010) :

1. Input data yang akan di-*cluster* X berupa matriks berukuran $n \times m$, dimana n adalah jumlah sampel data dan m adalah atribut setiap data. X_{ij} = data sampel ke- i ($i = 1, 2, 3, \dots, n$), atribut ke- j ($j = 1, 2, 3, \dots, m$).
2. Menentukan :
Jumlah *cluster* = c ;

Pangkat = w ;
 Maksimum iterasi = MaxIter ;
 Error terkecil yang diharapkan = ξ ;
 Fungsi obyektif awal = $P_0 = 0$;
 Iterasi awal = $t = 1$;

3. Membangkitkan bilangan random μ_{ik} , $i = 1, 2, \dots, n$; $k = 1, 2, \dots, c$; sebagai elemen-elemen matriks partisi awal U dengan range antara 0 sampai dengan 1. Kemudian, menghitung jumlah setiap kolom. Jumlah setiap kolom bilangan random yang terbentuk dapat dihitung dengan Persamaan (1).

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

Selanjutnya menghitung matriks partisi awal U atau normalisasi nilai μ_{ik} . Matriks partisi awal U dapat dihitung dengan Persamaan (2).

$$\mu_{ik} = \frac{\mu_{ik}}{Q_i} \quad (2)$$

4. Menghitung pusat cluster ke- k (V_{kj}) berdasarkan Persamaan (3), dengan $k = 1, 2, \dots, c$; dan $j = 1, 2, \dots, m$.

$$V_{kj} = \frac{\sum_{i=1}^n ((\mu_{ik})^w * X_{ij})}{\sum_{i=1}^n (\mu_{ik})^w} \quad (3)$$

5. Menghitung fungsi obyektif iterasi ke- t (P_t) berdasarkan Persamaan (4).

$$P_t = \sum_{i=1}^n \sum_{k=1}^c \left(\left[\sum_{j=1}^m (X_{ij} - V_{ij})^2 \right] (\mu_{ik})^w \right) \quad (4)$$

6. Menghitung perubahan matriks partisi berdasarkan Persamaan (5).

$$\mu_{ik} = \frac{\left[\sum_{j=1}^m (X_{ij} - V_{ij})^2 \right]^{-1}}{\sum_{k=1}^c \left[\sum_{j=1}^m (X_{ij} - V_{ij})^2 \right]^{-1}} \quad (5)$$

7. Memeriksa kondisi berhenti :
 Jika : Selisih nilai fungsi obyektif pada iterasi ke- t (P_t) dengan fungsi obyektif pada iterasi ke- $t-1$ (P_{t-1}) atau iterasi kurang dari maksimum iterasi maka kondisi berhenti yang dapat ditunjukkan dengan Persamaan (6) dan Persamaan (7).

$$(|P_t - P_{t-1}| < \xi) \quad (6)$$

$$\text{atau} \\ (t > \text{MaxIter}) \quad (7)$$

Jika tidak : iterasi dilanjutkan $t = t+1$ dan ulangi dari langkah ke-4.

2.2. Analisis Varian

Varian dalam *clustering* menurut Dr. Daniel T. Larose ada dua, yaitu varian didalam cluster (*variance within cluster*) dan varian antar cluster (*variance between cluster*). Cluster yang ideal dapat dilihat dari nilai varian yang kecil semakin kecil nilai varian maka semakin ideal cluster tersebut. Berikut Persamaan 8 untuk *variance cluster*,

Persamaan 9 untuk *variance within cluster*, dan Persamaan 10 untuk *variance between cluster*.

$$V_c^2 = \frac{1}{n_c - 1} \sum_{i=1}^{n_c} (d_i - \bar{d}_i)^2 \quad (8)$$

Dimana:

$V_c^2 = V_i^2$ atau varian pada cluster

$c = 1 \dots k$, dimana k = jumlah cluster

n_c = jumlah data pada cluster

d_i = data ke- i pada suatu cluster

$\bar{d}_i$ = rata-rata data pada suatu cluster

$$V_w = \frac{1}{N - k} \sum_{i=1}^k (n_i - 1) V_c^2 \quad (9)$$

Dimana:

V_w = *variance within cluster*

N = jumlah semua data

k = jumlah cluster

n_i = jumlah data pada cluster ke- i

Untuk mencari V_b^2 digunakan persamaan sebagai berikut:

$$V_b = \frac{1}{k - 1} \sum_{i=1}^k n_i (\bar{d}_i - \bar{d})^2 \quad (10)$$

Dimana :

$\bar{d}_i$ = rata-rata data pada cluster ke- i

$\bar{d}$ = rata-rata dari $\bar{d}_i$

Sehingga, untuk menilai cluster dikatakan baik atau tidak bisa dilihat dari kepadatan sebaran datanya, dengan menggunakan Persamaan (11).

$$V = \frac{V_w}{V_b} \quad (11)$$

2.3. Fuzzy Inference System Sugeno Orde Satu

Untuk membentuk *Fuzzy Inference System* dari hasil *Clustering* ini, kita dapat menggunakan metode inferensi *Fuzzy* sugeno orde-1. Sebelumnya, data yang ada dipisahkan terlebih dahulu antara data pada variabel-variabel *input* dengan data pada variabel *output*. Misalkan jumlah variabel *input* adalah m , dan variabel *output* biasanya bernilai 1. Pada metode ini, akan diperoleh kumpulan aturan yang berbentuk seperti pada Persamaan 12 berikut:

[R1] IF (x_1 is A_{11}) $\circ$ (x_2 is A_{12}) $\circ \dots \circ$ (x_n is A_{1m})

THEN ($z = k_{11}x_1 + \dots + k_{1m}x_m + k_{10}$);

[R2] IF (x_1 is A_{21}) $\circ$ (x_2 is A_{22}) $\circ \dots \circ$ (x_n is A_{2m})

THEN ($z = k_{21}x_1 + \dots + k_{2m}x_m + k_{20}$);

...

[Rr] IF (x_1 is A_{r1}) $\circ$ (x_2 is A_{r2}) $\circ \dots \circ$ (x_n is A_{rm})

THEN ($z = k_{r1}x_1 + \dots + k_{rm}x_m + k_{r0}$); (12)

Dengan:

- A_{ij} adalah himpunan *fuzzy* aturan ke- i variabel ke- j sebagai anteseden,
- k_{ij} adalah koefisien persamaan *output fuzzy* aturan ke- i variabel ke- j ($i = 1, 2, \dots, r$; $j = 1, 2, \dots, m$), dan k_{i0} adalah konstanta persamaan *fuzzy* aturan ke- i ,
- Tanda $\circ$ menunjukkan operator yang digunakan.

2.4. MAPE

MAPE atau *Mean Absolute Percentage Error* merupakan ukuran ketepatan relatif yang digunakan untuk mengetahui presentase penyimpangan hasil peramalan, dengan Persamaan (12):

$$\text{MAPE} = \sum \frac{|X-F|}{X} \times 100\% \quad (12)$$

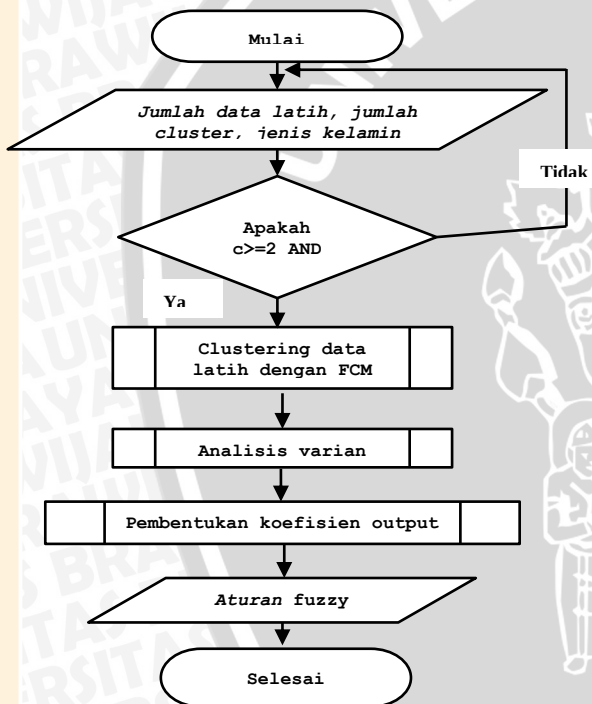
Keterangan :

X : nilai nyata

F : nilai hasil peramalan

3. METODE

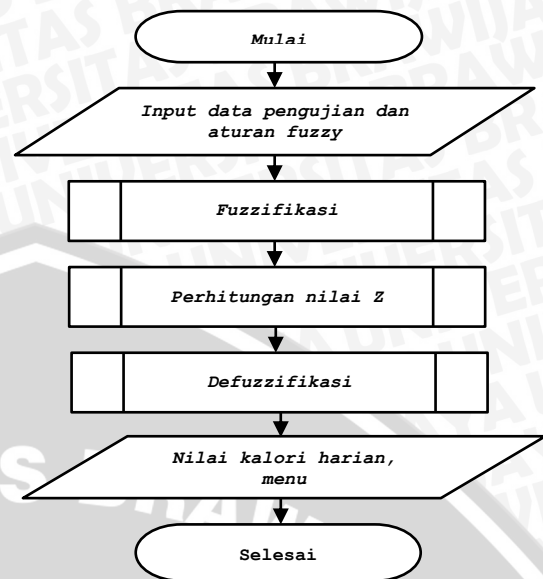
Sistem ini terdiri atas dua proses utama yaitu pelatihan dan pengujian. Alur proses pelatihan terdiri atas 3 proses, yaitu proses *clustering* data kalori harian, proses analisis varian, dan proses pembentukan koefisien output.



Gambar 1. Diagram Alir Proses Pelatihan

Proses pengujian dapat dijalankan apabila proses pelatihan telah dilakukan. Jadi, jika sistem belum pernah melakukan proses pelatihan maka sistem tidak bisa melakukan proses pengujian. Perancangan proses pengujian ditunjukkan gambar Gambar 2.

Alur proses pengujian terdiri atas 3 proses, yaitu proses fuzzifikasi, perhitungan nilai Z, dan defuzzifikasi. Input proses pengujian berupa data pengujian dan aturan hasil pelatihan, sedangkan output proses berupa hasil nilai kalori.



Gambar 2. Diagram Alir Proses Pengujian

Data yang digunakan dalam proses-proses ini adalah data yang didapat dari pakar. Selain itu data menu makanan sehat didapat dari DKBM (Daftar Komposisi Bahan Makanan) oleh PERSAGI (Persatuan Ahli Gizi Indonesia). Atribut data dalam penelitian ini terdiri dari umur, jenis kelamin, berat badan, tinggi badan, dan aktivitas harian.

Jumlah data yang digunakan dalam penelitian ini adalah 140 data latih yang dibagi menjadi dua, yaitu 70 data perempuan dan 70 data laki-laki, serta 15 data uji perempuan dan 15 data uji laki-laki, dimana data latih dan data uji merupakan data yang berbeda. Data latih digunakan untuk bahan pembelajaran bagi algoritma *clustering* dalam pembangkitan aturan *fuzzy*, sedangkan data uji merupakan data yang akan digunakan untuk menghitung nilai kalori harian dan menghitung seberapa besar nilai MAPE yang dihasilkan sistem.

4. HASIL DAN PEMBAHASAN

Untuk pengujian terdiri dari 3 kali uji coba dengan jumlah data latih yang berbeda. Pada proses pelatihan, iterasi maksimum yang ditetapkan adalah 100 dengan kesalahan minimum sebesar 0,0001. Uji coba dilakukan terhadap beberapa jumlah aturan yang berbeda, yaitu pada jumlah *cluster* ideal dan beberapa jumlah *cluster* lain sebagai pembanding.

Proses pengujian pada setiap uji coba menggunakan 15 data uji pada tiap parameter jenis kelamin. Proses ini akan menghasilkan nilai kalori. Hasil proses tersebut kemudian dibandingkan dengan parameter acuan yang ada sehingga dapat diketahui nilai MAPE nya.

Pelatihan dan pengujian pada setiap jumlah aturan dilakukan sebanyak 3 kali percobaan. Perhitungan MAPE pada setiap aturan dilakukan

sebanyak 5 kali. Nilai akhir MAPE akan didapat berdasarkan rata-rata akhir semua percobaan.

1. Uji coba 1, menggunakan 50 data latih.
2. Uji coba 2, menggunakan 60 data latih.
3. Uji coba 3, menggunakan 70 data latih.

2.1 Hasil Uji Coba

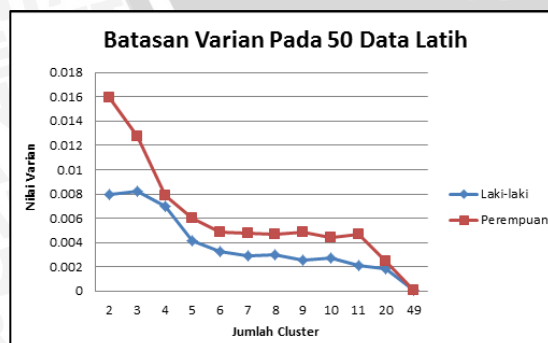
2.1.1 Uji Coba 1

Pada uji coba 1, pelatihan dilakukan mulai dari jumlah *cluster* 2 hingga 49. Untuk menentukan jumlah *cluster* yang ideal, maka dilakukan perhitungan varian untuk setiap jumlah *cluster*. Jumlah *cluster* ideal dipilih berdasarkan nilai varian terkecil. Hasil perhitungan nilai varian mulai dari jumlah *cluster* 2 hingga 49 ditunjukkan dalam Tabel 1.

Tabel 1. Nilai batasan varian uji coba 1

No.	Jumlah Cluster	Batasan Varian	
		Laki-laki	Perempuan
1	2	0.007956	0.015936
2	3	0.008241	0.012707
3	4	0.006954	0.007877
4	5	0.004188	0.006023
5	6	0.003223	0.004869
6	7	0.00295	0.004761
7	8	0.00302	0.004685
8	9	0.00251	0.004839
9	10	0.002739	0.004395
10	11	0.002103	0.004692
11	20	0.0018806	0.0024498
...	Cluster + 1	...	...
48	49	4.893E-05	7.933E-05

Berdasarkan Tabel 1, maka grafik hubungan antara nilai batasan varian dengan jumlah *cluster* dapat dibuat. Grafik nilai batasan varian pada uji coba 1 ditunjukkan oleh Gambar 3.



Gambar 3 Grafik nilai varian pada uji coba 1

Gambar 3 menunjukkan pergerakan nilai varian mulai dari jumlah *cluster* 2 hingga 49. Pada grafik tersebut, nilai varian minimum untuk jenis kelamin laki-laki terjadi pada cluster 49 ($V =$

0.0000489), sedangkan untuk jenis kelamin perempuan terjadi pada cluster 49 ($V = 0.0000793$).

Setelah itu proses pengujian MAPE dilakukan terhadap jumlah cluster 2, 6, 8, 9, 10, 11, 20, 49. Masing-masing sebanyak 5 percobaan. Nilai MAPE ditunjukkan oleh Tabel 2.

Tabel 2 Hasil pengujian MAPE uji coba 1

C	n	$\bar{X}$ Laki-laki (%)	$\bar{X}_L$ total (%)	$\bar{X}$ Perempuan (%)	$\bar{X}_P$ total (%)
2	1	23.41		42.49	
	2	24.7		42.98	
	3	24.71	24.448	42.98	42.882
	4	24.71		42.98	
	5	24.71		42.98	
4	1	30.56		30.47	
	2	29.72		30.41	
	3	29.72	29.888	30.41	30.422
	4	29.72		30.41	
	5	29.72		30.41	
8	1	29.22		48.74	
	2	30.45		48.83	
	3	30.46	30.21	48.83	48.812
	4	30.46		48.83	
	5	30.46		48.83	
9	1	52.1		58	
	2	52.71		58.04	
	3	52.71	52.588	58.04	58.032
	4	52.71		58.04	
	5	52.71		58.04	
10	1	100		100	
	2	100		100	
	3	100	100	100	100
	4	100		100	
	5	100		100	
11	1	100		100	
	2	100		100	
	3	100	100	100	100
	4	100		100	
	5	100		100	
20	1	100		100	
	2	100		100	
	3	100	100	100	100
	4	100		100	
	5	100		100	
49	1	100		100	
	2	100		100	
	3	100	100	100	100
	4	100		100	
	5	100		100	

Keterangan :

C : Cluster ke-

n : Percobaan ke-

$\bar{X}$: Rata-rata MAPE

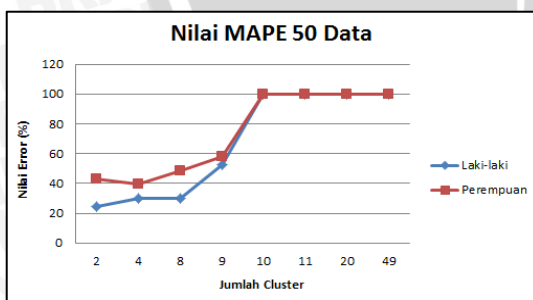
$\bar{X}_L$: Rata-rata MAPE Laki-laki

$\bar{X}_P$: Rata-rata MAPE Perempuan

Berdasarkan Tabel 2 tampak bahwa hasil perhitungan 5 kali percobaan pembentukan aturan pada *cluster* tertentu tidak stabil. Oleh karena itu, dihitung rata-rata nilai MAPE setiap percobaan untuk mewakili nilai MAPE setiap jumlah *cluster*. Jumlah nilai MAPE terkecil untuk laki-laki adalah berada di *cluster* 2 dengan nilai 24,448%, sedangkan untuk perempuan dengan nilai MAPE terendah 30,422% berada di *cluster* 4. Pergerakan nilai MAPE pada uji coba 1 ditunjukkan oleh Gambar 4.4.

Grafik pada Gambar 4.4 menunjukkan bahwa nilai MAPE terendah untuk laki-laki terjadi pada jumlah *cluster* 2 sebesar 24,448 %, sedangkan perempuan pada jumlah *cluster* 4 sebesar 30,422 %. Kemudian nilai MAPE berangsur menaik untuk jenis kelamin laki-laki, dimana pada *cluster* 4 nilai error meningkat menjadi 29,888%, pada *cluster* 8 meningkat menjadi 30,21%, dan pada *cluster* 9 meningkat menjadi 52,588 %, pada *cluster* 10 meningkat menjadi 100 %, dan untuk *cluster* 20 dan 49 nilai MAPE nya tetap 100 %. Untuk jenis kelamin perempuan, dimana pada *cluster* 2 nilai MAPE sebesar 42,882 %, pada *cluster* 4 menurun menjadi 30,422%, pada *cluster* 8 meningkat menjadi 48,812 %, pada *cluster* 9 meningkat menjadi 58,032 %, pada *cluster* 10 meningkat menjadi 100 %, dan pada *cluster* 11, 20 dan 49 nilai MAPE nya tetap sebesar 100 %.

Pada uji coba 1, jenis kelamin laki-laki dengan *cluster* 2 memiliki nilai MAPE lebih rendah dibandingkan dengan jumlah *cluster* ideal (*cluster* 49) yang menghasilkan nilai MAPE tertinggi yaitu 100%. Begitu juga dengan jenis kelamin perempuan, nilai MAPE pada *cluster* 4 lebih rendah dibandingkan dengan *cluster* ideal (*cluster* 49) yang menghasilkan nilai MAPE 100 %.



Gambar 4 Grafik nilai MAPE pada uji coba 1

Kemudian pengujian dengan cara yang sama dilakukan untuk uji coba 2 dan 3. Untuk hasil akhirnya dijelaskan pada analisis hasil.

2.2 Analisis Hasil

Jumlah *cluster* ideal didapat dari melihat nilai varian terendah. Pada semua uji coba, nilai MAPE terendah tidak terjadi pada jumlah *cluster* ideal. Perbandingan nilai error terendah sistem dari semua hasil uji coba 1, 2, dan 3 dijelaskan dalam Tabel 4.3

untuk jenis kelamin laki-laki dan Tabel 4.4 untuk jenis kelamin perempuan.

Tabel 3. Hasil perbandingan nilai MAPE laki-laki

Laki-laki						
Uji Coba	A	Varian	$\bar{X}$ (%)	B	Varian	$\bar{X}$ (%)
1	49	0,0000 489	100	2	0,0079 56	24,44 8
2	59	0,0000 509	100	2	0,0083 54	25,52 8
3	69	0,0003 65	100	4	0,0066 36	17,23 8

Keterangan :

A: nilai error berdasarkan nilai varian terendah

B: nilai error berdasarkan nilai error terendah

$\bar{X}$: Rata-rata MAPE

Tabel 4 Hasil perbandingan nilai MAPE perempuan

Perempuan						
Uji Coba	A	Varian	$\bar{X}$ (%)	B	Varian	$\bar{X}$ (%)
1	49	0,0000 793	48,81 2	4	0,0078 77	30,42 2
2	59	0,0000 625	29,07 2	1 0	0,0035 57	28,29 6
3	69	0,0000 856	24,44 8	4	0,0070 22	23,47 2

Keterangan :

A: nilai error berdasarkan nilai varian terendah

B: nilai error berdasarkan nilai error terendah

$\bar{X}$: Rata-rata MAPE

Berdasarkan Tabel 3 dan 4, dapat dilihat bahwa baik uji coba 1, 2, dan 3 untuk jenis kelamin laki-laki memberikan hasil yang berbeda antara *cluster* ideal menurut sistem dan menurut nilai varian terendah. Semua hasil dari uji coba berdasarkan sistem atau nilai MAPE terendah memiliki hasil lebih baik. Hal ini membuktikan bahwa penentuan jumlah *cluster* atau aturan dengan analisis varian kurang optimal. Tabel 6.7 juga menunjukkan bahwa uji coba 3 dengan jumlah data 70 memberikan hasil nilai MAPE terendah dibandingkan dengan yang lain. Baik jenis kelamin laki-laki maupun perempuan memiliki nilai MAPE terendah pada jumlah *cluster* 4. Maka jumlah *cluster* 4 dianggap sebagai jumlah aturan yang ideal.

5. KESIMPULAN

Hasil proses pelatihan menghasilkan suatu aturan atau *cluster* ideal berdasarkan nilai varian terkecil. Berdasarkan hasil pengujian, pembangkitan

aturan berdasarkan cluster ideal belumlah optimal. Karena nilai MAPE yang dihasilkan oleh aturan dari *cluster* ideal belum tentu merupakan nilai MAPE terkecil, masih ada nilai MAPE yang lebih kecil yang dihasilkan oleh *cluster* lain. Nilai MAPE terendah untuk jenis kelamin laki-laki adalah 17,238 % yang dimiliki oleh uji coba 3 (jumlah data 70) dengan jumlah *cluster* 4. Untuk jenis kelamin perempuan, nilai MAPE terendah yaitu 23,472 % yang dimiliki oleh uji coba 3 (jumlah data 70) dengan jumlah *cluster* 4.

Penyakit Stroke. Universitas Brawijaya, Malang.

6. DAFTAR PUSTAKA

- Abas, F.S. & Martinez, K., 2003. *Classification of Painting Cracks for Content-Based Analysis*. Tersedia di: http://eprints.soton.ac.uk/257294/1/spie_cracks.pdf [Diakses 2 Februari 2016]
- Hardiyansyah & Briawan, D., 1990. *Penilaian dan Perencanaan Konsumsi Pangan*. Bogor: Institut Pertanian Bogor.
- Jang, J. S., Sun, C. T., Mizutani, E., 1997. *Neuro Fuzzy and Soft Computing*. New York: Prentice Hall.
- Kusumadewi, S. & Purnomo, H., 2010. *Aplikasi Logika Fuzzy untuk Pendukung Keputusan*. Edisi Kedua. Yogyakarta: Graha Ilmu.
- Larose, Daniel T., 2005. *Discovering Knowledge in Data: an Introduction to Data Mining*. New Jersey: John & Wiley & Sons, Inc.
- Ludviani, R., 2011. *Pembangkitan Aturan Fuzzy menggunakan Fuzzy C-Means (FCM) Clustering untuk Diagnosa Risiko Penyakit Jantung Koroner (PJK)*. Tersedia di: <http://server3.docfoc.com/uploads/Z2016/05/18/dTTbXHiat9/166ce86374c9d098122f5d988655b13a.pdf> [Diakses 2 Februari 2016]
- Priyono, A., Ridwan, M., Alias, A. J., Rahmat, R. A. O. K., Hassan, A., Ali, M. A. M., *Generation of Fuzzy Rules with Subtractive Clustering*. Tersedia di: <http://www.jurnalteknologi.utm.my/index.php/jurnalteknologi/article/viewFile/782/766> [Diakses 2 Februari 2016]
- Sayekti, E. R., 2014. *Implementasi Algoritma Fuzzy C-Means untuk Pembangkitan Aturan Fuzzy pada Pengelompokan Tingkat Risiko Penyakit Kanker Payudara*. Universitas Brawijaya, Malang.
- Sulistyoningsih, H., 2011. *Gizi Untuk Kesehatan Ibu dan Anak*. Yogyakarta: Graha Ilmu.
- Valensia, S. A., 2015. *Implementasi Algoritma Fuzzy C-Means untuk Pembangkitan Aturan Fuzzy pada Deteksi Dini Risiko*