

PERBANDINGAN ALGORITMA K-MEANS CLUSTERING DENGAN FUZZY C-MEANS DALAM MENGUKUR TINGKAT KEPUASAN TERHADAP TELEVISI DAKWAH SURAU TV

Rio Andika Malik¹⁾ Sarjon Defit²⁾ Yuhandri³⁾

Magister Ilmu Komputer, Universitas Putra Indonesia YPTK Padang
Jl Raya Lubuk Begalung, Padang, Sumatera Barat
Telp (+62)81385636981

Alamat Email : rioandikamalik@gmail.com, sarjondefit@upiyptk.org, yuhandri.yunus@gmail.com

ABSTRAK

Media Televisi Dakwah Surau TV merupakan sebuah media penyiaran yang menyajikan siaran seputar Agama Islam. Media ini akan cepat berkembang karena menyajikan materi penyiaran dalam memenuhi kebutuhan spritual pemirsanya. Peningkatan perkembangan media ini sangat bergantung kepada kepuasan pemirsanya dalam segala bentuk aspek pendukung siaran yang disajikan. Maka diperlukan pengukuran tingkat kepuasan pemirsa sebagai upaya untuk melahirkan peningkatan kualitas siaran yang berkelanjutan. Penelitian ini melakukan perbandingan algoritma clustering dengan pemodelan K-Means Clustering dengan pemodelan Fuzzy C-Means dalam mengelompokkan dan pemetaan dataset yang paling tepat sehingga dapat membantu analisa atau mengukur tingkat kepuasan penonton terhadap media televisi dakwah Surau TV. Perbandingan kinerja algoritma clustering dengan pemodelan K-Means Clustering dan pemodelan Fuzzy C-Means adalah berdasarkan kecepatan proses dan penelusuran nilai parameter RMSE masing-masing algoritma clustering. Hasil penelitian menunjukkan nilai RMSE clustering menggunakan algoritma dengan pemodelan K-Means Clustering adalah sebesar 2.09879 dan besaran nilai penelusuran dari RMSE clustering menggunakan algoritma dengan pemodelan Fuzzy C-Means adalah sebesar 2.07911 dan kecepatan pemodelan Fuzzy C-Means lebih cepat dalam melakukan proses klasterisasi dibandingkan dengan pemodelan K-Means Clustering. Dapat disimpulkan bahwa pengelompokan hasil cluster dengan pemodelan Fuzzy C-Means mampu menghasilkan keakuratan klaster yang lebih tepat dibandingkan dengan hasil cluster pemodelan K-Means Clustering.

Kata kunci: Clustering; K-Means; Fuzzy C-Means; Survey Tingkat Kepuasan; RMSE

ABSTRACT

Da'wah Television Surau TV is a broadcasting media that presents broadcasts around Islam. This media will quickly develop as it presents broadcasting material in meeting the spiritual needs of its viewers. To Increased media development is highly dependent on the satisfaction of the audience in all aspects of broadcast supporting. It is therefore, to measure the level of audience satisfaction as an effort to generate continuous broadcast quality improvement. This research is performing of algorithm clustering comparison with K-Means Clustering modeling and Fuzzy C-Means modeling to classify and mapping the most appropriate dataset so that it can assist analysing or measuring the level of audience satisfaction toward the da'wah television Surau TV. Comparison of clustering algorithm performance with K-Means Clustering modeling and Fuzzy C-Means modeling is based on processing speed and trace value of each RMSE parameter of clustering algorithm. The RMSE result of clustering research using algorithm with K-Means Clustering is 2.09879 and by using algorithm with Fuzzy C-Means model is 2.07911. Fuzzy C-Means modeling speed is faster in conducting the clustering process compared with K-Means Clustering modeling. It can be concluded that clustering with Fuzzy C-Means modeling is able to produce more accurate cluster compared to clustering with K-Means Clustering modeling accuracy

Keywords: Clustering; K-Means; Fuzzy C-Means; Satisfaction rate survey; RMSE

I. PENDAHULUAN

1. Latar Belakang Masalah

Media Televisi Dakwah Surau TV merupakan sebuah media penyiaran yang

menyajikan siaran seputar Agama Islam. Media ini akan cepat berkembang karena menyajikan materi penyiaran dalam memenuhi kebutuhan spritual pemirsanya. Peningkatan perkembangan media ini

sangat bergantung kepada kepuasan pemirsanya dalam segala bentuk aspek pendukung siaran yang disajikan. Dalam perspektif kritis, media televisi tidak dapat dipisahkan dengan nilai, ideologi, ras budaya serta kultur yang terkonstruksi dalam setiap programnya. Program dipandang sebagai pranata dominan yang dapat memainkan kontribusi yang cukup vital dalam kehidupan bermasyarakat. Saat ini ketidakkritisan penonton atau penikmat siaran televisi selama ini, semakin memperlemah posisi tawar penonton di era industri televisi. Pihak pengelola dan pemilik stasiun televisi terus menerus melakukan hal yang stagnan, sehingga berdampak buruk dan merugikan penonton atau penikmat siaran televisi itu sendiri.

Maka diperlukan pengukuran tingkat kepuasan pemirsa sebagai upaya untuk melahirkan peningkatan kualitas siaran yang berkelanjutan. Data Mining dengan metode clustering dapat dimanfaatkan sebagai dalam mengelompokkan komponen apa saja yang bisa memberi pengaruh terhadap penerimaan siaran televisi di masyarakat untuk mengukur tingkat kepuasan penonton terhadap siaran televisi. Pemanfaatan data mining dapat mendeskripsikan tingkat kepuasan terhadap televisi dakwah islam Surau TV tersebut sejauh mana penerimaan siaran berdasarkan hasil pemetaan dan pengelompokan data mentah menjadi garis besar informasi yang bisa dijadikan sebagai aspek penunjang keputusan dalam peningkatan kualitas siaran secara menyeluruh.

Perbandingan kinerja algoritma clustering dengan pemodelan K-Means Clustering dan pemodelan Fuzzy C-Means adalah berdasarkan kecepatan proses dan penelusuran nilai parameter RMSE masing-

masing algoritma clustering. Dalam penelitian sebelumnya analisis perbandingan K-Means dan Fuzzy C-Means mampu menunjukkan seberapa efektif metode yang diusulkan terhadap hasil pengelompokan pemetaan motivasi belajar mahasiswa. Penelitian serupa dengan mengenai perbandingan pengklusteran data menggunakan metode k-means dan fuzzy c-means hasilnya sangat relevan dalam memberikan hasil *cluster* terbaik dari permasalahan yang dihadapi.

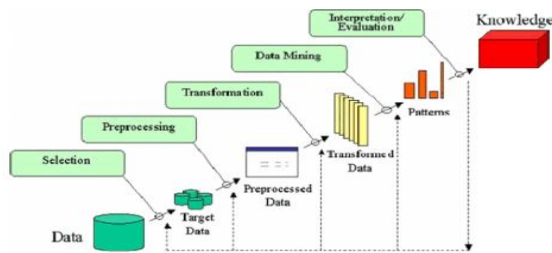
Implementasi data mining dengan menggunakan memanfaatkan algoritma K-means Clustering dan Fuzzy c-means jika dilihat dari beberapa penelitian terdahulu mampu memberikan hasil pengelompokan (*cluster*) data terbaik. Kedua pemodelan tersebut memiliki hasil yang saling signifikan serta juga memiliki beberapa perbedaan dalam hal bentuk dan pola *cluster*, oleh karena itu uji perbandingan antara dua metode data mining pada pemodelan K-means clustering dan Fuzzy C-means untuk menentukan metode algoritma yang terbaik dalam menganalisa tingkat kepuasan terhadap siaran televisi dakwah islam Surau TV dengan cara melihat hasil pemetaan dan pengelompokan data dari klasterisasi yang disajikan oleh masing-masing metode algoritma.

II. STUDI PUSTAKA

1. *Knowledge Discovery in Database (KDD)*

Pada dasarnya data adalah entitas atau objek umum yang tidak mempunyai arti, meskipun sangat memungkinkan memiliki suatu nilai di dalamnya. Data mining merupakan disiplin ilmu yang mempelajari metode untuk mengekstrak

pengetahuan atau menemukan pola dari suatu data (*Knowledge Discovery in Database (KDD)*) [4].



Gambar 1. Proses *Knowledge discovery in database (KDD)*

Pada Gambar 1 proses ekstraksi pengetahuan dari data (KDD) dalam menghasilkan pengetahuan dan terdiri dari beberapa acuan, tahapan atau langkah yang secara garis besar diinterpretasikan sebagai berikut [5]:

1. Data Selection

Pemilihan atau seleksi data dari kumpulan data mentah yang terhimpun perlu dilakukan sebelum memulai tahapan penggalian informasi dalam Knowledge Discovery in Database (KDD). Data yang dihasilkan dari proses seleksi data yang akan dimanfaatkan dalam sebuah proses memaknai data atau data mining, kemudian dimuat dalam sebuah berkas terpisah dari basis data secara keseluruhan yang terhimpun.

2. Pre-processing / Cleaning

Tahapan berikutnya dalam proses penemuan knowledge agar proses data mining dapat dilakukan, perlu dilaksanakan proses pembersihan terhadap data hasil seleksi yang menjadi acuan utama. Proses pembersihan data mencakup antara lain mencari serta menghilangkan duplikasi data, pemeriksaan terhadap data yang inkonsisten, serta memperbaiki kesalahan pada data misalnya perbaikan terhadap kesalahan cetak (typografi).

3. Transformation

Dalam tahapan ini dilakukan pengkodean sebagai bentuk transformasi terhadap data hasil keluaran dari proses cleaning pada tahap 2, sehingga data bersih tersebut selaras untuk kemudian selanjutnya akan dilakukan proses data mining. Proses coding dalam penemuan knowledge merupakan proses yang tidak baku/proses kreatif dan sangat tergantung pada pola atau jenis informasi yang akan dicari dalam keseluruhan basis data.

4. Data Mining

Data mining adalah tahapan paling utama dalam proses penemuan knowledge (KDD) karena merupakan proses mencari informasi atau pola menarik dalam data terseleksi dengan menggunakan suatu pemodelan, metode, teknik, atau algoritma tertentu. Metode, teknik atau algoritma maupun pemodelan dalam sebuah proses memaknai data sangat beragam. Pemilihan metode, teknik atau algoritma maupun pemodelan yang akurat sangat bersandar kepada proses dan tujuan dari knowledge discovery in database (KDD) secara keseluruhan.

5. Interpretation / Evaluation

Pola atau gambaran informasi yang diperoleh dari keseluruhan proses data mining sangat perlu disajikan dalam potret yang mudah dimengerti dan dipahami oleh pihak yang berkepentingan. Dalam tahap ini juga meliputi pemeriksaan apakah informasi atau pengetahuan yang ditemukan/dihasilkan bertentangan dengan berbagai fakta maupun hipotesis yang ada sebelumnya.

Konsep utama dari sebuah transformasi dalam tahapan data mining adalah suatu kumpulan data yang bersumber dari database yang berukuran besar yang diekstrak dan dirangkum untuk menemukan

suatu pola berupa informasi yang berguna dan mengandung sebuah pengetahuan yang bisa dimanfaatkan untuk membantu lini strategis/pengambil keputusan dalam mengambil kebijakan yang tepat untuk kepentingan sebuah organisasi.

2. K-Means Clustering

Metode, algoritma, atau pemodelan *K-Means clustering* merupakan metode *clustering* yang paling umum dan sederhana hal ini disebabkan K-means mempunyai kemampuan mengelompokkan dan memetakan data dalam jumlah yang cukup besar dengan waktu komputasi yang relatif efisien dan cepat [6]. Pengertian dari *K* dalam *K-Means clustering* dimaksudkan sebagai sebuah konstanta jumlah kluster yang dibentuk atau diinginkan, sedangkan *Means* adalah nilai suatu rata-rata dari sebuah kelompok data yang dalam hal ini didefinisikan sebagai *cluster*, sehingga dapat *K-Means Clustering* dapat didefinisikan sebagai suatu metode pemodelan data mining dan metode dalam penganalisaan data yang mengelompokkan data dengan sistem partisi atau yang melakukan proses pemetaan tanpa supervisi (*unsupervised*). Metode *K-Means Clustering* berusaha mengelompokkan atau memetakan data yang ada kedalam tiap kelompok, dimana data dari satu kelompok memiliki karakteristik yang sama dengan data lainnya serta memiliki karakteristik yang berbeda dengan data yang terdapat didalam kelompok yang lainnya [7].

Tahapan dari proses analisis data set menggunakan pemodelan algoritma K-Means Clustering ini adalah dengan menggambarkan langkah-langkah pemodelan k-means clustering dari awal algoritma dimulai hingga menghasilkan

pengelompokan data akhir (pada saat iterasi ke-*n* saat tidak terjadi perubahan pusat cluster/centroid) dengan taksiran bahwa tolak ukur terhadap input adalah jumlah data set sebanyak *n* data dan jumlah inisialisasi centroid (pusat kluster) hingga pada saat iterasi ke-*n* saat tidak terjadi perubahan pusat cluster/centroid menggunakan persamaan Ecludien Distance.

Berikut merupakan algoritma K-Means Clustering

1. Masukkan data yang akan dikluster atau dikelompokkan.
2. Tentukan nilai *K* sebagai jumlah kluster yang akan dibentuk.
3. Inisialisai *K* dari data sebanyak jumlah kluster secara acak sebagai pusat kluster (*centroid*).
4. Hitung jarak antara masing-masing data dengan pusat kluster (*Centroid*), dengan menggunakan persamaan *Euclidean Distance*.

$$d(x, y) = \left[\sum_{i=1}^n |x_n - y_n|^2 \right]^{1/2}$$

Dimana:

$d(x,y)$ = jarak data *x* ke pusat kluster *j*

x_n = data ke *n* pada atribut ke *x*

y_n = titik pusat ke *n* pada atribut

y

n = banyaknya objek

5. Kelompokkan setiap data berdasarkan jarak terdekat antara data dengan *centroid*nya.
6. Tentukan posisi pusat kluster (*centroid*) baru (*k*)

Jika pusat *cluster* tidak berubah maka proses kluster telah selesai, jika belum maka ulangi langkah ke-4 sampai pusat *cluster* (*centroid*) tidak berubah lagi.

3. Fuzzy C-Means

Fuzzy C-Means didasarkan pada teori logika fuzzy. Lotfi Zadeh (1965) memperkenalkan teori dari pemodelan ini pertama kali dimana keanggotaan dari data tidak secara tegas diberi nilai dengan 0 (tidak menjadi anggota kluster) dan nilai 1 (menjadi anggota kluster), namun dengan sebuah nilai derajat keanggotaan yang batasan nilainya 0 sampai 1. Tahapan awal dari konsep dasar *Fuzzy C-Means* yang paling awal adalah dengan menentukan pusat *cluster* (*centroid*) yang akan mengidentifikasi lokasi/ruang rata-rata untuk tiap-tiap kluster. Dalam kondisi awal, pusat kluster ini belum dapat dikatakan akurat hal ini diakibatkan oleh setiap data memiliki derajat keanggotaan untuk masing masing kluster. Perbaikan terhadap pusat *cluster*(*centroid*) dan masing-masing nilai keanggotaan data dengan perulangan, akan terlihat bahwa pusat *cluster*(*centroid*) akan bergerak mendekati ruang/lokasi yang tepat [8].

Berdasarkan pada minimisasi terhadap fungsi rasiona yang mencitrakan jarak yang diberikan ke *centroid* atau pusat cluster dari titik data dengan cara memperbaiki *centroid* (pusat kluster) dan nilai keanggotaan masing-masing data secara *repetitive* atau berulang maka posisi pusat kluster (*centroid*) yang tepat akan dapat ditemukan [9].

Tahapan pemodelan Fuzzy C-Means dari awal algoritma dimulai seperti menentukan masing-masing jumlah cluster, fungsi objektif awal, iterasi awal, iterasi maksimum, pangkat, error terkecil yang diharapkan, memunculkan bilangan acak, menghitung jumlah dari masing-masing kolom dan kemudian menghitung pusat kluster ke-k hingga menghasilkan

pengelompokan data akhir.

Berikut merupakan algoritma Fuzzy C-Means:

1. Masukkan data yang akan dilakukan klusterisasi.
2. Tentukan masing-masing jumlah cluster (c), , fungsi objektif awal ($P_0 = 0$), iterasi awal ($t = 1$), iterasi maksimum ($MaxIter$), pangkat (w), error terkecil yang diharapkan (ϵ), Munculkan bilangan acak (μ_{ik}) Hitung jumlah dari masing-masing kolom.

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

Dengan $j=1,2,\dots,m$; maka hitung :

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

3. Menghitung pusat kluster ke- k (V_{kj}) menggunakan persamaan berikut:

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

4. Lakukan perhitungan terhadap fungsi rasional pada Iterasi ke- t (P_t) menggunakan persamaan berikut:

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

5. Melakukan perhitungan terhadap perubahan matriks partisi

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

6. Cek kondisi berhenti:
 - a. Jika ($|P_t - P_{t-1}| < \epsilon$) atau ($t > MaxIter$) maka berhenti
 - b. Jika tidak : $t=t+1$, maka ulangi langkah ke-4

4. Root Means Square Error

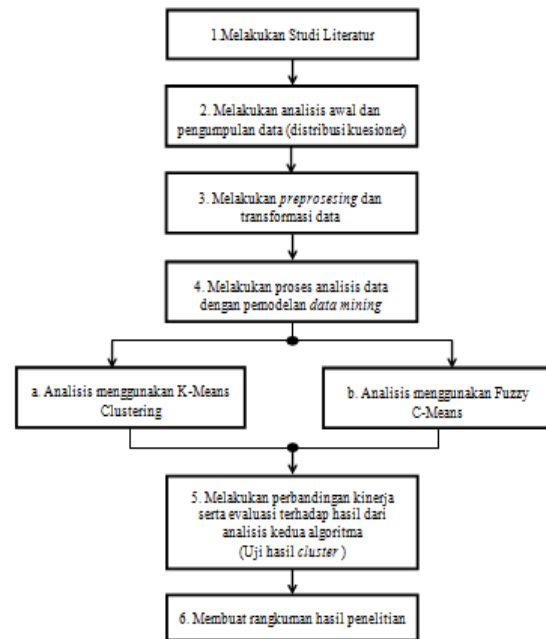
Untuk melakukan evaluasi terhadap sebuah nilai hasil dari pengukuran dengan

nilai sebenarnya atau nilai dianggap benar maka diperlukan menggunakan parameter *root mean ssquare error (RMSE)*. Menurut Febrianti,F.dkk (2016) semakin kecil nilai RMSE, maka pengklasteran data semakin mendekati benar. Perbandingan dari pemodelan *K-Means Clustering* dan pemodelan *Fuzzy C-Means* ini tidak berhenti pada perhitungan dari algoritmanya saja, akan tetapi perbandingan kedua pemodelan *K-Means Clustering* dan pemodelan *Fuzzy C-Means* ini akan dapat terlihat ketika dilakukan perhitungan terhadap nilai RMSE-nya melalui perhitungan dengan persamaan sebagai berikut :

$$RMSE = \sqrt{\sum_{n=1}^i \frac{(data - predicted)^2}{i}}$$

III. METODE

Dalam melaksanakan sebuah penelitian, dibutuhkan sebuah metode penelitian sehingga dapat meningkatkan kemungkinan berjalannya sebuah kegiatan penelitian dengan lancar. Metode penelitian berperan sebagai kerangka kerja dalam melaksanakan penelitian, sehingga dengan mengikuti kerangka kerja tersebut maka penelitian dapat berjalan secara sistematis dan diharapkan penelitian bisa lebih tepat guna dan dalam jangka waktu yang baik. Berikut merupakan kerangka kerja penelitian ini:



Gambar 2. Metodologi Penelitian

Dari gambar 3.1 Kerangka Kerja Penelitian, uraian dalam sebuah kerangka pemikiran studi ilmiah ini dijelaskan dan dipertegas secara komprehensif terhadap asal-usul faktor peubah (variable) yang diteliti, sehingga variabel-variabel yang tetera di dalam rumusan masalah semakin jelas sumber serta riwayat asal-usulnya, maka masing-masing tahapan dari penelitian ini dapat diuraikan sebagai berikut :

1. Melakukan Studi Literatur

Kajian terhadap literatur terkait penelitian yang dilakukan ini berfungsi sebagai landasan ilmiah untuk memperkuat pendapat peneliti dalam penelitian ini. Studi literatur terhadap penelitian serupa dengan metode yang sama atau penelitian berbeda dengan metode yang sama atau penelitian sama dengan metode yang berbeda juga dijadikan sebagai pedoman maupun pembanding untuk nantinya penelitian ini dapat bernilai lebih dan menjadi penyempurna penelitian sebelumnya atau bahkan dapat dijadikan acuan untuk

penelitian berikutnya. Literatur yang dijadikan acuan pada studi ilmiah ini adalah berupa jurnal ilmiah, karya ilmiah dan beberapa buku yang ilmiah.

2. Melakukan Analisis Awal dan Pengumpulan Data

Dalam tahapan analisis awal terhadap aspek apa saja yang berpengaruh terhadap penelitian ini menjadi langkah awal sebelum dilakukannya pengumpulan data. Analisis awal yang baik sebelum dilakukannya pengumpulan data disini dimaksudkan untuk memperoleh kumpulan data yang tepat sehingga mampu menjadi jawaban dari perumusan masalah penelitian yang telah ditetapkan sebelumnya untuk kemudian dilakukan proses pendistribusian kuesioner sebagai langkah pengumpulan data.

3. Melakukan *Preprocessing* dan Tranformasi Data

Agar proses penelitian berjalan dengan baik, kebutuhan data yang akan diklasterisasi perlu diterjemahkan dan dikonversi kedalam bentuk data yang sesuai dengan metode klasterisasi yang digunakan yaitu algoritma dengan pemodelan *K-Means Clustering* dan pemodelan *Fuzzy C-Means*.

4. Melakukan Proses Analisis Data dengan Pemodelan *data mining*

Pada tahap ini akan diterapkan pengolahan dan pengujian data dari hasil kegiatan 3 yaitu *preprocessing* dan tranformasi data menggunakan *tools / software* MATLAB dengan menerapkan kedua teknik klasterisasi menggunakan algoritma pemodelan *K-Means Clustering* dan menggunakan pemodelan *Fuzzy C-Means*.

5. Melakukan Perbandingan Kinerja Serta Evaluasi Terhadap Hasil dari Analisis Kedua Algoritma (uji hasil *cluster*).

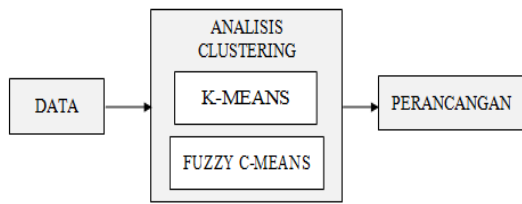
Setelah dilakukannya implementasi terhadap masing-masing metode yaitu algoritma dengan pemodelan *K-Means Clustering* dan pemodelan *Fuzzy C-Means* menggunakan *software* MATLAB, selanjutnya akan dilakukan perbandingan *rules* yang dihasilkan serta kinerja kedua algoritma yang didapat dari kedua metode klasterisasi.

6. Membuat Rangkuman Hasil Penelitian.

Pada tahapan ini kegiatan yang dilakukan adalah membuat kesimpulan menyeluruh baik itu kelebihan dan kekurangan setiap metode serta beberapa saran yang mungkin muncul untuk penelitian berikutnya dari perbandingan kedua metode yaitu algoritma *Clustering* menggunakan parameter *root mean ssquare error (RMSE)* dalam pengelompokan data untuk mengukur tingkat kepuasan terhadap siaran televisi dakwah Surau TV.

IV. HASIL PEMBAHASAN

Penelitian ini melakukan tahapan dari analisa mulai dari data secara keseluruhan serta proses manual pembahasan dari pengolahan data yang akan dilakukan klasterisasi menggunakan pemodelan *K-Means Clustering* dan *Fuzzy C-Means* berdasarkan kerangka kerja penelitian yang terdapat pada bab III Metodologi Penelitian. Secara garis besar tahapan dari proses analisa dan pembahasan pada bab ini dapat digambarkan dalam bagan alir sebagai berikut:



Gambar 3 Bagan Alir Tahapan Analisa dan Pembahasan

Alat bantu analisis, data yang digunakan sebagai data input dalam penelitian ini adalah data hasil pendistribusian kuesioner yang telah disetujui sebelumnya oleh tim redaksi Surau TV yang menjadi objek penelitian. Pengumpulan data primer melalui penyebaran kuesioner merupakan teknik pengumpulan atau penghimpunan data yang dilakukan melalui media tulis, survei online maupun wawancara dengan cara mendistribusikan seperangkat pertanyaan kepada responden. Data dari kuesioner adalah jawaban yang diperoleh dan diberikan oleh responden yang akan ditransformasikan kedalam variabel penelitian.

Data variabel diperoleh dengan menggunakan angket model skala Likert yang disebar dan dinyatakan kepada responden sebagai objek dalam penelitian ini. Berbagai respon yang dihimpun dimaksudkan untuk menjangkau opini, sikap dan pendapat berbentuk opini yang terdiri dari lima skala likert

Pada penelitian ini, skala likert digunakan untuk mengukur tingkat kepuasan menggunakan dua variabel yaitu aspek non teknis dan aspek teknis terhadap kepuasan penonton siaran dakwah Surau TV yang dilakukan oleh peneliti dalam membantu tim redaksi Surau TV adalah dengan dataset sebanyak 50 kuesioner. Kedua aspek yang dijadikan variabel

penelitian diharapkan mampu mewakili keseluruhan kebutuhan dalam penghimpunan data pada penelitian ini. Berikut merupakan notasi profil data hasil pengumpulan kuesioner penelitian ini dapat terlihat adalah data nilai dari profil penonton siaran dakwah Surau TV hasil kolektif pendistribusian kuesioner terhadap penonton siaran dakwah Surau TV yang menjadi objek penelitian tingkat kepuasan pada penelitian ini:

Tabel 1. Notasi Nilai Profil berdasarkan jenis kelamin

Jenis	
Kelamin	Jumlah
Laki-Laki	34
Perempuan	16

Tabel 2 Notasi Nilai Profil berdasarkan Usia

Usia	Jumlah
18-24 Tahun	10
25-32 Tahun	27
33-40 Tahun	6
>40 Tahun	7

Tabel 3 Notasi Nilai Profil berdasarkan Status Pernikahan

Status	
Pernikahan	Jumlah
Menikah	25
Single	25

Tabel 4 Notasi Nilai Profil berdasarkan Pekerjaan

Pekerjaan	Jumlah
Ibu Rumah Tangga	2
Wiraswasta	9
PNS	8
TNI / Polri	2

Karyawan Swasta	18
Pegawai BUMN	2
Pegawai BUMD	3
Dosen	1
Mahasiswa	5

Tabel 5 Notasi Nilai Profil berdasarkan Pendidikan Terakhir

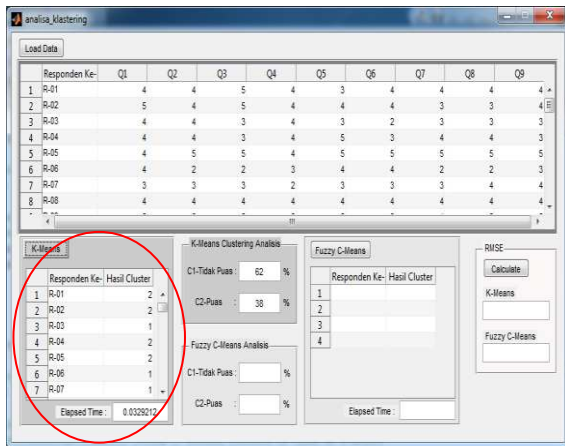
Pendidikan Terakhir	Jumlah
SMA / SMK	6
Diploma	13
Sarjana (S1)	24
Pascasarjana (S2)	7

Data hasil pengumpulan secara kolektif terhadap tanggapan dari responden mengenai aspek teknis maupun aspek non teknis data hasil kuesioner yang telah didistribusikan kemudian dilakukan prapengolahan berupa konversi data nilai atribut penonton tersebut, maka diperoleh notasi nilai profil penonton dapat ditransformasikan kedalam sebuah variable yang mewakili nilai secara keseluruhan berupa id R-*i* dimana *i* adalah data responden ke-*i*. Hasil proses prapengolahan untuk selanjutnya dapat dimanfaatkan sebagai data input dalam analisis *clustering* sebagai berikut:

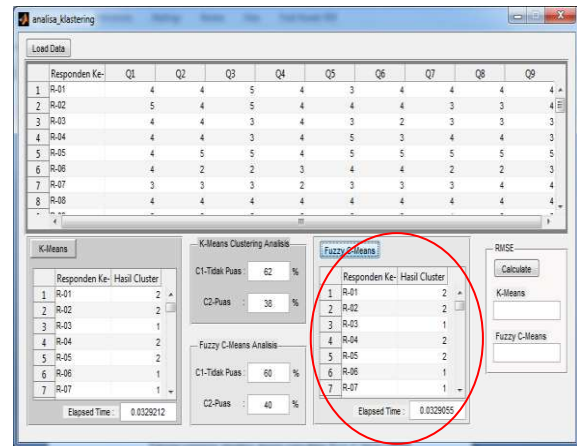
Responden ke	Indikator								
	Q1	Q2	Q3	Q4	Q5	Q6	Q7	Q8	Q9
R-01	4	4	5	4	3	4	4	4	4
R-02	5	4	5	4	4	4	3	3	4
R-03	4	4	3	4	3	2	3	3	3
R-04	4	4	3	4	5	3	4	4	3
R-05	4	5	5	4	5	5	5	5	5
.	.	.	.	.	.	.	.	.	.
.	.	.	.	.	.	.	.	.	.
.	.	.	.	.	.	.	.	.	.
R-49	5	5	4	3	4	4	4	3	3
R-50	3	3	4	3	4	5	4	3	3

1. Pengujian K-Means

Pengujian algoritma dengan pemodelan K-Means Clustering adalah dengan melakukan interaksi dengan GUI analisis_klastering.fig adalah melalui tombol K-Means. Tombol ini merupakan hasil penterjemahan algoritma dan flowchart diagram Algoritma K-Means Clustering kedalam software Matlab 2013a untuk selanjutnya menjadi acuan dalam pembahasan perbandingan dengan algoritma clustering dengan pemodelan Fuzzy C-Means nantinya. Adapun cara kerja dari proses implementasi ke dalam bahasa pemrograman menggunakan Matlab 2013a adalah sesuai dengan Output yang dihasilkan dari interaksi terhadap tombol ini adalah dapat terlihat pada uitable3 dan pada elapsed time akan menampilkan waktu yang dibutuhkan dalam 1x proses clustering. Berikut merupakan tampilan GUI Analisis_klastering.fig melalui interaksi dengan tombol K-Means terlihat pada gambar 4 berikut :



Gambar 4 Tampilan GUI Pengujian Algoritma K-Means Clustering



Gambar 5 Tampilan GUI Pengujian Algoritma Fuzzy C-Means

2. Pengujian Fuzzy C-Means

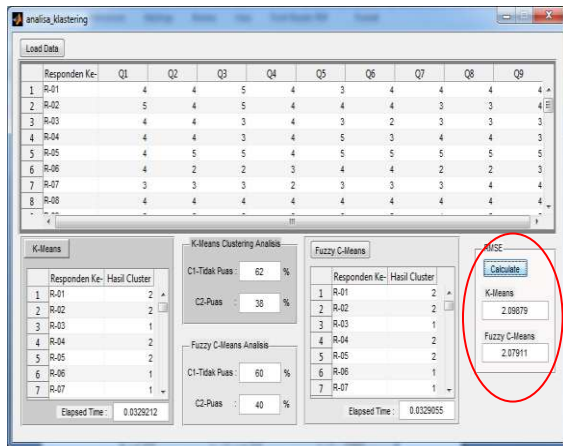
Pengujian algoritma dengan pemodelan Fuzzy C-Means adalah dengan melakukan interaksi dengan GUI analisis_klustering.fig melalui tombol Fuzzy C-Means. Tombol ini merupakan representasi dari pembahasan proses komputasi dari algoritma Fuzzy C-Means. Adapun cara kerja dari proses implementasi ke dalam bahasa pemrograman menggunakan Matlab 2013a adalah sesuai dengan pembahasan mengenai algoritma dan flowchart diagram Fuzzy C-Means yang telah didefinisikan pada BAB sebelumnya. Output yang dihasilkan dari interaksi terhadap tombol ini adalah dapat terlihat pada uitable2 dan pada elapsed time akan menampilkan waktu yang dibutuhkan dalam 1x proses clustering yang untuk selanjutnya dapat menjadi acuan dalam pembahasan perbandingan dengan algoritma clustering dengan pemodelan K-Means Clustering. Berikut merupakan tampilan GUI Analisis_klustering.fig melalui interaksi dengan tombol Fuzzy C-Means terlihat pada gambar 5 berikut :

3. Pengujian RMSE

Pengujian parameter Root Means Square Error (RMSE) adalah dengan melakukan interaksi dengan GUI analisis_klustering.fig adalah melalui tombol RMSE. Tombol ini merupakan representasi dari proses komputasi menggunakan parameter Root Means Square Error (RMSE). Dimana dalam penelitian ini besaran nilai RMSE didapat dari akar penjumlahan kuadrat nilai rata-rata setiap data dikurangi hasil cluster setiap data kemudian dibagi dengan jumlah data dapat terlihat persamaan berikut:

$$RMSE = \sqrt{\frac{\sum_{n=1}^{50} (average - cluster)^2}{50}}$$

Hasil penelusuran nilai Root Means Square Error (RMSE) dari masing-masing algoritma pemodelan K-Means Clustering dan Algoritma pemodelan Fuzzy C-Means yang kemudian nilai tersebut merupakan hasil dari perbandingan kedua algoritma clustering ini berdasarkan parameter dimana semakin kecil nilai RMSE, maka pengklasteran data semakin mendekati benar.



Gambar 6 Tampilan GUI Penelusuran RMSE

Pada Gambar 6 terlihat besaran nilai penelusuran dari RMSE proses clustering menggunakan algoritma dengan pemodelan K-Means Clustering adalah sebesar 2.09879 dan besaran nilai penelusuran dari RMSE proses clustering menggunakan algoritma dengan pemodelan Fuzzy C-Means adalah sebesar 2.07911

V. KESIMPULAN

Setelah melakukan proses pengujian algoritma clustering dengan pemodelan K-Means Clustering dan Fuzzy C-Means yang menghasilkan data hasil klasterisasi dari kedua algoritma clustering tersebut serta perhitungan parameter Root Means Square Error (RMSE) maka dalam penelitian ini dapat ditarik kesimpulan :

1. Kesimpulan

1. Penelitian ini menghasilkan pemetaan tingkat kepuasan penonton terhadap televisi dakwah islam Surau TV tergolong rendah dimana pemodelan K-Means Clustering menghasilkan jumlah anggota klaster C1 (tidak puas) sebesar 62% dan anggota klaster C2(puas) sebesar 38%. Sedangkan menggunakan Fuzzy C-Means menghasilkan jumlah anggota klaster C1 (tidak puas) sebesar

60% dan anggota klaster C2 (puas) sebesar 40%..

2. Pada penelitian ini terlihat bahwa pada data responden ke-50 perbedaan hasil cluster dimana pada pengelompokan data menggunakan pemodelan K-Means Clustering merupakan anggota Cluster ke-1 (Tidak Puas) sedangkan pada pemodelan pengelompokan data menggunakan Fuzzy C-Means merupakan anggota Cluster ke-2 (Puas). Berdasarkan nilai indikator kuesioner data ke-50 terlihat bahwa nilai kepuasan apabila diterjemahkan kedalam skala likert seharusnya lebih tepat menjadi anggota cluster-2 (Klaster dengan kategori Puas) karena apabila diteliti tidak terdapat nilai indikator dari data responden ke-50 yang bernilai tidak puas (2) atau sangat tidak puas (1). Hal ini sesuai dengan hasil pengelompokan dengan menggunakan pemodelan dengan algoritma Fuzzy C-Means. Oleh sebab itu hasil klasterisasi menggunakan pemodelan dengan algoritma Fuzzy C-Means lebih tepat.
3. Perbandingan kecepatan proses masing-masing algoritma clustering dimana terlihat waktu yang dibutuhkan oleh algoritma clustering menggunakan algoritma dengan pemodelan K-Means Clustering adalah 0.0329212 detik dan waktu yang dibutuhkan oleh algoritma clustering menggunakan algoritma dengan pemodelan Fuzzy C-Means adalah 0.0329055 detik. Dapat disimpulkan bahwa pemodelan Fuzzy C-Means dalam penelitian ini lebih cepat dalam melakukan proses klasterisasi dibandingkan dengan pemodelan K-Means Clustering.
4. Besaran nilai penelusuran parameter

RMSE menggunakan bahasa pemrograman Matlab 2013a dimana semakin kecil nilai RMSE, maka pengklasteran data semakin mendekati benar, clustering menggunakan algoritma dengan pemodelan K-Means Clustering adalah sebesar 2.09879 dan besaran nilai penelusuran dari RMSE proses clustering menggunakan algoritma dengan pemodelan Fuzzy C-Means adalah sebesar 2.07911. Dapat disimpulkan bahwa pengelompokan hasil cluster dengan pemodelan Fuzzy C-Means mampu menghasilkan keakuratan klaster yang lebih tepat dibanding dengan pengelompokan hasil cluster dengan pemodelan K-Means Clustering. algoritma lainnya dalam proses KDD.

DAFTAR PUSTAKA

- [1] Rahmat Edi Irawan.(2015), “*Sikap Penonton Dalam Program Televisi Indonesia Saat ini*”, Jurnal Sosiologi Selektif, Vol. 9, No.2, ISSN: 1978-0362 Hal.139-153.
- [2] Nur Indah Selviana dan Mustakim.(2016), “*Analisis Perbandingan K-Means dan Fuzzy C-Means untuk Pemetaan Motivasi Balajar Mahasiswa*”, SNTIKI VII, ISSN: 2085-9902 Hal.95-105.
- [3] Fitria Febrianti, Moh.Hafiyusholeh, Ahmad Hanif Ashyar.(2016), “*Perbandingan Pengklusteran Data Iris Menggunakan Metode K-Means dan Fuzzy C-Means*”, Jurnal Matematika “MANTIK”, Vol. 2 No.1, ISSN: 2527-3159 Hal. 7-13.
- [4] Selviana,Nur Indah dan Mustakim. (2016), “*Analisis Perbandingan K-Means dan Fuzzy C-Means untuk Pemetaan Motivasi Balajar Mahasiswa*”, SNTIKI VII, ISSN: 2085-9902 Hal.95-105.
- [5] Nasari, F. dan Surya Darma (2015), “*Penerapan K-Means Clustering Pada Data Penerimaan Mahasiswa Baru (Studi Kasus: Universitas Potensi Utama)*”, Seminar Nasional Teknologi Informasi dan Multimedia STMIK AMIKOM Yogyakarta, ISSN: 2302-3805, Hal 73-78.
- [6] Dimas Wahyu Wibowo, M. Aziz Muslim, M. Sarosa.”*Perhitungan Jumlah dan Jenis Kendaraan Menggunakan Metode Fuzzy C-means dan Segmentasi Deteksi Tepi Canny*” Jurnal EECCIS Vol. 7, No. 2, Desember 2013. Hal 103-110.
- [7] Aldi Nurzahputra, Much Aziz Muslim,Miranita Khusniati (2017). “Penerapan Algoritma K-Means Untuk Clustering Penilaian Dosen Berdasarkan Indeks Kepuasan Mahasiswa”.Techno.COM, Vol. 16, No. 1, Februari 2017 : 17-24.
- [8] Siska Kurnia Gusti (2012). “Analisis Sebaran Puskesmas Untuk Peningkatan Pelayanan Kesehatan Dengan Metode Fuzzy C-Means”. Seminar Nasional Teknologi Informasi Komunikasi dan Industri (SNTIKI) 4 ISSN : 2085-9902. Hal 78-84.
- [9] Nelson Butarbutar, Agus Perdana Windarto, Dedi Hartama, Solikhun. “Komparasi Kinerja Algoritma Fuzzy C-Means Dan K-Means Dalam Pengelompokan Data Siswa Berdasarkan Prestasi Nilai akademik Siswa (Studi Kasus : Smp Negeri 2 Pematangsiantar)”.JURASIK (Jurnal Riset Sistem Informasi & Teknik Informatika) Volume 1, Nomor 1, Juli 2016. ISSN 2527-5771. Hal 46-55