

Implementasi Algoritma K-Means Untuk Klasterisasi Besarnya Penggunaan Listrik Rumah Tangga Di Pagar Alam

Febriansyah¹, Efan², Arinda Purnama Sari³

^{1,2,3} Program Studi Teknik Informatika, Institut Teknologi Pagar Alam

Email: ¹febriansyahh1213@gmail.com, ²efanzakie85@gmail.com, ³arindapga20@gmail.com

Email Penulis Korespondensi: febriansyahh1213@gmail.com

Abstrak

PLN (Perusahaan Listrik Negara) merupakan perusahaan yang bertanggung jawab atas penyediaan listrik di Indonesia. Salah satunya adalah PT PLN Persero yang ada di Kota Pagar Alam. Penelitian ini bertujuan untuk mengimplementasikan algoritma K-Means dalam mengklasifikasi penggunaan listrik rumah tangga di Kota Pagar Alam. PT PLN Persero, sebagai penyedia utama listrik di Indonesia, menghadapi tantangan dalam memastikan ketersediaan energi yang memadai. Penelitian ini dilatar belakangi dengan belum adanya pengelompokan atau cluster dari data penggunaan listrik rumah tangga. Penelitian ini menggunakan metode CRISP-DM, dimana tahapan meliputi *Business Understanding*, *Data Understanding*, *Data Preparation*, *Modelling*, *Evaluation*, dan *Deployment*. Metode pengujian dengan *elbow method*. Hasil dari penelitian ini didapatkan 3 cluster, yaitu $c_0=377$, $c_1= 34$ dan $c_2= 129$. Hasil perhitungan *Elbow Method* = 276486734427.498. Maka dapat disimpulkan bahwa algoritma K-Means pada kasus pengelompokan penggunaan listrik di kota pagar alam dapat dikatakan sesuai.

Kata Kunci: K-Means; Clustering; CRISP-DM; Elbow Method; PLN.

Abstract

PLN (State Electricity Company) is the company responsible for providing electricity in Indonesia. One of them is PT PLN Persero in Pagar Alam City. This research aims to implement the K-Means algorithm in classifying household electricity use in Pagar Alam City. PT PLN Persero, as the main provider of electricity in Indonesia, faces challenges in ensuring adequate energy availability. This research is motivated by the absence of groupings or clusters of household electricity usage data. This research uses the CRISP-DM method, where the stages include *Business Understanding*, *Data Understanding*, *Data Preparation*, *Modeling*, *Evaluation*, and *Deployment*. Test method using *elbow method*. The results of this research obtained 3 clusters, namely $c_0=377$, $c_1= 34$ and $c_2= 129$. Elbow Method calculation results = 276486734427.498. So it can be concluded that the K-Means algorithm in the case of grouping electricity use in the city of Pagar Alam can be said to be appropriate.

Keywords: Keyword1, Keyword2, Keyword3, Keyword4, Keyword5 (at least 5 words related to the research content separated by commas)

1. PENDAHULUAN

Di era teknologi pada saat ini telah mengalami kondisi kemajuan yang cukup pesat dan mengarah ke bentuk digital. Era digital menyediakan banyak peluang bagi pembangunan bangsa [1]. Salah satu pemanfaatan teknologi saat ini ialah pemanfaatan Big data dengan data mining. Data mining adalah langkah analisis terhadap proses penemuan didalam basis data atau *knowledge discovery in databases* [2]. Data mining juga merupakan proses mencari pola atau informasi menarik dalam data terpilih dengan menggunakan teknik atau metode tertentu. Teknik-teknik, metode-metode, atau algoritma dalam data mining sangat bervariasi. Pemilihan metode atau algoritma yang tepat sangat bergantung pada tujuan dan proses *Knowledge Discovery in Database (KDD)* secara keseluruhan [3]

Algoritma K-Means adalah algoritma Clustering yang paling sederhana dibanding algoritma Clustering yang lain. Kelebihan dari K-Means adalah kesederhanaannya yang membuatnya mudah digunakan dan efisien dalam menemukan Cluster dengan bentuk elips atau bola. Namun, K-Means juga memiliki kelemahan yaitu tidak efektif dalam menanggapi Cluster dengan ukuran atau bentuk yang tidak terdefinisi dengan jelas. [4]. Algoritma juga merupakan inti dari ilmu computer atau komputasi. Berbagai bidang ilmu computer berawal pada istilah algoritma itu sendiri. Akan tetapi, jangan berasumsi bahwa algoritma harus identik dengan komputasi belaka. Pada kehidupan sehari-hari manusia, kita juga memiliki proses algoritmik [5]. Meskipun begitu, algoritma ini masih menjadi pilihan utama dalam banyak kasus Clustering karena kepraktisannya

Clustering adalah teknik pengelompokan data non-hierarki yang memisahkan data ke dalam cluster, mengelompokkan data dengan fitur yang sama bersama-sama dan mengelompokkan data dengan karakteristik yang berbeda ke dalam kelompok yang berbeda [6]. Dengan algoritma ini, kita dapat mengklaster sebuah data salah satunya adalah data tentang penggunaan listrik rumah tangga.

Penggunaan listrik dalam rumah tangga pada saat ini telah menjadi bagian terpenting dari kehidupan sehari-hari. Dalam beberapa tahun terakhir, pabrik akan telah menjadi pusat industri dan komersial yang penting. Dalam meningkatkannya kualitas hidup penduduk, pemerintah setempat telah berupaya meningkatkan

infrastruktur, termasuk jaringan listrik, untuk memenuhi kebutuhan energi rumah tangga. Namun, pemerintah telah mengizinkan masyarakat dalam mengatur penggunaan listrik rumah tangga di Paigair Ailaim.

CRISP-DM merupakan metode yang menggunakan model proses pengembangan data yang banyak digunakan para ahli untuk memecahkan masalah terbukti 3 sampai 4 kali lebih banyak digunakan dibandingkan dengan standar lain yang digunakan [7].

RapidMiner adalah aplikasi atau perangkat lunak yang berfungsi sebagai alat pembelajaran dalam ilmu data mining. Platform dikembangkan oleh perusahaan yang didedikasikan untuk semua langkah yang melibatkan sejumlah besar data dalam bisnis komersial, penelitian, pendidikan, dan pembelajaran. RapidMiner memiliki sekitar 100 solusi pembelajaran untuk pengelompokan, klasifikasi dan analisis regresi [8].

Google Colab adalah sebuah IDE untuk pemrograman Python dimana pemrosesan akan dilakukan oleh server Google yang memiliki perangkat keras dengan performa yang tinggi [9].

Pengertian Python (bahasa pemrograman) merupakan bahasa pemrograman tinggi yang bisa melakukan eksekusi sejumlah instruksi multi guna secara langsung (interpretatif) dengan metode Object Oriented Programming dan juga menggunakan semantik dinamis untuk memberikan tingkat keterbacaan syntax [10].

Berdasarkan hasil observasi dan wawancara di PT PLN Persero yang ada di Kota Paigair Ailaim, selain untuk rumah tangga, Listrik juga banyak digunakan untuk industri, pendidikan, kesehatan atau untuk rumah sakit dan sebagainya. Ada beberapa kasus yang sering terjadi di PLN, Penggunaan listrik yang tidak terkontrol telah menyebabkan beban listrik yang tidak seimbang, sehingga mengakibatkan gangguan pada jaringan listrik. Hal ini juga telah meningkatkan biaya operasional dan pemeliharaan jaringan listrik, serta meningkatkan risiko keamanan dan keselamatan. Salah satu cara untuk mengatasi permasalahan tersebut yang dilakukan adalah dengan menggunakan algoritma k-means untuk klasifikasi penggunaan listrik rumah tangga. Sedangkan belum ada pengelompokan atau cluster dari data penggunaan listrik rumah tangga dengan cluster penggunaan listrik yang tinggi, sedang atau rendah. Dengan itu, dibutuhkan cluster data tersebut agar PLN bisa membuat suatu kebijakan keputusan dalam mengatasi permasalahan energi listrik di Kota Paigair Ailaim. Dibalik ini adalah penelitian terdahulu yang terkait dengan penelitian ini.

[11] Dengan judul “Hubungan Antara Besarnya Daya Listrikterpasang Dengan Banyaknya Pemakaian Listrik Dalam Skala Rumah Tangga”, yaitu tentang banyaknya penggunaan atau pemakaian listrik rumah tangga oleh karena itu pelanggan listrik mempunyai kecenderungan sifat konsumtif terhadap pemakaian listrik.

[12] Dengan Judul “Klasifikasi Data Tenaigai Kerja Terbuka Menurut Provinsi Dengan Penggunaan Algoritma K-Means” yaitu tentang mengklasifikasi data tenaga kerja terbuka menurut provinsi.

[7] Penerapan Algoritma K-Means untuk Klasifikasi Penduduk Miskin pada Kota Paigair Ailaim. Yaitu tentang langkah-langkah mengklasifikasikan penduduk miskin di Kota Paigair Ailaim dengan Metode pengujian sistem menggunakan black box testing.

[13] Dengan judul “Analisis Peluang Penghematan Konsumsi Energi Listrik Pada Pelanggan Rumah Tangga” yaitu tentang Ketersediaan sumber energi yang semakin langka menyebabkan kemungkinan terjadinya krisis energi. Dari hasil analisis tersebut dapat disusun langkah-langkah dan strategi yang akan dilakukan guna mendapatkan penghematan konsumsi energi listrik pada rumah tangga.

[14] Dengan judul “Penerapan Algoritma C4.5 Untuk Memprediksi Besarnya Penggunaan Listrik Rumah Tangga di Kota Baitam”, Kepada penduduk menyebabkan kebutuhan akan energi listrik menjadi sangat tinggi, perhitungan Algoritma C4.5 membentuk pohon keputusan dimana variabel jumlah anggota keluarga, luas bangunan rumah dan lama waktu di rumah menjadi variabel penting dari prediksi besarnya penggunaan listrik.

Berdasarkan penelitian terdahulu diatas, dapat disimpulkan bahwa solusi yang didapatkan ialah akan dilakukannya klasifikasi menggunakan algoritma k-means untuk mengklasifikasi penggunaan listrik rumah tangga di Paigair Ailaim, guna untuk menganalisis pola penggunaan listrik rumah tangga di Paigair Ailaim dan untuk menentukan klastr yang memiliki karakteristik penggunaan listrik yang sama. Dengan demikian, kita dapat mengetahui bagaimana penggunaan listrik rumah tangga di Paigair Ailaim dan bagaimana kita dapat meningkatkan efisiensi penggunaan listrik rumah tangga

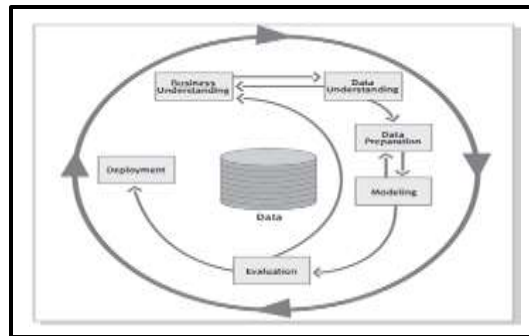
Berdasarkan permasalahan diatas, maka penulis akan melakukan Teknik Klasifikasi dengan menggunakan algoritma K-Means untuk pembaharuan dan juga ingin mengetahui bagaimana penggunaan algoritma k-means dapat membantu dalam mengatasi masalah penggunaan listrik rumah tangga di Paigair Ailaim. Dengan demikian, kita dapat mengetahui apakah penggunaan algoritma k-means dapat membantu dalam meningkatkan efisiensi penggunaan listrik rumah tangga di kota paigair ailaim ini dan bagaimana kita dapat menggunakan algoritma ini untuk mengatasi masalah yang terkait dengan penggunaan listrik rumah tangga di Paigair Ailaim.

Berdasarkan latar belakang diatas tersebut penulis mengambil Judul “Implementasi Algoritma K-Means Untuk Klasifikasi Besarnya Penggunaan Listrik Rumah Tangga Di Paigair Ailaim”. Dari judul tersebut diharapkan data yang di cluster akan bermanfaat bagi semua serta sekaligus memberikan solusi dari permasalahan.

2. METODOLOGI PENELITIAN

2.1 Tahapan Penelitian

Pada penelitian ini penulis menggunakan metode penelitian dengan *CRISP-DM*. *CRISP DM* atau lengkapnya disebut *The Cross Industry Standard Process for Data Mining* merupakan standard yang dikembangkan sejak tahun 1996 di Eropa. *CRISP DM* terdiri dari 6 tahap yang mana 3 tahap paling awal berdasarkan pengalaman penulis dapat *Non Mutually Exclusive* [15]. Imana tahapan meliputi *Business Understanding*, *Data Understanding*, *Data Preparation*, *Modelling*, *Evaluation*, dan *Deployment*.



Gambar 1. Tahapan CRISP-DM

- Business Understanding (Pemahaman Bisnis)**
Ini adalah fase awal di mana pemahaman mendalam tentang proses data mining yang direncanakan dikembangkan, serta penjelasan tujuan penelitian dari sudut pandang bisnis. Tugas pada tahap ini termasuk menetapkan tujuan dan sasaran penelitian, serta memahami lingkup dan konteks di mana data mining akan diterapkan [16].
- Data Understanding (Pemahaman Data)**
Fase pemahaman data dimulai dengan pengumpulan data awal untuk memperdalam pemahaman tentang karakteristik data, menemukan potensi masalah dalam kualitas data, mendapatkan wawasan awal dari data, atau menemukan segmen data yang menarik untuk membentuk hipotesis tentang informasi tersembunyi. Proses ini melibatkan empat langkah, yakni pengumpulan data awal, deskripsi data, eksplorasi data, dan pengecekan kualitas data [17].
- Data Preparation (Persiapan Data)**
Tahapan berikutnya adalah data preparation, yang merupakan serangkaian langkah-langkah untuk mengubah data mentah menjadi data berkualitas, serta menyiapkan data SMS agar sesuai untuk pelatihan model klasifikasi [18].
- Modelling (Permodelan)**
Modelling merupakan tahap pemilihan teknik data mining dengan menentukan algoritma klasifikasi yang akan digunakan untuk membuat model [19].
- Evaluation (Evaluasi)**
Tujuan utama dari fase evaluasi adalah untuk memastikan akurasi dan kebenaran dari tahap pemodelan, sehingga menghasilkan output yang sesuai dengan harapan yang diinginkan. Hasil yang diperoleh terdiri dari dataset yang telah menjalani proses modelling, khususnya menerapkan metode FPGrowth dengan parameter nilai patterns dan rules sesuai dengan yang diharapkan [20].
- Deployment (Penyebaran)**
Menggunakan model yang dihasilkan, pembentukan model tidak menandakan selesainya proyek [21].

2.2 Metode Pengujian

Pada penelitian ini menggunakan metode pengujian *elbow*. Metode *Elbow* merupakan suatu metode yang digunakan untuk menghasilkan informasi dalam menentukan jumlah *cluster* terbaik dengan cara melihat persentase hasil perbandingan antara jumlah *cluster* yang akan membentuk siku pada suatu titik [22]. Metode *elbow* biasa disajikan dalam bentuk grafik untuk mengetahui lebih jelas siku yang terbentuk. Tujuan dari metode *elbow* adalah untuk memilih nilai *k* yang kecil dan masih memiliki nilai *withinss* yang rendah. Nilai *k* pada kombinasi siku dengan *k-means* adalah grafik hubungan *cluster* dengan penurunan error. Jumlah *cluster* yang dihasilkan dari pengujian dengan *k-means* dievaluasi dengan teknik *SSE*. *SSE* (*Sum of Square Error*) merupakan rumus yang digunakan untuk mengukur perbedaan antara data yang telah dilakukan sebelumnya [23].

3. HASIL DAN PEMBAHASAN

Penulis melakukan perhitungan dengan bantuan aplikasi *Rapid Miner*. Hasil dari indikator-indikator tersebut merupakan proses panjang dari perhitungan algoritma *k-means*, mulai dari menentukan jumlah *k*, menghitung jarak menggunakan *euclidean distance* kemudian menentukan titik *centroid* secara acak, dan perbaikan nilai *centroid*. Dengan menggunakan 3 *cluster*, *cluster* pertama yaitu *cluster_0* yang berjumlah 377 data dengan tingkat penggunaan Rendah, *cluster_1* berjumlah 34 data dengan tingkat penggunaan Sedang dan *cluster_2* berjumlah 129 data dengan tingkat penggunaan Tinggi dan dilanjutkan dengan pengukuran jarak menggunakan *device bouldin index performance vector* rata_rata dalam *centroid* yang didapat dengan nilai 512012471.162.

Kemudian hasil dari tahapan pengujian *Elbow* pada *Google Colaboration (Google Colab)* dengan bahasa pemrograman Python untuk menghitung hasil *Sum Of Square Error (SSE)* diperoleh $K=3$ dengan Nilai 276486734427.49817. K adalah *cluster_0* yaitu berjumlah 377 data memiliki tingkat penggunaan Rendah, *cluster_1* berjumlah 34 data dengan keterangan tingkat penggunaan Sedang dan *cluster_2* berjumlah 129 data dengan keterangan tingkat penggunaan Tinggi. Berdasarkan hasil pengujian tersebut dari hasil clustering *k-means* pada *Rapid Miner* dengan jumlah 3 dapat dikatakan *valid* atau sesuai dengan *cluster k-means* pada *python* menggunakan *elbow method*.

3.1 Tahapan CRISP-DM pada Rapidminer

a. Business Understanding (Pemahaman Bisnis)

Tahap pertama subjek pada penelitian ini adalah mengacu pada data penggunaan listrik rumah tangga dimana pada tujuan penelitian ini adalah untuk menghasilkan beberapa *cluster*. Pada tahap ini diperlukan pemahaman tentang pentingnya data penggunaan listrik rumah tangga, agar dapat digunakan untuk membuat kebijakan yang lebih efektif untuk menanggulangi pemborosan penggunaan listrik. Oleh karena itu, diperlukan strategi dengan *clustering* atau pengelompokan data penggunaan listrik rumah tangga.

b. Data Understanding (Pemahaman Data)

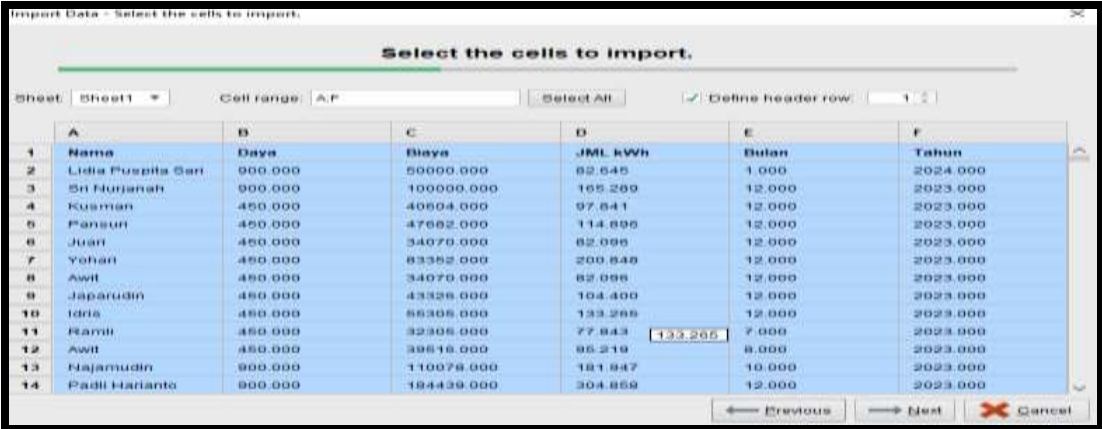
Peneliti melakukan pemahaman terhadap data penggunaan listrik rumah tangga. Pemahaman data mengacu pada klasterisasi besarnya penggunaan listrik rumah tangga data yang diambil yaitu data kec. Dempo Utara ada 540 record dengan atribut Nama, Daya, Biaya, Jml kWh, Bulan dan Tahun. Data yang diterima dalam bentuk *struck / token* lalu disalin ke *excel* data yang didapat semuanya terisi dan tidak ada yang kosong jadi peneliti perlu melakukan *missing data* dan atributnya semua terguna dalam penelitian ini atribut yang digunakan dapat dilihat pada gambar grafik dibawah ini.

c. Data Preparation (Persiapan Data)

Tahapan pengolahan data langsung diimplementasikan pada *RapidMiner* dengan tiga tahapan sebagai berikut.

1. Data Selection

Pada tahapan ini ada 6 atribut yang didapat dari data penggunaan listrik rumah tangga diantaranya Nama, Daya, Biaya, JML kWh, Bulan dan Tahun. Dilakukan seleksi menggunakan *excel* dengan cara *filter*. Setelah dilakukan *filter*, data yang peneliti gunakan untuk *clustering* berdasarkan pemilik usaha ada 6 atribut pada tahap *selection* data diantaranya Nama, Daya, Biaya, JML kWh, Bulan dan Tahun.

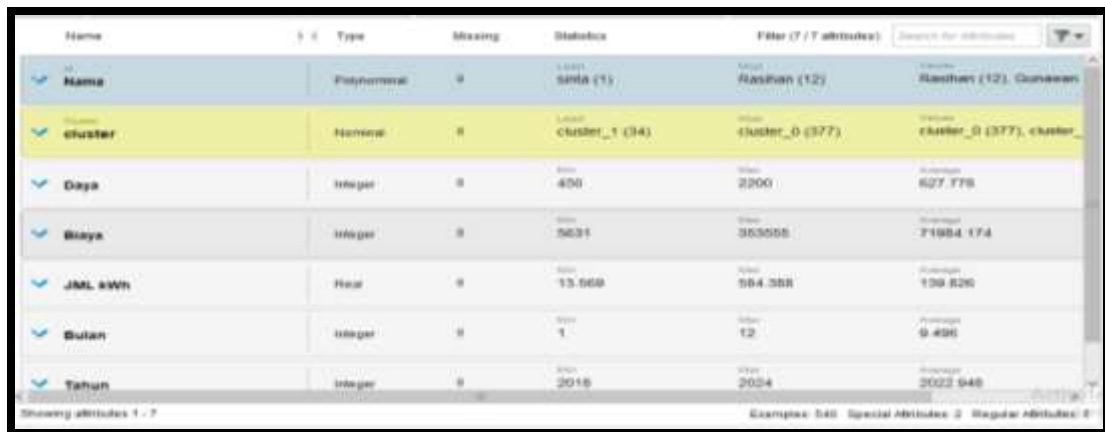


	A	B	C	D	E	F
	Nama	Daya	Biaya	JML kWh	Bulan	Tahun
1	Lidia Puspa Sari	900.000	50000.000	82.645	1.000	2024.000
2	Si Nurjanah	900.000	100000.000	165.269	12.000	2023.000
3	Kusman	450.000	40604.000	97.641	12.000	2023.000
4	Pansuri	450.000	47602.000	114.896	12.000	2023.000
5	Juan	450.000	34070.000	82.096	12.000	2023.000
6	Yohari	450.000	83352.000	200.848	12.000	2023.000
7	Awit	450.000	34070.000	82.096	12.000	2023.000
8	Japardudin	450.000	43326.000	104.400	12.000	2023.000
9	Idris	450.000	85305.000	133.265	12.000	2023.000
10	Ramli	450.000	32305.000	77.843	7.000	2023.000
11	Awit	450.000	39616.000	95.219	8.000	2023.000
12	Hajamudin	900.000	110076.000	181.847	10.000	2023.000
13	Padli Harianto	900.000	194439.000	304.859	12.000	2023.000

Gambar 2. Data Selection

2. Data Processing

Proses *processing* merupakan proses yang mencakup *cleaning* dan transformasi data. Pada tahap *data processing* ini data awal 540 data dan diproses tetap 540 data yang diolah pada tahap *data processing*, peneliti sudah memastikan tidak ada lagi *missing value* atau data yang kosong, jika atribut yang digunakan data *missing value* nya sudah 0, maka data dapat digunakan pada tahap selanjutnya, sehingga dapat dilihat pada gambar dibawah.



Name	Type	Mixing	Statistics	File (7 / 7 attributes)	Search for Attributes
Nama	Polynomial	0	Local: SMA (1)	Local: Rashedan (12)	Local: Rashedan (12), Consonant
cluster	Normal	0	Local: cluster_1 (34)	Local: cluster_0 (377)	Local: cluster_0 (377), cluster_1
Daya	Integer	0	Min: 400	Max: 2200	Average: 827.778
Biaya	Integer	0	Min: 5631	Max: 363555	Average: 71984.174
JML kWh	Real	0	Min: 13.568	Max: 584.388	Average: 139.826
Bulan	Integer	0	Min: 1	Max: 12	Average: 9.496
Tahun	Integer	0	Min: 2015	Max: 2024	Average: 2022.948

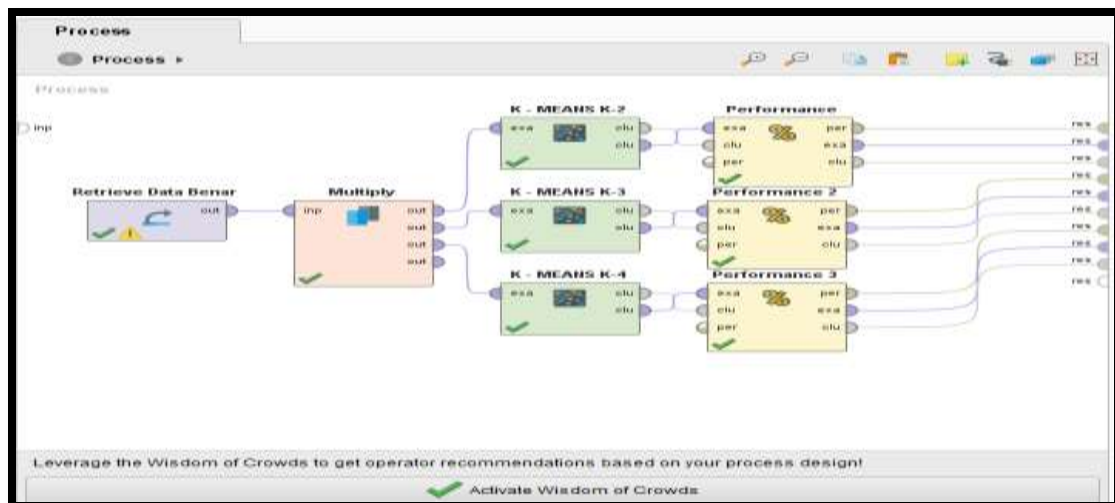
Gambar 3. Data Processing

3. Transformation

Pada fase transformasi data, data yang diproses menunjukan atribut yang dipilih dan diproses di *rapid miner*, data yang diolah harus berupa angka dan jika bukan angka maka harus diubah dengan proses *transormation* data. Dikarenakan data yang didapat oleh penulis sudah berbentuk angka semua, jadi penulis tidak melakukan transformasi data.

d. Modelling (Permodelan)

Pada tahap ini proses dimana model klasterisasi dengan menggunakan algoritma *k-means* dengan modelnya *clustering* operator ini mengambil objek dari port input dan mengirimkan salinannya ke *port output* dan hubungkan dengan *performance davies bouldin*. Setiap *port* yang terhubung membuat salinan yang independen (tidak terikat). Jadi ketika mengubah suatu salinan tidak berpengaruh pada salinan yang lainnya sehingga dapat dihubungkan dengan *performance* yaitu untuk mengetahui suatu model algoritma *k-means*. Sehingga menghasilkan suatu pola informasi yang dapat memudahkan pihak yang berkepentingan.



Gambar 4. Desain Modelling

1. Tentukan jumlah cluster

Untuk menentukan jumlah cluster dilakukan percobaan jumlah cluster 3. Percobaan dengan 3 cluster itu ada C0, C1 dan C2 dengan atribut Daya, Biaya, JML kWh, Bulan dan Tahun. Dengan pengukuran *performance vector* rata-rata dalam *centroid distance* yang didapat dengan nilai 512012471.162, kemudian rata-rata *centroid cluster_0* dengan nilai 351786318.910 rata-rata *centroid cluster_1* dengan nilai 1823535798.660, rata-rata *centroid cluster_2* dengan nilai 634597480.961 dan nilai *davies bouldin index* 0.498.

Attribute	cluster_0	cluster_1	cluster_2
Daya	548.806	900	786.822
Biaya	41756.419	231262.324	118343.915
JML kWh	91.887	382.252	216.030
Bulan	9.422	9.882	9.612
Tahun	2022.931	2023	2022.984

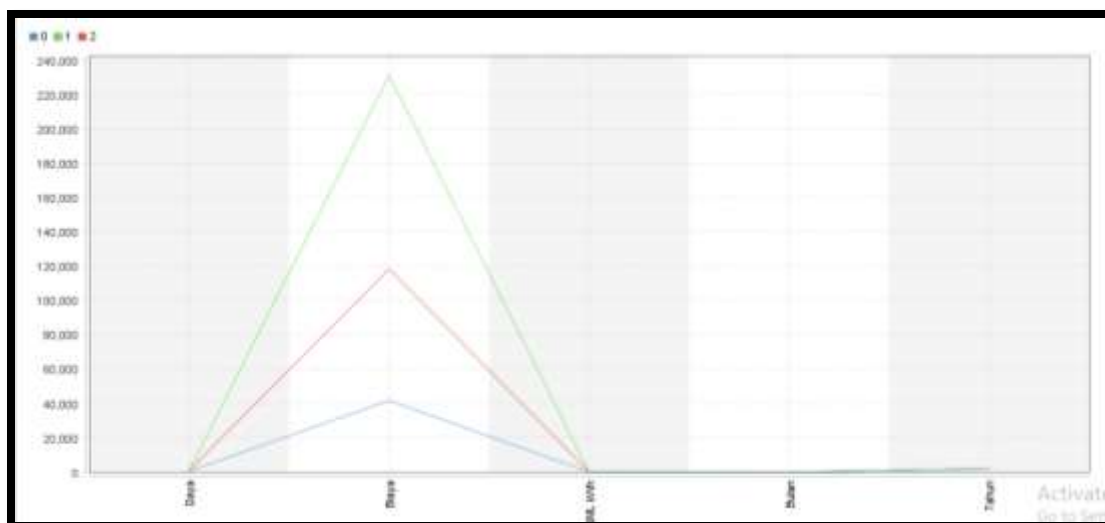
Gambar 5. Titik *Centroid*

- Menentukan titik *centroid* secara acak

Titik *centroid* secara acak ditentukan berdasarkan beberapa percobaan *cluster* di *RapidMiner*. Dari 3 percobaan tersebut dipilih titik *centroid* terkecil untuk mendapatkan hasil terbaik.

- Hitung jarak data ke *centroid*

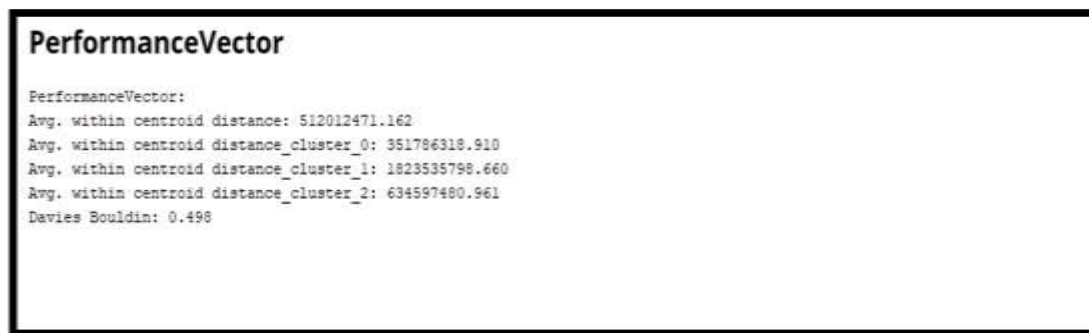
Setelah memasukkan dataset dan algoritma *K-Means* pada *RapidMiner* maka proses data dengan klik run untuk melihat jarak data ke *centroid*. Grafik jarak data ke *centroid* dapat dilihat pada Gambar 6.



Gambar 6. Grafik Titik *Centroid*

- Evaluation* (Evaluasi)

Berdasarkan pengujian yang telah dilakukan dengan menggunakan *davies bouldin index*, supaya menghasilkan nilai yang baik dengan melakukan sebanyak 3 *cluster* dengan keseluruhan data 540 *record performance vector* rata-rata *centroid distance* yang didapat dengan nilai 512012471.162, kemudian rata-rata *centroid cluster_0* dengan nilai 351786318.910 rata-rata *centroid cluster_1* dengan nilai 1823535798.660, rata-rata *centroid cluster_2* dengan nilai 634597480.961 dan nilai *davies bouldin index* 0.498, berikut hasil *davies bouldin index* algoritma tersebut.



Gambar 7. Performance Vector

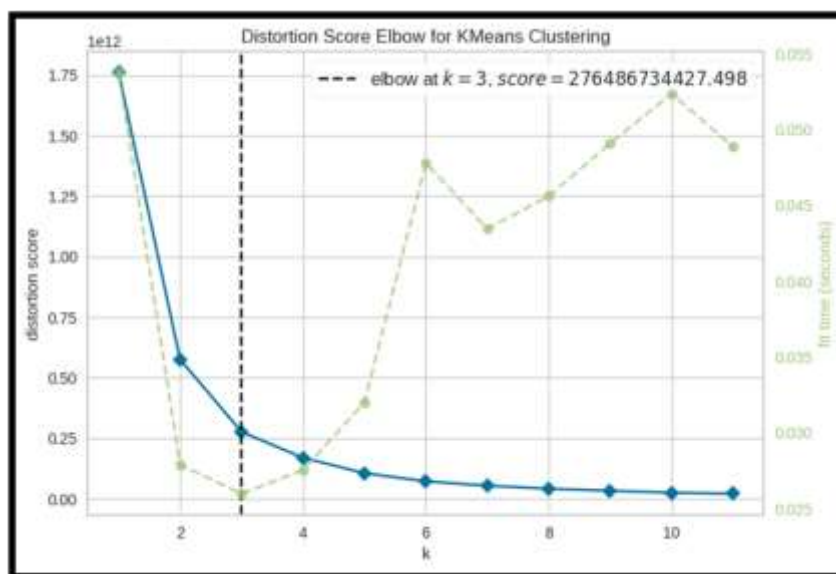
f. Deployment (Penyebaran)

Pada tahapan ini deployment ini ialah tahapan terakhir berupa pengetahuan atau informasi yaitu mengenai pola penggunaan listrik rumah tangga di Kec. Dempo Utara Pagar Alam dan klaster yang memiliki karakteristik penggunaan listrik yang sama beberapa tahun terakhir, sehingga dapat diketahui ada 3 cluster bahwa yang memiliki keterangan cluster yang tingkat pembayaran listrik Rendah ialah cluster_0 dengan total 377 data, cluster_1 memiliki keterangan cluster yang tingkat pembayaran listrik Tinggi dengan total 34 data dan cluster_2 memiliki keterangan tingkat pembayaran listrik Sedang dengan total 129 data. Dari proses clustering dapat dilihat tingkat penggunaan listrik dominan berada pada tingkat pembayaran listrik Rendah itu artinya sesuai dengan ekonomi yang ada di pagar alam karena mayoritas pendapatan masyarakatnya adalah dari petani dan UMKM. Sehingga dapatlah sebuah pola dari 3 cluster tersebut. Dibawah ini adalah tabel dari hasil cluster.

Tabel 1. Jumlah Cluster		
Cluster	Tingkat	Jumlah
Cluster_0	Rendah	377
Cluster_1	Tinggi	34
Cluster_2	Sedang	129

3.2 Pengujian metode elbow google colab

Langkah selanjutnya akan dilakukan pengujian menggunakan *Elbow Method*. Sebuah plot yang menunjukkan nilai *inertia* (sum squared distances of samples to their closest cluster center) untuk setiap nilai k dari 1 hingga 12. Hasil *Elbow* dapat dilihat pada gambar dibawah ini.



Gambar 8. Hasil Pengujian Elbow Method

Kemudian didapatkan hasil *cluster* yang berbentuk siku, dari gambar diatas data penggunaan listrik rumah tangga sebanyak 540 *record* maka didapatkan hasil *cluster* yang tepat berjumlah 3 *cluster* ($K=3$) dengan nilai *sum of square error* yaitu 276486734427.498, bahwasannya hasil tersebut adalah Jumlah *Cluster* yang optimal yaitu biasanya terletak di titik *centorid* terdekat atau siku terdekat pada plot *elbow*.

4. KESIMPULAN

Penelitian ini menghasilkan 3 pola *cluster* dengan *cluster_0* yaitu berjumlah 377 data memiliki tingkat penggunaan Rendah, *cluster_1* berjumlah 34 data dengan keterangan tingkat penggunaan Sedang dan *cluster_2* berjumlah 129 data dengan keterangan tingkat penggunaan Tinggi.

Kemudian hasil dari tahapan pengujian Elbow pada Google Colaboration (Google Colab) dengan bahasa pemrograman Python untuk menghitung hasil Sum Of Square Error (SSE) diperoleh $K=3$ dengan Nilai 276486734427.49817. K adalah *cluster_0* yaitu berjumlah 377 data memiliki tingkat penggunaan Rendah, *cluster_1* berjumlah 34 data dengan keterangan tingkat penggunaan Sedang dan *cluster_2* berjumlah 129 data dengan keterangan tingkat penggunaan Tinggi. Berdasarkan hasil pengujian tersebut dari hasil clustering k-means pada Rapid Miner dengan jumlah 3 dapat di katakan sesuai dengan cluster k-means pada python menggunakan elbow method.

Dari hasil tersebut Diharapkan data yang di cluster akan dapat bermanfaat bagi instansi dan mempermudah seseorang dalam membaca sebuah data, serta dapat membantu menyederhanakan data dengan menggambarkan kelompok-kelompok umum, memungkinkan analis atau pengambil keputusan untuk bekerja dengan data yang lebih terstruktur.

UCAPAN TERIMAKASIH

Terima kasih disampaikan kepada pihak-pihak yang telah mendukung terlaksananya penelitian ini.

DAFTAR PUSTAKA

- [1] K. Hidayah, "Membangun Bangsa pada Era Digital," *lan.go.id*, 2023. <https://lan.go.id/?p=12800>
- [2] D. Suyanto, *Data Mining untuk Klasifikasi dan Klasterisasi Data*, Pertama. INFORMATIKA Bandung, 2019.
- [3] Yuli Mardi, "Data Mining: Klasifikasi Menggunakan Algoritma C4.5 Data mining merupakan bagian dari tahapan proses Knowledge Discovery in Database (KDD). Jurnal Edik Informatika," *J. Edik Inform.*, vol. 2, no. 2, pp. 213–219, 2019.
- [4] B. D. Laraswati, "Kelebihan dan Kekurangan Jenis-Jenis Clustering dalam Data Science," *blog.algoritma*, 2023. <https://blog.algoritma.com/jenis-clustering/>
- [5] A. Rangkuti, M. Prodi, P. Matematika, S. Utara, and Y. Yahfizham, "Pengenalan Algoritma Pemrograman Dasar Dalam Konteks Pembelajaran Pemrograman Awal," *J. Mat. dan Ilmu Pengetah. Alam*, vol. 1, no. 4, pp. 2987–5315, 2023, [Online]. Available: <https://doi.org/10.59581/konstanta.v1i4.1714>
- [6] T. Amalina, D. Bima, A. Pramana, and B. N. Sari, "Metode K-Means Clustering Dalam Pengelompokan Penjualan Produk Frozen Food," *J. Ilm. Wahana Pendidik.*, vol. 8, no. 15, pp. 574–583, 2022, [Online]. Available: <https://doi.org/10.5281/zenodo.7052276>
- [7] F. Febriansyah and S. Muntari, "Penerapan Algoritma K-Means untuk Klasterisasi Penduduk Miskin pada Kota Pagar Alam," *JISKA (Jurnal Inform. Sunan Kalijaga)*, vol. 8, no. 1, pp. 66–77, 2023, doi: 10.14421/jiska.2023.8.1.66-77.
- [8] V. R. Prasetyo, H. Lazuardi, A. A. Mulyono, and C. Lauw, "Penerapan Aplikasi RapidMiner Untuk Prediksi Nilai Tukar Rupiah Terhadap US Dollar Dengan Metode Linear Regression," *J. Nas. Teknol. dan Sist. Inf.*, vol. 7, no. 1, pp. 8–17, 2021, doi: 10.25077/teknosi.v7i1.2021.8-17.
- [9] R. Gelar Guntara, "Pemanfaatan Google Colab Untuk Aplikasi Pendeteksian Masker Wajah Menggunakan Algoritma Deep Learning YOLOv7," *J. Teknol. Dan Sist. Inf. Bisnis*, vol. 5, no. 1, pp. 55–60, 2023, doi: 10.47233/jteksis.v5i1.750.
- [10] M. K. Rahmadhika and A. M. Thantawi, "Rancang Bangun Aplikasi Face Recognition Pada Pendekatan CRM Menggunakan Opencv Dan Algoritma Haarcascade," *IKRA-ITH Inform. J. Komput. dan Inform.*, vol. 5, no. 1, pp. 109–118, 2021, [Online]. Available: <https://journals.upi-yai.ac.id/index.php/ikraith-informatika/article/view/921>
- [11] R. F. Setiono, I. Setiani, "Hubungan Antara Besarnya Daya Listrik Terpasang Dengan Banyaknya Pemakaian Listrik Dalam Skala Rumah Tangga," *Stud. P., Elektro, T., Vokasi, S., Diponegoro, U., Pascasarjana, P., Diponegoro, U*, pp. 978–979, 2019.
- [12] A. Aziz et al., "Klasterisasi Data Tenaga Kerja Terbuka Menurut Provinsi Dengan Penggunaan Algoritma K - Means 1,2," vol. 10, no. 3, pp. 444–456, 2023.
- [13] S. Setya Wiwaha, "Analisis Peluang Penghematan," pp. 47–58, 2017.
- [14] Y. Yulia and N. Azwanti, "Penerapan Algoritma C4.5 Untuk Memprediksi Besarnya Penggunaan Listrik Rumah Tangga di Kota Batam," *J. RESTI (Rekayasa Sist. dan Teknol. Informasi)*, vol. 2, no. 2, pp. 584–590, 2018, doi: 10.29207/resti.v2i2.503.
- [15] I. K. Sukesu, "CRISP DM Sebagai Salah Satu Standard untuk Menghasilkan Data Driven Decision Making yang Berkualitas," *djkn.kemenkeu.go.id*, 2022. <https://www.djkn.kemenkeu.go.id/artikel/baca/15134/CRISP-DM-Sebagai-Salah-Satu-Standard-untuk-Menghasilkan-Data-Driven-Decision-Making-yang-Berkualitas.html>
- [16] R. Mubarak, M. Hanafi, and D. Sasongko, "Komparasi Performa Naive Bayes Gaussian dan K-NN Untuk Prediksi Kelulusan Mahasiswa dengan CRISP-DM," vol. 4, no. 6, pp. 2982–2991, 2024, doi: 10.30865/klik.v4i6.1924.
- [17] F. Yanti, B. N. Sari, and S. Defiyanti, "IMPLEMENTASI ALGORITMA LSTM PADA PERAMALAN STOK OBAT (STUDI KASUS: PUSKESMAS BEBER)," vol. 8, no. 4, pp. 6082–6089, 2024.
- [18] M. A. Sofyan, N. Rahaningsih, and R. D. Dana, "Deteksi Sms Spam Berbahasa Indonesia Menggunakan Algoritma Support Vector Machine," *J. Mhs. Tek. Inform.*, vol. 8, no. 3, pp. 3071–3079, 2024.
- [19] R. H. Rachmat Hidayat, "Perbandingan Metode Naive Bayes Dan Decision Tree C4.5 untuk Analisis Sentimen Produk Es Teh Indonesia di Media Sosial Twitter," *J. SISKOM-KB (Sistem Komput. dan Kecerdasan Buatan)*, vol. 7, no. 2, pp. 88–98, 2024, doi:

10.47970/siskom-kb.v7i2.607.

- [20] D. Rafly, R. Insani, and A. Dzulkarnain, “Penerapan FP-Growth untuk Menentukan Rekomendasi Produk pada Kogu Coffee Shop Malang,” *INTEGER J. Inf. Technol.*, vol. 9, no. 1, pp. 27–36, 2024.
- [21] D. A. Dzulhijjah, M. B. Herlambang, M. Haifan, P. Studi, P. Insinyur, and T. Selatan, “Implementasi Framework CRISP-DM untuk Proses Data Mining Aplikasi Credit Scoring PT.XYZ,” *Pros. Semin. Nas. Teknol. Inf. dan Bisnis*, vol. 1, no. 1, pp. 238–251, 2023.
- [22] N. Putu, E. Merliana, and A. J. Santoso, “PROSIDING SEMINAR NASIONAL MULTI DISIPLIN ILMU & CALL FOR PAPERS UNISBANK (SENDI_U) Kajian Multi Disiplin Ilmu untuk Mewujudkan Poros Maritim dalam Pembangunan Ekonomi Berbasis Kesejahteraan Rakyat ANALISA.PENENTUAN JUMLAH CLUSTER TERBAIK PADA METODE K-ME,” pp. 978–979, 2019, [Online]. Available: <https://www.unisbank.ac.id/ojs/index.php/sendu/article/view/3333>
- [23] N. A. Maori, “METODE ELBOW DALAM OPTIMASI JUMLAH CLUSTER PADA K-MEANS CLUSTERING,” vol. 14, no. 2, pp. 277–287, 2023.