

PENERAPAN DATA MINING UNTUK MENENTUKAN PENERIMA BANTUAN BLT MENGGUNAKAN METODE CLUSTERING K-MEANS PADA DESA PAMULIHAN

Eddiwan Dwiguna, Agus Bahtiar

Manajemen Informatika, STMIK IKMI Cirebon
Jalan Perjuangan No 10 B Majasem Kota Cirebon, Indonesia
eddiwandwiguna@gmail.com

ABSTRAK

Penggunaan metode Clustering *K-means* dalam Data Mining untuk menentukan penerima Bantuan Langsung Tunai (BLT) terus berkembang seiring dengan kemajuan ketersediaan data, teknologi, dan kesadaran akan efisiensi dan akurasi dalam pengelolaan program bantuan sosial. Proses manual dalam menentukan penerima BLT di Desa Pamulihan dinilai tidak transparan, kurang efisien, dan dapat menimbulkan ketidakpuasan di kalangan penduduk. Keberlanjutan program BLT juga terkendala oleh ketidakjelasan kriteria penentuan penerima, keterbatasan sumber daya finansial, dan potensi risiko kecurangan. Penelitian ini difokuskan pada penerapan *K-means* Clustering dengan mempertimbangkan penentuan kriteria penerima, optimalitas pengelompokan, evaluasi model, etika dan keadilan proses, serta pengelolaan data dan keamanan privasi. Tujuan utama dari penelitian ini adalah membangun model Data Mining menggunakan metode Clustering *K-means* untuk mengidentifikasi calon penerima bantuan BLT di Desa Pamulihan. Hasil penelitian di peroleh pengelompokkan terbaik adalah K9 dengan nilai *Davies bouldin index* (DBI) sebesar 0.745, nilai tersebut paling optimal karena mendekati angka 0 serta menghasilkan cluster data dengan jumlah data terbanyak adalah cluster 0, yang terdiri dari 50 data yang terdiri dari 14 data penerima, 15 data bukan penerima, dan 21 data dipertimbangkan jumlah data paling sedikit adalah cluster 2, yang terdiri dari 16 data yaitu penerima 15 dan bukan penerima 1.

Kata kunci : Data Mining, Clustering *K-means*, Bantuan Langsung Tunai, Bantuan Sosial Desa Pamulihan

1. PENDAHULUAN

Data Mining, khususnya penggunaan metode Clustering *K-means*, dalam menentukan penerima bantuan Bantuan Langsung Tunai (BLT) terus mengalami perkembangan terkini yang menarik. Perkembangan ini berakar dari peningkatan ketersediaan data, teknologi, dan kesadaran akan kebutuhan untuk mengelola program bantuan sosial dengan lebih efisien dan akurat. Berikut adalah deskripsi tentang perkembangan terkini dalam penggunaan Data Mining untuk menentukan penerima bantuan BLT dengan metode Clustering *K-means*.

Penelitian ini meliputi penerapan Data Mining untuk mengidentifikasi penerima Bantuan Langsung Tunai (BLT) dengan menggunakan metode clustering *K-means* yang melibatkan beberapa konsep dan definisi yang relevan serta permasalahan penelitian yang harus diuraikan. Penelitian ini fokus pada penerapan *K-means* clustering. Permasalahannya meliputi penentuan kriteria penerima, optimalitas pengelompokan, evaluasi model, etika dan keadilan proses, serta pengelolaan data dan keamanan privasi. Tujuannya adalah mencari solusi yang efisien dan adil dalam mengidentifikasi penerima bantuan BLT dengan menggunakan pendekatan data mining.

Data yang dapat mendukung urgensi penelitian mengenai penerapan Data Mining dengan metode Clustering *K-means* untuk menetapkan calon penerima Bantuan Langsung Tunai (BLT) di Desa Pamulihan dapat bersumber dari berbagai sumber, baik berupa informasi lapangan yang diperoleh secara langsung di

lokasi penelitian maupun dari literatur yang telah dijelaskan dalam jurnal atau penelitian sebelumnya.

Oleh karena itu diusulkan penelitian dengan judul Penerapan Data Mining Untuk Menentukan Penerima Bantuan BLT Menggunakan Metode Clustering K – MEANS Pada Desa Pamulihan. Adapun yang menjadi alasan dilakukannya penelitian ini dengan judul tersebut karena judul sudah dipilih setelah pertimbangan yang matang karena mencerminkan inti dari penelitian ini. Pertama, judul ini mencakup fokus utama penelitian, yaitu pemanfaatan teknik Data Mining dengan metode Clustering *K-means* untuk proses seleksi penerima Bantuan Langsung Tunai (BLT) di Desa Pamulihan. Kedua, judul ini dengan jelas menggambarkan solusi yang diusulkan, yaitu penerapan metode Clustering *K-means* untuk meningkatkan tingkat akurasi dan efisiensi dalam penentuan penerima BLT. Ketiga, judul ini mencantumkan lokasi penelitian, yaitu Desa Pamulihan, memberikan konteks yang spesifik dan relevan yang menunjukkan permasalahan yang dihadapi di tingkat lokal. Oleh karena itu, judul ini secara komprehensif mencerminkan tujuan dan metode penelitian serta memberikan informasi yang memadai kepada pembaca mengenai latar belakang penelitian ini.

2. TINJAUAN PUSTAKA

2.1. Penelitian Terdahulu

Penelitian yang dilakukan Yosep Filki yang berjudul Algoritma *K-means* Clustering dalam

Memprediksi Penerima Bantuan Langsung Tunai (BLT) Dana Desa, untuk menentukan penerima BLT-DD yang tepat sasaran di Nagari Taluk Kecamatan Lintau Buo Kabupaten Tanah Datar menggunakan metode *K-means* Clustering. Hasil penelitian ini menunjukkan bahwa penggunaan Metode *K-means* Clustering efektif dalam memprediksi penerima Bantuan Langsung Tunai Dana Desa (BLT-DD) di Nagari Taluk Kecamatan Lintau Buo Kabupaten Tanah Datar. Penggunaan metode *K-means* Clustering dalam pengolahan data usulan penerima BLT-DD di Nagari Taluk tahun 2022 menghasilkan tiga cluster, yaitu layak, dipertimbangkan, dan tidak layak. Dari 25 data sampel, 11 data termasuk dalam cluster 1 dengan status layak, 5 data termasuk dalam cluster 2 dengan status dipertimbangkan, dan 9 data termasuk dalam cluster 3 dengan status tidak layak [1].

Penelitian yang dilakukan oleh Sarmila Sari dan Joy Nashar Utamajaya berjudul Sistem Pendukung Keputusan Penerima Bantuan Langsung Tunai Dana Desa Menggunakan Metode Algoritma K-means Clustering, bertujuan untuk menetapkan prioritas penerima bantuan utama dari sejumlah warga yang mengajukan permohonan bantuan. Algoritma K-means digunakan untuk ini. Dengan menerapkan metode K-means clustering menggunakan aplikasi RapidMiner Studio pada 336 data penduduk yang memenuhi kriteria sebagai calon penerima BLT-DD, ditemukan tiga cluster. Cluster 1, dengan 87 data, tergolong dalam kategori penerima bantuan yang tepat sasaran. Cluster 2, dengan 193 data, tergolong dalam kategori penerima bantuan yang masih perlu dipertimbangkan. Sementara itu, Cluster 3, dengan 56 data, tergolong dalam kategori penerima bantuan yang tidak tepat sasaran. [2].

Penelitian yang dilakukan Aviv Fitria Yulia, Ratnasari, Panji Bintoro yang berjudul *K-means* Clustering Dalam Menentukan Kelayakan penerima bantuan untuk warga miskin studi kasus pekon Sukoharjo III, penelitian tersebut menghadapi masalah serupa dengan penelitian ini, yaitu identifikasi calon penerima bantuan sosial. Penelitian tersebut menggunakan metode clustering K-means untuk mengelompokkan calon penerima bantuan, yang menghasilkan peningkatan yang signifikan dalam tingkat akurasi identifikasi penerima bantuan. Kesimpulan yang dapat diambil dari penelitian sebelumnya adalah bahwa metode clustering K-means telah terbukti efektif dalam meningkatkan akurasi dan efisiensi proses identifikasi penerima bantuan sosial. Temuan ini memberikan dukungan kuat terhadap penelitian yang sedang berlangsung, menegaskan bahwa penerapan data mining dan metode clustering K-means memiliki potensi besar untuk memberikan solusi efisien terhadap tantangan identifikasi penerima bantuan BLT di Desa Pamulihan, sesuai dengan hasil positif yang tercatat dalam penelitian terdahulu [3].

Penelitian yang dilakukan Arisman Waruwu, Milfa Yetri, Feri Setiawan yang berjudul Implementasi Data Mining Dalam Mengelompokkan Data penduduk

Kurang Mampu Menggunakan Metode *K-means* Clustering, untuk melakukan analisis dan pengelompokan data penduduk kurang mampu di Desa Hiliwale II menggunakan metode *k-means* clustering, agar data penduduk kurang mampu tersebut dapat dikelompokkan berdasarkan karakteristik yang sama. Data penduduk kurang mampu di Desa Hiliwale II berhasil dikelompokkan menjadi 2 cluster menggunakan metode *k-means* clustering. Cluster 1 memiliki jumlah anggota sebanyak 15 orang. Cluster ini terdiri dari penduduk yang memiliki rata-rata pendapatan, status kepemilikan rumah, dan jumlah tanggungan di atas rata-rata. Cluster 2 memiliki jumlah anggota sebanyak 19 orang. Cluster ini terdiri dari penduduk yang memiliki rata-rata pendapatan, status kepemilikan rumah, dan jumlah tanggungan di bawah rata-rata [4].

Penelitian yang dilakukan Fadia Azzahra, Ahmad Zakir, Dedy Irwan yang berjudul Pemanfaatan Algoritma *K-means* Untuk Menentukan Penerimaan Bantuan Langsung Tunai Pada Kecamatan Medan Area, tujuan dari penelitian ini yaitu untuk membangun suatu sistem data mining penentuan calon penerima bantuan tunai langsung berdasarkan kriteria yang sudah ditentukan oleh kantor camat medan area. Hasil yang dicapai adalah algoritma *K-means* terbukti efektif digunakan untuk mengelompokkan data bantuan langsung tunai sehingga diketahui masyarakat yang membutuhkan dan tidak membutuhkan bantuan tersebut [5].

2.2. Data Mining

Data mining merupakan proses ekstraksi pola atau pengetahuan yang berguna dari suatu set data yang besar. Proses analisis data yang dikenal sebagai data mining bertujuan untuk mengekstraksi pola atau informasi berharga dari kumpulan data besar atau kompleks. Proses ini memanfaatkan teknik seperti matematika, statistika, dan kecerdasan buatan untuk mengidentifikasi pola yang tidak terlihat secara langsung, memprediksi hasil masa depan, membagi data menjadi kelompok, menemukan anomali, dan mengoptimalkan proses bisnis. Data mining, yang memanfaatkan berbagai algoritma seperti regresi, klasifikasi, pengelompokan, dan asosiasi, memungkinkan organisasi untuk membuat pilihan yang lebih baik dan mengoptimalkan penggunaan sumber daya mereka [6].

2.3. Clustering

Dalam analisis data, clustering adalah cara untuk mengelompokkan objek atau data ke dalam kelompok atau klaster berdasarkan hal-hal yang serupa. Kelompok yang memiliki banyak kesamaan internal sementara sedikit perbedaan antar mereka adalah tujuan utama clustering. Dengan metode ini, data dapat dikelompokkan tanpa membutuhkan informasi sebelumnya tentang kategori atau kelompok yang diharapkan. Algoritma clustering digunakan untuk mencoba menemukan struktur yang tersembunyi

dalam data dan membagi objek menjadi kelompok yang berbeda berdasarkan sejumlah fitur atau atribut yang diukur. Segmentasi pelanggan, pengelompokan dokumen dalam analisis teks, dan klasifikasi genetik dalam studi bioinformatika adalah beberapa contoh penggunaan clustering [7].

2.4. Algoritma K-means

Algoritma clustering K-means digunakan dalam analisis data dan pembelajaran mesin dengan tujuan mengelompokkan data ke dalam k kelompok atau klaster berdasarkan kesamaan karakteristik atau fitur. Proses dimulai dengan menginisialisasi titik pusat klaster secara acak, lalu data dikelompokkan ke dalam klaster yang memiliki pusat terdekat. Titik pusat klaster dihitung ulang menggunakan data yang telah dikelompokkan, dan proses ini diulang sampai titik pusat stabil. Walaupun K-means memiliki kelebihan dalam skalabilitas dan efisiensi, algoritma ini sensitif terhadap inisialisasi awal, yang dapat memengaruhi hasil clustering tergantung pada titik awal yang dipilih. Penting untuk menghitung jumlah klaster sebelumnya sebelum menggunakan K-means, dan interpretasi hasil clustering perlu dilakukan dengan hati-hati. Beberapa contoh penggunaan algoritma ini melibatkan segmentasi pelanggan, analisis teks, dan pengelompokan data dalam berbagai bidang ilmu [8].

2.5. Kriteria Penerima Bantuan BLT

Kriteria untuk penerima Bantuan Langsung Tunai (BLT) bergantung pada berbagai hal, seperti keadaan ekonomi, status sosial, dan kondisi individu. BLT biasanya diberikan kepada masyarakat dengan pendapatan rendah atau di bawah garis kemiskinan. Faktor-faktor seperti status keluarga, jumlah tanggungan, dan keberlanjutan dalam kasus usia lanjut atau disabilitas juga dapat menjadi pertimbangan. Selain itu, keadaan khusus, seperti konsekuensi langsung dari bencana alam atau krisis ekonomi, dapat menentukan penerima BLT. Untuk memastikan bahwa penerima benar-benar memenuhi persyaratan yang telah ditetapkan, pendaftaran dan verifikasi sering menjadi bagian dari proses. Ada kemungkinan bahwa program BLT akan difokuskan pada area tertentu yang dianggap membutuhkan bantuan tambahan. Selain itu, pemerintah dapat mempertimbangkan perubahan dalam kondisi kehidupan seseorang sebagai faktor penentu kelayakan penerima, seperti kehilangan pekerjaan atau kondisi kesehatan yang buruk [9].

2.6. Implementasi K-means dalam Penentuan Penerima Bantuan

Penjelasan tentang bagaimana metode K-means diimplementasikan dalam penelitian ini. Ini mencakup langkah-langkah pemrosesan data, pengaturan parameter K, dan analisis hasil untuk menentukan kelompok penerima bantuan BLT. Dalam penentuan penerima bantuan, algoritma clustering digunakan untuk mengelompokkan orang atau keluarga berdasarkan fitur tertentu. Selama proses ini, semua

data yang relevan diukur, termasuk tingkat pendapatan, jumlah tanggungan, dan status pekerjaan. Kemudian, algoritma K-means membagi data tersebut menjadi kelompok, atau klaster, yang memiliki kesamaan karakteristik. Hasilnya membantu pemerintah atau lembaga bantuan sosial menemukan kelompok masyarakat yang membutuhkan bantuan ekonomi berdasarkan tren atau pola tertentu, seperti tingkat pendapatan rendah atau jumlah tanggungan yang tinggi. Meskipun K-means memberi wawasan awal, interpretasi hasil dan kebijakan bantuan harus mempertimbangkan faktor sosial, ekonomi, dan etika. Sebelum membuat keputusan akhir tentang siapa yang akan menerima bantuan, hasil clustering juga perlu diverifikasi dan divalidasi dengan data tambahan [10].

2.7. Davis Bouldin Index

Davies Bouldin Index (DBI) adalah metode validasi cluster yang dirancang oleh D.L. Davies. DBI merupakan fungsi perbandingan dari jumlah distribusi di dalam setiap cluster terhadap pemisahan antar cluster. Penggunaan DBI sebagai alat pengukuran bertujuan untuk memaksimalkan jarak antar-cluster. Dalam konteks penelitian ini, DBI diterapkan sebagai metode validasi data pada masing-masing cluster untuk memastikan kualitas pengelompokan yang optimal. Dengan memanfaatkan Davies Bouldin Index, sebuah skema pengelompokan dianggap optimal jika memiliki nilai Davies Bouldin Index yang minimal.[11].

3. METODE PENELITIAN

3.1. Sumber Data

Sumber Data pada penelitian ini yaitu Data primer, Data primer adalah jenis informasi yang dikumpulkan secara langsung dari sumbernya tanpa melalui interpretasi atau pengolahan sebelumnya. Ini termasuk data yang dikumpulkan melalui metode seperti wawancara, survei, observasi langsung, dan eksperimen. Jenis data ini unik karena dikumpulkan khusus untuk tujuan penelitian atau analisis dan dapat memberikan wawasan langsung tentang subjek atau konteks yang sedang diteliti.

Data primer pada penelitian ini bersumber dari database yang diperoleh dari Desa Pamulihan dengan nama data data penghasilan aset penduduk desa Pamulihan.

3.2. Teknik Pengumpulan Data

Untuk mendapatkan data atau informasi dari berbagai sumber, teknik pengumpulan data Memilih teknik pengumpulan yang tepat, untuk memastikan bahwa data yang dikumpulkan relevan, akurat, dan dapat diandalkan. Berikut ini adalah beberapa metode yang saya gunakan untuk mengumpulkan data.

1. Wawancara

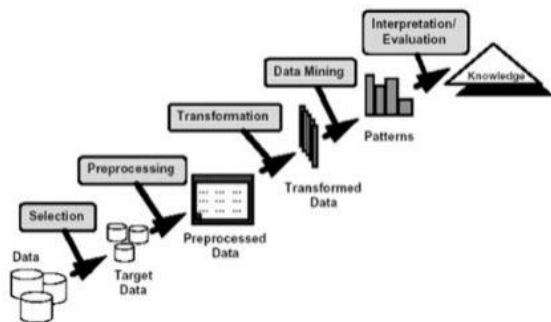
Untuk mendapatkan informasi yang diinginkan, metode wawancara melibatkan interaksi langsung antara peneliti atau pewawancara dan

responden. Wawancara dapat dilakukan dalam berbagai bentuk dan pendekatan

2. Observasi

Observasi adalah metode pengumpulan data yang melibatkan pengamatan langsung terhadap objek, kejadian, atau perilaku untuk mendapatkan informasi yang diperlukan. Observasi dapat dilakukan dalam berbagai konteks, seperti penelitian ilmiah, penelitian sosial, atau pemantauan aktivitas tertentu.

3.3. Tahapan Perancangan



Gambar 1. Proses KDD

Pada gambar 1, terlihat tahapan *Knowledge Discovery in Database (KDD)* yang menjelaskan beberapa langkah yang diambil untuk menghasilkan klasifikasi dari data mining menggunakan metode *K-means Clustering*.

3.3.1. Data Selection

Data selection, juga dikenal sebagai pemilihan data, adalah proses memilih subset atau bagian tertentu dari data yang akan digunakan untuk analisis lebih lanjut. Ini adalah tahap awal dalam proses penemuan pengetahuan dalam database (KDD), atau analisis data secara keseluruhan. Kualitas dan relevansi hasil dapat sangat dipengaruhi oleh pemilihan data yang baik.

3.3.2. Data Preprocessing/Cleaning Data

Data Cleaning adalah proses analisis data yang bertujuan untuk menemukan, memperbaiki, dan menghilangkan kesalahan atau inkonsistensi dari kumpulan data. Tujuan utama pembersihan data adalah meningkatkan kualitas data sehingga analisis yang dilakukan dapat menghasilkan hasil yang lebih akurat dan dapat diandalkan. Proses ini melibatkan tindakan untuk mengatasi berbagai masalah yang dapat muncul dengan dataset.

3.3.3. Data Transformation

adalah proses mengubah atau mengolah data dari bentuk atau struktur aslinya menjadi bentuk yang lebih sesuai atau bermanfaat untuk digunakan untuk analisis atau keperluan pengolahan data lainnya. Tujuan transformasi data adalah untuk meningkatkan kualitas data, membuatnya lebih mudah dipahami, dan

membuatnya lebih mudah untuk mengekstrak pengetahuan yang bermanfaat.

3.3.4. Data Mining

Data mining merupakan proses ekstraksi pola atau informasi penting dari kumpulan data yang berskala besar. Tujuan utama dari proses ini adalah untuk menemukan hubungan yang tidak terlihat atau mendapatkan pengetahuan baru dari data yang ada. Metode data mining melibatkan penerapan teknik matematika, statistik, dan kecerdasan buatan untuk menganalisis serta menginterpretasi pola yang mungkin sulit atau bahkan tidak dapat terlihat melalui analisis manusia biasa.

3.3.5. Data interpretation (evaluation)

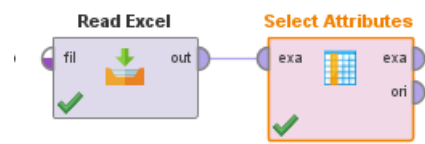
Interpretasi data adalah proses menganalisis dan memberikan makna atau pemahaman terhadap informasi yang terkandung dalam data. Proses ini melibatkan ekstraksi pola, signifikansi, dan pengetahuan yang dapat membantu dalam pengambilan keputusan atau memahami lebih lanjut tentang fenomena atau masalah. Interpretasi data seringkali merupakan bagian penting dari siklus analisis data dan dapat mencakup berbagai aspek.

4. HASIL DAN PEMBAHASAN

4.1. Hasil Nilai K Optimal Pada Data Penerima Bantuan BLT

4.1.1. Data Selection

Untuk memasukkan dataset Penerima bantuan BLT xlsx ke dalam repository Rapidminer, perlu dilakukan proses import menggunakan operator read excel. Operator ini dapat ditambahkan ke workflow Rapidminer seperti yang ditunjukkan pada gambar 2 read excel dan selection data berikut.



Gambar 2. Read Excel dan Selection Data

Hasil Read Excel dari data Penerima (BLT), yang sudah dibaca, menunjukkan jumlah 280 data dengan 6 atribut sebelum dilakukan selection data.

Format your columns.

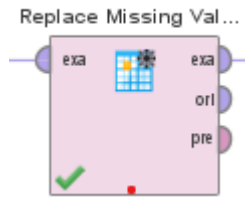
TANGGUNGAN	0 -	PEKERJAAN	0 -	GAJI	0 -	STATUS
integer		polymemo		integer		polymemo
1	1	buruh	2	ponore		
2	1	petani	2	bukan penerima		
3	1	petani	5	bukan penerima		
4	4	petani	5	diperibadikan		
5	5	petani	5	ponore		
6	6	petani	2	ponore		
7	7	petani	2	ponore		
8	8	petani	2	ponore		
9	9	petani	2	ponore		
10	10	buruh	3	diperibadikan		
11	11	petani	3	diperibadikan		
12	12	petani	5	bukan penerima		
13	13	petani	4	bukan penerima		
14	14	petani	3	diperibadikan		
15	15	ad-cara	3	ponore		
16	16	petani	5	bukan penerima		
17	17	petani	2	ponore		
18	18	buruh	2	ponore		
19	19	ad-cara	4	bukan penerima		

Gambar 3. Setelah Proses Slection Data

Sebelum dilakukan proses selection data terdapat 6 atribut yaitu NO, NAMA, TANGGUNGAN, PEKERJAAN, GAJI, STATUS dan setelah dilakukan proses selection data menjadi 4 atribut yaitu TANGGUNGAN, PEKERJAAN, GAJI serta STATUS seperti yang terlihat pada gambar 3 setelah proses selection data di atas.

4.1.2. Data Preprocessing

Setelah data Penerima Bantuan BLT diimpor, langkah selanjutnya adalah membersihkan data tersebut.



Gambar 4. Proses *Preprocessing*

Pada gambar 4, terlihat operator pada Rapid Miner untuk melakukan pembersihan data.

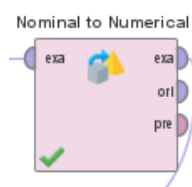
Name	Type	Missing	Statistics
STATUS	Nominal	0	count: 280 (100%)
cluster	Nominal	0	count: 280 (100%)
PEKERJAAN = buruh	Nominal	0	count: 280 (100%)
PEKERJAAN = petani	Nominal	0	count: 280 (100%)
PEKERJAAN = guru	Nominal	0	count: 280 (100%)
PEKERJAAN = pedagang	Nominal	0	count: 280 (100%)
PEKERJAAN = wiraswasta	Nominal	0	count: 280 (100%)
PEKERJAAN = PNS	Nominal	0	count: 280 (100%)
PEKERJAAN = IRT	Nominal	0	count: 280 (100%)

Gambar 5. Hasil *Preprocessing*

Berdasarkan yang didapat dari proses *preprocessing*, seperti terlihat pada Gambar 5 hasil *preprocessing* di atas, dapat disimpulkan bahwa tidak ada atribut yang memiliki nilai yang *missing* sehingga datanya tetap 280 data.

4.1.3. Data Transformation

Setelah dilakukan preprocessing dilanjutkan proses transformasi data pada tahap ini dilakukan transformasi data sebanyak 2 kali, yang pertama adalah seperti yang terlihat pada gambar 6 proses transformasi data nominal to numerical berikut.

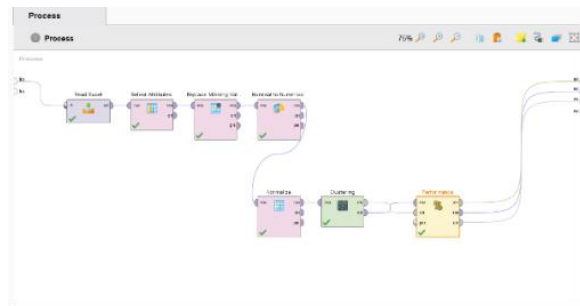


Gambar 6. Proses *Transformation* Data Nominal to Numerical

Transformasi data nominal to numeric dapat dilakukan untuk mengubah data kategorikal menjadi data numerik. Transformasi ini dapat meningkatkan akurasi model data mining, performa algoritma pembelajaran mesin, dan interpretabilitas model.

4.1.4. Data Mining

Untuk membuat cluster *K-means*, langkah pertama yang harus dilakukan adalah menentukan jumlah cluster. Jumlah cluster atau K yang digunakan adalah jumlah cluster optimal, yaitu jumlah cluster yang dapat mengelompokkan data Penerima Bantuan BLT dengan baik. Berikut adalah desain proses untuk membuat cluster *K-means* data Penerima Bantuan BLT menggunakan tools Rapid Miner versi 10.1.



Gambar 7. Model data mining

Proses pengolahan data mining dengan model *decision tree* dimulai dengan pemilihan data yang relevan. Data yang dipilih kemudian diproses melalui tahapan seleksi atribut, *preprocessing*, dan transformasi data. Tahapan-tahapan ini saling berhubungan dan saling mendukung satu sama lain.

4.1.5. Interpretation/Evaluation

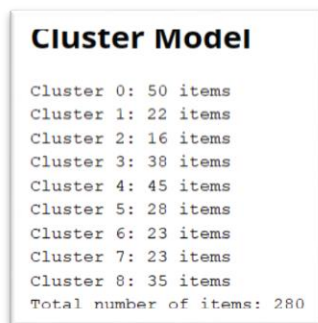
Berikut merupakan tabel hasil percobaan sampai dengan 10 cluster.

Tabel 1. Hasil Percobaan Nilai K

CLUSTER	DBI
2	0.945
3	0.840
4	0.865
5	0.899
6	0.822
7	0.832
8	0.842
9	0.745
10	0.753

Berdasarkan hasil DBI yang diperoleh dalam tabel 1, hasil percobaan nilai K di atas dapat disimpulkan hasil evaluasi dari clustering pada tabel diatas bahwa pengelompokan terbaik adalah cluster 9 dengan *Davies bouldin index* (DBI) sebesar 0.745, nilai tersebut paling optimal karena mendekati angka 0. Penelitian ini mengelompokkan data Penerima Bantuan BLT dengan tingkat kesamaan karakteristik yang mirip antara satu sama lain. Berikut hasil cluster

model yang bisa dilihat pada gambar 8 hasil cluster model pada halaman berikutnya.

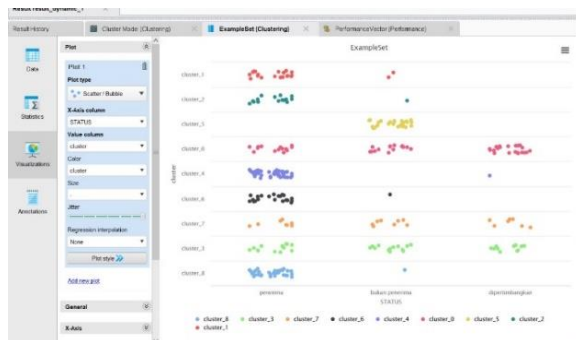


Gambar 8. Hasil Cluster Model

Hasil cluster model data Penerima Bantuan BLT terdiri dari 280 data yang terbagi menjadi 9 cluster yaitu cluster 0, 50 data, cluster 1, 22 data, cluster 2, 16 data, cluster 3, 38 data, cluster 4, 45 data, cluster 5, 28 data, cluster 6, 23 data, cluster 7, 23 data, dan cluster 8, 35 data.

4.2. Hasil Cluster Penerima Bantuan BLT Desa Pamulihan

Berikut adalah hasil cluster dari data Penerima Bantuan BLT Desa Pamulihan yang bisa dilihat pada gambar 9 hasil cluster penerima bantuan BLT berikut.



Gambar 9. Hasil Cluster Penerima Bantuan BLT

Terdapat 9 cluster yang dihasilkan dari proses clustering. Masing-masing cluster memiliki karakteristiknya masing-masing, yang dapat dilihat dari jumlah data dalam cluster tersebut, serta status data dalam cluster tersebut.

1. Cluster 0 adalah cluster dengan jumlah data yaitu 50 data. Cluster ini terdiri dari 14 data penerima, 15 data bukan penerima, dan 21 data dipertimbangkan.
2. Cluster 1 adalah cluster dengan jumlah 22 data. Cluster ini terdiri dari 20 data penerima dan 2 data bukan penerima.
3. Cluster 2 adalah cluster dengan jumlah yaitu 16 data. Cluster ini terdiri dari 15 data penerima dan 1 data bukan penerima.
4. Cluster 3 adalah cluster dengan jumlah yaitu 38 data. Cluster ini terdiri dari 13 data penerima, 13

data bukan penerima, dan 12 data dipertimbangkan.

5. Cluster 4 adalah cluster dengan jumlah data yaitu 45 data. Cluster ini terdiri dari 44 data penerima dan 1 data dipertimbangkan.
6. Cluster 5 adalah cluster dengan jumlah yaitu 28 data. Cluster ini terdiri dari 28 data bukan penerima.
7. Cluster 6 adalah cluster dengan jumlah data yaitu 23 data. Cluster ini terdiri dari 22 data penerima dan 1 data bukan penerima.
8. Cluster 7 adalah cluster dengan jumlah data yaitu 23 data. Cluster ini terdiri dari 8 data penerima, 7 data dipertimbangkan, dan 8 data bukan penerima.
9. Cluster 8 adalah cluster dengan jumlah data yaitu 35 data. Cluster ini terdiri dari 34 data penerima dan 1 data bukan penerima.

Berdasarkan penjelasan di atas, dapat disimpulkan bahwa hasil cluster menunjukkan bahwa terdapat 9 cluster data yang berbeda berdasarkan status datanya. Kelompok data dengan jumlah data terbanyak adalah cluster 0, yang terdiri dari 50 data. Kelompok data ini terdiri dari 14 data penerima, 15 data bukan penerima, dan 21 data dipertimbangkan dan kelompok data dengan jumlah data paling sedikit adalah cluster 2, yang terdiri dari 16 data yaitu penerima 15 dan bukan penerima 1.

5. KESIMPULAN DAN SARAN

Berdasarkan hasil penelitian di Desa Pamulihan dengan 280 data Penerima Bantuan BLT, disimpulkan bahwa nilai K optimal untuk algoritma *K-means* adalah 9, dengan DBI sebesar 0.745, menunjukkan evaluasi clustering yang baik. Saran untuk penelitian selanjutnya mencakup penggunaan dataset yang lebih besar dan beragam guna mendapatkan hasil clustering yang lebih akurat serta pemilihan algoritma clustering yang berbeda untuk membandingkan kinerja algoritma dan menentukan yang terbaik sesuai karakteristik data. Kombinasi dari data yang lebih besar dan variasi algoritma dapat memberikan pemahaman yang lebih mendalam terhadap pola-pola dalam data Penerima Bantuan BLT, meningkatkan kehandalan hasil penelitian.

DAFTAR PUSTAKA

- [1] Y. Filki, "Algoritma K-Means Clustering dalam Memprediksi Penerima Bantuan Langsung Tunai (BLT) Dana Desa," *J. Inform. Ekon. Bisnis*, vol. 4, pp. 166–171, 2022, doi: 10.37034/infekon.v4i4.166.
- [2] S. Sari and J. N. Utamajaya, "Sistem Pendukung Keputusan Penerima Bantuan Langsung Tunai Dana Desa Menggunakan Metode Algoritma K-Means Clustering," *J. JUPITER*, vol. 14, no. 1, pp. 150–160, 2022.
- [3] A. Fitria Yulia and P. Bintoro, "K-Means Clustering in Determining the Eligibility of

- Recipients of Assistance for the Poor Case Study of Village Sukoharjo III,” *Int. J. Softw. Eng. Informatics*, vol. 1, no. 1, pp. 63–70, 2023, [Online]. Available: <https://journal.aisyahuniversity.ac.id/index.php/IJosei>
- [4] A. Waruwu, M. Yetri, and F. Setiawan, “Implementasi Data Mining Dalam Mengelompokkan Data penduduk Kurang Mampu Menggunakan Metode K-Means Clustering,” *J. Sist. Inf. Triguna Dharma (JURSI TGD)*, vol. 2, no. 6, p. 945, 2023, doi: 10.53513/jursi.v2i6.8965.
- [5] F. Azzahra, A. Zakir, and D. Irwan, “Pemanfaatan Algoritma K Means Untuk Menentukan Penerimaan Bantuan Langsung Tunai Pada Kecamatan Medan Area,” *Djtechno J. Teknol. Inf.*, vol. 3, no. 2, pp. 260–267, 2022, doi: 10.46576/djtechno.v3i2.2739.
- [6] W. Febrianti and Mustopa Husein Lubis, “Data Mining dalam Mengidentifikasi Calon Penerima Bantuan Satu Keluarga Satu Sarjana (SKSS) dengan Menggunakan Algoritma K-Means,” *J. Inform. Ekon. Bisnis*, vol. 4, pp. 53–58, 2022, doi: 10.37034/infek.v4i2.117.
- [7] L. G. Rady Putra and A. Anggrawan, “Pengelompokan Penerima Bantuan Sosial Masyarakat dengan Metode K-Means,” *MATRIK J. Manajemen, Tek. Inform. dan Rekayasa Komput.*, vol. 21, no. 1, pp. 205–214, 2021, doi: 10.30812/matrik.v21i1.1554.
- [8] A. M. Sholihah, N. Suarna, G. Dwilestari, and N. R, “Implementasi Metode K-means Clustering Untuk Menganalisa Penerima Bantuan di Desa Palasah,” *J. Inform. dan Teknol. Inf.*, vol. 1, no. 2, pp. 111–117, 2023, doi: 10.56854/jt.v1i2.121.
- [9] L. Anggraeni and M. Muslihudin, “Sosialisai Dan Pendampingan Pengelolaan Website Desa Kepada Aparatur Desa,” *J. PkM ...*, vol. 1, no. 2, pp. 41–50, 2020, [Online]. Available: <http://jurnalpkmpemberdayaan.com/index.php/PkMLP3K/article/view/8>
- [10] C. Suhaeni, A. Kurnia, and R. Ristiyanti, “Perbandingan Hasil Pengelompokan menggunakan Analisis Cluster Berhierarchy, K-Means Cluster, dan Cluster Ensemble (Studi Kasus Data Indikator Pelayanan Kesehatan Ibu Hamil),” *J. Media Infotama*, vol. 14, no. 1, 2018, doi: 10.37676/jmi.v14i1.469.
- [11] Z. Mustofa and Iman Saufik Suasana, “Algoritma Clustering K-Medoids Pada E-Government Bidang Information and Communication Technology Dalam Penentuan Status Edgi,” *J. Teknol. Inf. Dan Komun.*, vol. 9, no. 1, pp. 1–10, 2020, doi: 10.51903/jtikp.v9i1.162.