

Implementasi Algoritma K-Means Classifier Sebagai Pendukung Keputusan Penerima Dana Bantuan Siswa Miskin (Studi Kasus : SMKN Sukoharjo)

Aviv Fitria Yulia¹ Handoyo Widi Nugroho²

Institut Informatika dan Bisnis Darmajaya

¹aviv.fitriayulia.1821210010@mail.darmajaya.ac.id

²handoyo@darmajaya.ac.id

Abstract

Humans need education in their lives, education is an effort so that humans can develop their potential through the learning process. Poverty is a condition that is below the standard value line for minimum needs, both for food and non-food, the Poor Student Assistance Program (BSM) is a program that aims to eliminate barriers for poor students to attend school, the problem that occurs at SMKN SUKOHARJO is that the school has difficulty in Determination of recipients of Poor Student Assistance (BSM), this is due to the many criteria that must be considered in determining recipients of assistance such as: the number of students, which is 1044 students, parents' income, parental burden, distance traveled, and student scores. build a scholarship recipient selection system at SMKN SUKOHARJO with a selective and targeted process. Data utilization techniques are also called Data Mining. One of the most popular data mining methods is clustering using the K-Means algorithm. K-Means can process data without notifying the class label in advance. This research will produce three groups: eligible to receive assistance, considered to receive assistance, and not eligible to receive assistance. The data processing of the selection of grant receipts using the K-Means algorithm obtained a Davies bouldin index of 0.262. These results were considered quite good because the closer the results obtained were to zero, the better the similarity of the member data between clusters.

Keywords: Data Mining, K-Means

Abstrak

Manusia membutuhkan pendidikan dalam kehidupannya, Pendidikan merupakan usaha agar manusia dapat mengembangkan potensi dirinya melalui proses pembelajaran. Kemiskinan merupakan sebuah kondisi yang berada di bawah garis nilai standar kebutuhan minimum, baik untuk makanan dan non makanan, Program Bantuan Siswa Miskin (BSM) adalah program yang bertujuan untuk menghilangkan halangan siswa miskin untuk bersekolah, masalah yang terjadi di SMKN SUKOHARJO yaitu pihak sekolah mengalami kesulitan dalam penentuan penerima Bantuan Siswa Miskin (BSM), hal ini dikarenakan banyaknya kriteria yang harus di pertimbangkan dalam menentukan penerima bantuan seperti: jumlah siswa yaitu berjumlah 1044 siswa, penghasilan orang tua, beban orang tua, jarak tempuh, dan nilai siswa. membangun sistem seleksi penerima beasiswa di SMKN SUKOHARJO dengan proses selektif dan tepat sasaran. teknik pemanfaatan data disebut juga Data Mining. Salah satu metode Data mining yang cukup populer yaitu clustering dengan menggunakan algoritma K-Means. K-Means dapat mengolah data tanpa diberitahu lebih dahulu label kelasnya. Penelitian ini akan menghasilkan tiga kelompok : layak menerima bantuan, dipertimbangkan menerima bantuan, tidak layak menerima bantuan. Pengolahan data seleksi penerimaan dana bantuan menggunakan algoritma K-Means mendapatkan hasil Davies bouldin indeks sebesar 0,262. Hasil tersebut dinilai cukup baik sebab semakin dekat hasil yang diperoleh dengan angka nol, maka kemiripan data anggota antar cluster semakin baik.

Keywords: Data Mining, K-Means

1. PENDAHULUAN

Manusia membutuhkan pendidikan dalam kehidupannya, Pendidikan merupakan usaha agar manusia dapat mengembangkan potensi dirinya melalui proses pembelajaran atau cara lain yang dikenal dan diakui oleh masyarakat. Undang - Undang Dasar Negara Republik Indonesia Tahun 1945 Pasal 31 Ayat 1 Menyebutkan bahwa setiap warga negara berhak mendapat pendidikan, dan Ayat 3 Menegaskan bahwa Pemerintah mengusahakan dan menyelenggarakan satu sistem pendidikan nasional yang meningkatkan keimanan dan ketakwaan serta akhlak mulia dalam rangka mencerdaskan kehidupan bangsa yang diatur dengan undang-undang. Untuk itu, seluruh komponen bangsa wajib mencerdaskan kehidupan bangsa yang merupakan salah satu tujuan negara Indonesia.

Kemiskinan merupakan sebuah kondisi yang berada di bawah garis nilai standar kebutuhan minimum, baik untuk makanan dan non makanan, yang disebut garis kemiskinan (povertyline) atau batas kemiskinan (povertythreshold). Garis kemiskinan adalah sejumlah rupiah yang diperlukan oleh individu untuk dapat membayar kebutuhan makanan setara 2100 kilo kalori per orang per hari dan kebutuhan non-makanan yang terdiri dari perumahan, pakaian, pendidikan, transportasi, serta aneka barang dan jasa lainnya (BPS, Kemiskinan dan Ketimpangan, 2017).

Program Bantuan Siswa Miskin (BSM) adalah program yang bertujuan untuk menghilangkan halangan siswa miskin untuk berpartisipasi bersekolah dengan membantu siswa miskin memperoleh akses pendidikan yang layak, mencegah putus sekolah, menarik siswa miskin untuk kembali bersekolah. Melalui program BSM ini diharapkan anak usia sekolah keluarga miskin dapat terus sekolah dan menikmati pendidikan seperti anak yang lainnya.

Program Bantuan Siswa Miskin (BSM), masalah yang terjadi di SMKN SUKOHARJO yaitu pihak sekolah mengalami kesulitan dalam penentuan penerima Bantuan Siswa Miskin (BSM), hal ini dikarenakan banyaknya kriteria yang harus di pertimbangkan dalam menentukan penerima bantuan seperti: jumlah siswa yaitu berjumlah 1044 siswa, penghasilan orang tua, beban orangtua, jarak tempuh, dan nilai siswa.

Pengklasifikasian data siswa ini memiliki peran penting, karena dalam penentuan rekomendasi beasiswa kriteria yang perlu dipertimbangkan dan memakan waktu yang cukup lama tetapi belum tentu mendapatkan keputusan yang tepat dan akurat, maka pihak sekolah memerlukan sistem yang efektif dan efisien dalam waktu cepat dan dilakukan penilaian dengan mempertimbangkan berbagai kriteria tersebut.

Tidak semua siswa mendapatkan beasiswa ini, terdapat beberapa kriteria yang harus di penuhi berdasarkan peraturan yang telah di tetapkan oleh SMKN SUKOHARJO seperti nilai rata-rata rapor, penghasilan orang tua siswa, tanggungan orang tua siswa, dan jarak tempuh ke sekolah. Setiap siswa yang mengajukan beasiswa di lakukan proses seleksi terlebih dahulu berdasarkan kriteria tersebut. Sistem yang telah di lakukan selama ini masih menggunakan sistem manual atau belum di katakan sistem khusus yang digunakan untuk menentukan penerimaan bantuan, Sehingga masih sering menyebabkan terjadinya beberapa permasalahan, diantaranya yaitu masih sangat belum objektif dan tidak tepat sasaran dalam proses pemilihan siswa, mengakibatkan alokasi dana bantuan yang di berikan kepada siswa masih belum tepat dan belum membuat seluruh siswa yang memang benar-benar miskin yang berhak menerima dana tersebut. Oleh karena itu, untuk mengoptimalkan proses klasifikasi diterapkan sebuah metode Clustering K-Means.

Berdasarkan permasalahan tersebut membuat penulis untuk melakukan penelitian guna memberi solusi terhadap masalah yang terjadi dengan mengangkat judul, "IMPLEMENTASI ALGORITMA K-MEANS CLASSIFIER SEBAGAI PENDUKUNG KEPUTUSAN PENERIMA DANA BANTUAN SISWA MISKIN SMKN SUKOHARJO", kemudian dapat di lakukan pengalihan data untuk menentukan siapa yang berhak mendapatkan beasiswa dengan proses selektif dan tepat sasaran.

2. LANDASAN TEORI

2.1 Data Mining

Data mining adalah salah satu teknik penelusuran data untuk membangun sebuah model, kemudian menggunakan model tersebut agar dapat mengenali pola data yang lain yang tidak berada dalam basis data yang tersimpan. Dalam data mining, pengelompokan data juga dilakukan. Tujuannya adalah agar penulis dapat mengetahui pola dan tindak lanjut yang diambil. Semua hal tersebut bertujuan untuk mendukung kegiatan evaluasi agar sesuai dengan yang diharapkan (Prasetyo, 2012).

Data mining merupakan sebuah proses untuk menemukan pola atau pengetahuan yang bermanfaat secara

otomatis atau semi otomatis dari sekumpulan data dalam jumlah besar. Data mining hadir dianggap sebagai bagian dari Knowledge Discovery in Database (KDD) yaitu sebuah proses mencari pengetahuan yang bermanfaat dari data. KDD terdiri dari beberapa langkah (Santosa, 2007) yaitu:

- Pembersihan data (membuang noise dan data yang tidak konsisten).
- Integrasi data (penggabungan data dari beberapa sumber).
- Seleksi data (memilih data yang relevan yang akan digunakan untuk analisa).
- Data mining.
- Evaluasi model.
- Presentasi pengetahuan dengan teknik visualisasi.

2.2 Algoritma K-Means

Algoritma K-Means merupakan algoritma clustering yang berulang-ulang. Algoritma K-Means dimulai dengan pemilihan secara acak N , N disini merupakan banyaknya cluster yang ingin dibentuk. Kemudian tetapkan nilai- N secara random, untuk sementara nilai tersebut menjadi pusat dari cluster atau biasa disebut dengan centroid, mean atau “means”. Hitung jarak setiap data yang ada terhadap masing-masing centroid menggunakan rumus Euclidian hingga ditemukan jarak yang paling dekat dari setiap data dengan centroid. Lakukan langkah tersebut hingga nilai centroid tidak berubah, (Santosa, 2007).

Dari beberapa teknik clustering yang paling sederhana dan umum dikenal adalah clustering K-Means. Dalam teknik ini kita ingin mengelompokkan obyek kedalam N kelompok atau cluster. Untuk melakukan clustering, nilai N harus ditentukan terlebih dahulu. Biasanya user atau pemakai sudah mempunyai informasi awal tentang obyek yang sedang dipelajari, termasuk beberapa jumlah cluster yang paling tepat. Secara detail kita biasa menggunakan ukuran tidak miripan untuk mengelompokkan obyek kita. Tidak miripan bisa diterjemahkan dalam konsep jarak. Jika jarak dua obyek atau dua titik cukup dekat maka dua obyek itu mirip. Semakin dekat berarti semakin tinggi kemiripannya. Semakin tinggi jarak semakin tinggi tidak miripannya. Dalam penelitian yang mengungkapkan langkah-langkah pengerjaan algoritma K-Means (Santosa, 2007) yaitu:

Penentuan pusat cluster awal Dalam penentuan nilai N buah pusat cluster awal dilakukan pembangkitan bilangan random yang mempresentasikan urutan data input pusat awal cluster didapatkan dari data sendiri bukan dengan menentukan titik baru, yaitu dengan merandom pusat awal dari data. Perhitungan jarak dengan pusat cluster untuk mengukur jarak antara data dengan pusat cluster digunakan rumus Euclidian distance. Algoritma perhitungan jarak data dengan pusat cluster Ambil nilai data dan nilai pusat cluster. itung Euclidian distance data dengan tiap pusat cluster.

2.3 Cleaning Data

Data cleaning adalah serangkaian proses untuk mengidentifikasi kesalahan pada data dan kemudian mengambil tindakan lanjut, baik berupa perbaikan ataupun penghapusan data yang tidak sesuai. Prosedur data cleaning dilakukan untuk memastikan kualitas data yang digunakan, untuk memastikan kebenaran, konsistensi, dan kegunaan suatu data yang ada dalam dataset. Caranya adalah dengan mendeteksi adanya eror atau corrupt pada data, kemudian memperbaiki atau menghapus data jika memang diperlukan.

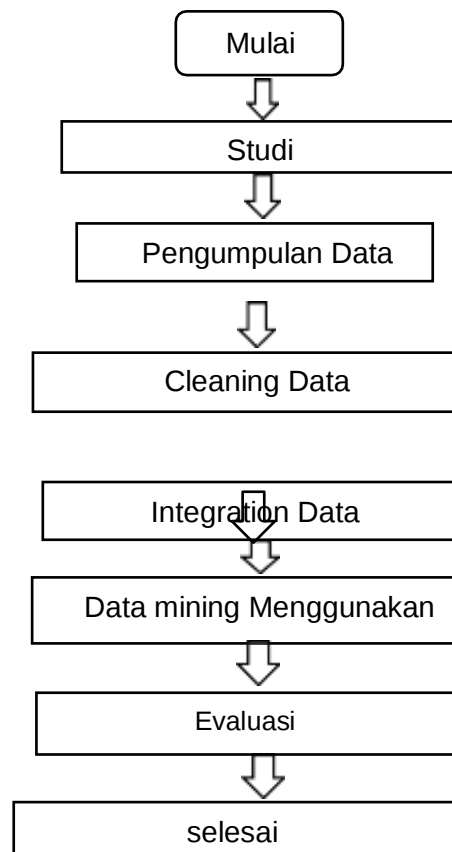
Terkadang, saat Anda menggabungkan beberapa data sources sekaligus, ada kemungkinan data terduplikasi atau bahkan salah label. Situasi seperti ini juga memerlukan data cleaning agar tidak muncul masalah yang lebih rumit.

2.4 Integration Data

Integration data adalah pengubahan data menjadi format yang sesuai untuk digunakan dalam proses data mining. Contoh : Setelah menemukan batu-batu yang cocok, selanjutnya penambang akan mulai mengkombinasikan untuk dijadikan batangan emas atau bentuk emas lainnya, Dalam data mining, data yang berhasil dibersihkan juga akan diintegrasikan.

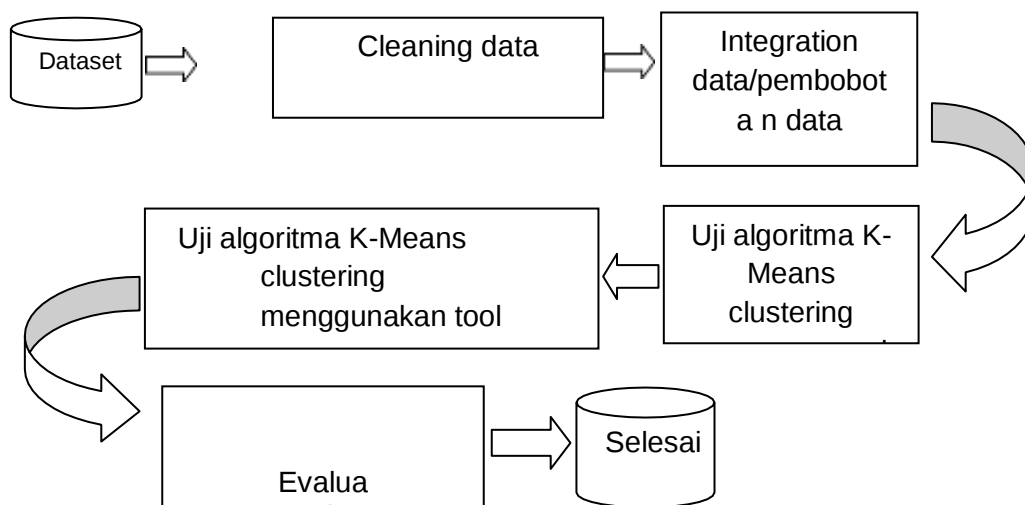
3 METODE PENELITIAN

Tahapan penelitian merupakan gambaran dari langkah-langkah yang dilakukan dalam penelitian. Tahapan penelitian yang dilakukan dapat dilihat pada gambar 3.1



Gambar 3.1 Tahapan Penelitian

SKEMA RISET / EKSPERIMEN



Gambar 3.2 Skema Riset / Eksperimen

3.1 Studi Awal

Studi awal dilakukan untuk mempelajari keadaan yang terjadi di SMKN SUKOHARJO, masalah-masalah yang terjadi di SMKN SUKOHARJO yaitu pihak sekolah mengalami kesulitan dalam penentuan penerima Bantuan Siswa Miskin(BSM), hal ini dikarenakan banyaknya kriteria yang harus di pertimbangkan dalam menentukan penerima bantuan seperti :jumlahsiswayaitu berjumlah 1044 siswa, penghasilan orang tua, beban orang tua, jarak tempuh, dan nilai siswa. Tidak semua siswa mendapatkan beasiswa ini, terdapat beberapa kriteria yang harus di penuhi berdasarkan peraturan yang telah di tetapkan oleh SMKN SUKOHARJO seperti nilai rata-rata rapor, penghasilan orang tua siswa, tanggungan orang tua siswa, dan jarak tempuh ke sekolah. Oleh karena itu, untuk mengoptimalkan proses klasifikasi diterapkan sebuah metode Clustering K-Means. Agar dapat di lakukan penggalan data untuk menentukan siapa yang berhak mendapatkan beasiswa dengan proses selektif dan tepat sasaran.

3.2 Pengumpulan Data

Pada tahap pengumpulan data ini dapat dijelaskan sebagai berikut:

Jenis data dan sumber data

Data yang digunakan dalam penelitian ini adalah sebanyak 1044 datasiswa. Data siswa merupakan jenis data kuantitatif, Data kuantitatif adalah data yang dapat dihitung, berupa angka. Data nilai yang dikumpulkan adalah data siswa semester tahun Pelajaran 2021/2022, Adapun sumber data diperoleh langsung dari operator SMKN SUKOHARJO. Berikut data mentah berupa nilai rata-rata rapor, penghasilan orang tua siswa, tanggungan orang tua siswa, dan jarak tempuh ke sekolah.

Nama	Penghasilan	Jml. tanggungan	Jarak Tempuh
1. DIO FIDIAN	Rp. 1.000.000 - Rp. 1.500.000	2	7
2. A. FIDIO FIDIAN	Rp. 100.000 - Rp. 100.000	1	7
3. Irena Aisy Saputra	Rp. 100.000 - Rp. 100.000	1	8
4. HANAN BRADY NURHARJANI	Raport dari Rp. 500.000	1	8
5. HEDUNANAMA ALIYAH	Rp. 100.000 - Rp. 100.000	1	8
6. Adhikyan Ananda Dhanu	Rp. 1.000.000 - Rp. 1.000.000	1	8
7. Adhikyan Ananda Dhanu	Raport dari Rp. 500.000	1	8
8. Alvin Rizki Al Auli	Rp. 300.000 - Rp. 100.000	1	8
9. Adhikyan Ananda Dhanu	Rp. 100.000 - Rp. 100.000	1	8
10. Alvin Rizki Al Auli	Raport dari Rp. 500.000	1	8
11. Adhikyan Ananda Dhanu	Rp. 1.000.000 - Rp. 1.000.000	1	8
12. Adhikyan Ananda Dhanu	Raport dari Rp. 500.000	1	8
13. Adhikyan Ananda Dhanu	Rp. 1.000.000 - Rp. 1.000.000	1	8
14. Adhikyan Ananda Dhanu	Rp. 1.000.000 - Rp. 1.000.000	1	8
15. Adhikyan Ananda Dhanu	Raport dari Rp. 500.000	1	8
16. Adhikyan Ananda Dhanu	Rp. 1.000.000 - Rp. 1.000.000	1	8
17. Adhikyan Ananda Dhanu	Rp. 1.000.000 - Rp. 1.000.000	1	8
18. Adhikyan Ananda Dhanu	Rp. 1.000.000 - Rp. 1.000.000	1	8
19. Adhikyan Ananda Dhanu	Rp. 1.000.000 - Rp. 1.000.000	1	8
20. Adhikyan Ananda Dhanu	Rp. 1.000.000 - Rp. 1.000.000	1	8
21. Adhikyan Ananda Dhanu	Raport dari Rp. 500.000	1	8

Gambar 3.3 Data Siswa Tahun Pelajaran 2021/2022

Metode pengumpulan data

Metode yang digunakan dalam proses pengumpulan data adalah sebagai berikut:

1. Observasi, Peneliti melakukan pengamatan secara langsung di lapangan untuk mempelajari permasalahan yang ada, yang erat kaitannya dengan objek yang diteliti. Pengamatan ini dilakukan untuk menambah pengetahuan mengenai topik yang diangkat peneliti.
2. Studi Pustaka, dilakukan dengan mengumpulkan dan mempelajari literatur yang berkaitan dengan penerapan metode K-Means Clustering. Sumber literatur berupa jurnal dan paper. Dengan metode ini diharapkan dapat mempertegas teori dan keperluan analisa.

3.3 Cleaning Data

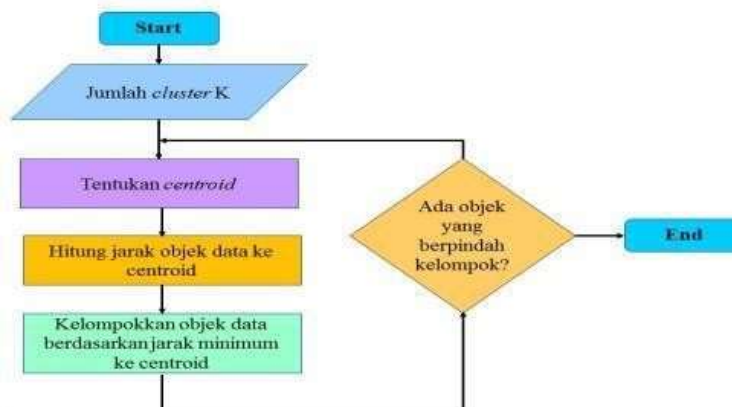
Setelah tahap seleksi data, dilakukan proses cleaning supaya data yang digunakan tidak terjadi duplikasi data, tidak ada data yang inkonsisten, dan tidak terdapat kesalahan pada data, seperti kesalahan cetak (typografi) dalam hal ini peneliti menggunakan manual (microsoftofficeexcel).

3.4 Integration Data

Pada tahap ini di lakukan konversi nilai atau memberikan pembobotan pada 1044 data asli oleh pihak terkait yang berwenang melakukan pembobotan. Dimana data di berikan bobot dari data nominal ke data numerik untuk mempermudah perhitungan K-MeansClustering.

3.5 Data mining Menggunakan K-Means Clustering

Data mining adalah proses mencari pola atau informasi menarik dalam data terpilih dengan menggunakan teknik atau metode tertentu. Pada proses data mining untuk menentukan kelompok belajar siswa berdasarkan kemampuan peneliti menggunakan K-Means Clustering. Secara umum urutan proses Clustering dengan algoritma K-Means ditunjukkan pada gambar 3.3



Gambar 3.4 Flowchart Algoritma K-Means Clustering(Eko Prasetyo, 2012)

Pada proses pengumpulan data ada 4 parameter yang akan digunakan dalam pengolahan data yaitu, Data nama siswa, Penghasilan orang tua, Jumlah tanggungan orang tua, Jarak tempuh siswa, Dan nilai rata-rata siswa.

- Proses Clustering Menggunakan Algoritma K-Means

Data yang sudah dijadikan sampel akan dilakukan pengolahan data dengan proses clustering dengan menggunakan algoritma K-Means sehingga didapatkanlah hasil pengelompokan data yang diinginkan. Adapun langkah dalam cluster dengan algoritma

K-Means yaitu :

1. Menentukan Jumlah Cluster

Menentukan jumlah cluster yang digunakan pada data siswa sebanyak 3 cluster.

2. Menentukan Centroid

Penentuan pusat awal cluster (centroid) ditentukan secara random atau acak yang diambil dari data yang telah didapatkan. Nilai cluster1 diambil dari baris ke-30, nilai cluster2 pada baris ke-8, nilai cluster 3 pada baris ke-5 seperti pada tabel 3.1 di bawah ini:

Tabel 3.1. Nilai Awal Pusat Cluster

Data ke 30	Cluster 1	5	4	5	2
Data ke 8	Cluster 2	2	4	1	2
Data ke 5	Cluster 3	1	3	4	1

Tabel 3.1 merupakan Nilai cluster1 diambil dari baris ke-30, nilai cluster2 pada baris ke-8, nilai cluster 3 pada baris ke-5

1. Menghitung jarak setiap data terhadap pusat cluster.

Misalnya untuk menghitung jarak data pertama dengan nilai pusat cluster pertama:

$$D = \sqrt{(x^1 - x^1)^2}$$

$$d = \sqrt{(3 - 5)^2 + (4 - 4)^2 + (5 - 5)^2 + (2 - 2)^2}$$

$$d = 2$$

Jarak data pertama dengan nilai pusat cluster kedua adalah:

$$d = \sqrt{(3 - 2)^2 + (4 - 4)^2 + (5 - 1)^2 + (2 - 2)^2}$$

$$= 4,123106$$

Jarak data pertama dengan nilai pusat cluster ketiga adalah:

$$d = \sqrt{(3-1)^2 + (4-3)^2 + (5-4)^2 + (2-1)^2}$$

$$d=2,645751$$

3.6 Evaluasi

Pola informasi yang dihasilkan dari proses data mining perlu ditampilkan dalam bentuk yang mudah dimengerti oleh pihak yang berkepentingan. Pada tahap ini setelah mendapatkan hasil Cluster dari masing-masing selanjutnya akan diinterpretasikan dan dievaluasi. Interpretasi dilakukan dengan menggunakan bahasa yang mudah dipahami.

4. HASIL DAN PEMBAHASAN

4.1 Hasil

a. Proses Clustering

Pada tahap ini akan dilakukan proses utama yaitu segmentasi atau pengelompokan data penerima dana bantuan siswa miskin. Berikut ini merupakan penerapan algoritma K-Means dengan asumsi bahwa parameter input adalah jumlah dataset sebanyak n data dan jumlah inisialisasi centroid k = 3 sesuai dengan penelitian. Data yang diambil untuk penelitian berjumlah 1044 data siswa untuk dijadikan penerapan ke dalam algoritma K-Means. Percobaan dilakukan dengan menggunakan parameter-parameter berikut :

Jumlah cluster :

3 Jumlah data :

1044 Jumlah

atribut :4

b. Pengujian Dengan Tool Rapid Miner

Pada penelitian ini penulis menggunakan tool rapid miner sebagai alat pengujian data set. Adapun tahapan pengujian yang dilakukan yaitu sebagai berikut :



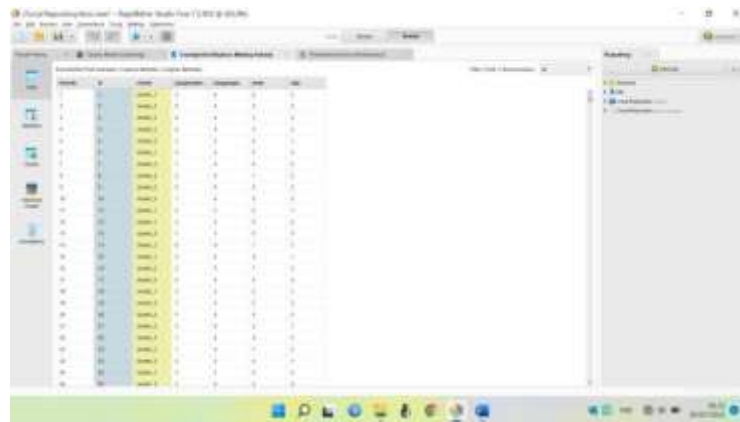
Gambar 4.1 Design Proses

Pada tahap ini dilakukan 4 proses yaitu :

- Read excel
- Tahapan ini dilakukan operasi penginputan dataset berupa file berekstensi.xls Data siswa calon penerima dana bantuan siswa miskin.
- Replace missing value
Tahapan ini dilakukan operasi pengisian nilai yang hilang dengan nilai maksimal.
- Clustering
Tahapan ini dilakukan operasi clustering sebagai algoritma yang digunakan pada penelitian ini.

- Performance

Tahapan ini dilakukan operasi pencarian nilai davies bouldin index.



Gambar 4.2 Example set result

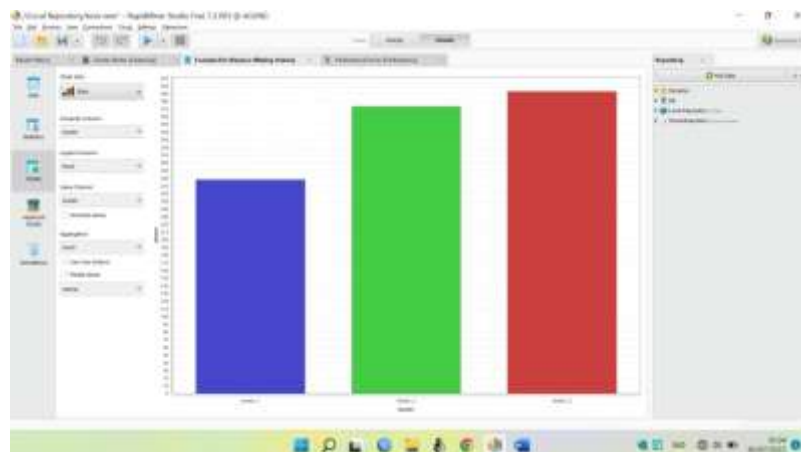
Pada tahapan ini ditampilkan hasil dari klasterisasi data . Label cluster terbagi menjadi tiga kelompok yaitu cluster 0, cluster 1, cluster 2. Pembagian ini berdasarkan hasil kedekatan tiap masing-masing data dengan jarak terdekat (k).



Gambar 4.3 Cluster Model

Pada tahapan ini ditampilkan hasil pembagian data terhadap tiap cluster.

Cluster 0 memiliki 393 anggota, Cluster 1 memiliki 278 anggota, Cluster 2 memiliki 373 anggota dari total 1044 dataset yang di uji.



Gambar 4.4 Chart clustering

Pada tahapan ini ditampilkan hasil pengelompokan data dalam bentuk diagram dengan warna. Warna merah mengartikan cluster 0, warna biru mengartikan cluster 1, dan warna hijau mengartikan cluster 2.



Gambar 4.5 Performance Vector

Pada tahapan ini adalah menghitung rata-rata jarak dari tiap-tiap cluster .

4.2 Pembahasan

a. hasil clustering K-Means

Setelah dilakukan pengujian dengan tool rapid miner , maka dapat di simpulkan sebagai berikut :

- Cluster 0 memiliki 393 anggota,
- Cluster 1 memiliki 278 anggota,
- Cluster 2 memiliki 373 anggota
- dari total 1044 dataset yang di uji

b. Hasil Performance Vektor

Evaluasi hasil dari average within centroid distance mendekati angka 0 mengartikan bahwa masing-masing anggota di dalam cluster berada dalam jarak yang berdekatan. Evaluasi menggunakan davies bouldin indeks memiliki skema internal cluster yang dilihat dari kuantitas dan kedekatan antar hasil cluster.

Semakin kecil nilai davies bouldin indeks yang diperoleh (non-negatif) ≥ 0 , maka semakin baik cluster yang diperoleh dari pengelompokan menggunakan metode clustering. Hasil perhitungan menggunakan algoritma K- Means menunjukkan nilai 0,262. Angka tersebut memiliki arti masing-masing objek dalam cluster tersebut memiliki kesamaan yang cukup baik karena mendekati angka 0.

5. KESIMPULAN

Berdasarkan hasil penelitian yang dilakukan oleh penulis, dapat diambil kesimpulan sebagai berikut :

- Penerapan algoritma K-Means membagi dataset menjadi tiga kelompok yaitu layak menerima bantuan siswa miskin ,dapat di pertimbangkan menerima bantuan siswa miskin dan tidak layak menerima bantuan siswa miskin.
- Hasil pengujian mendapatkan nilai devies bouldin indeks sebesar 0,262 yang memiliki arti kesamaan antar anggota cluster yang cukup baik.
- Hasil dari pengujian ini sangat layak untuk dijadikan rekomendasi penerima bantuan siswa miskin di sekolah (SMKN SUKOHARJO).

5.2 Saran

Mengingat masih banyaknya hal-hal yang belum dapat diimplementasikan dari penelitian ini, maka penulis mempertimbangkan beberapa saran yaitu :

1. Hasil clustering yang terbentuk dapat dikembangkan menjadi basis pengetahuan untuk sistem pendukung keputusan penerima dana bantuan di sekolah sesuai dengan kemiripanya.
2. Melakukan kombinasi dengan metode atau pendekatan yang lain guna mendapatkan hasil penelitian yang lebih baik.

DAFTAR PUSTAKA

- Andriani, A. (2013). Sistem Pendukung Keputusan Berbasis Decision Tree Dalam Pemberian Beasiswa Studi Kasus: Amik “ Bsi Yogyakarta .” Seminar Nasional Teknologi Informasi Dan Komunikasi 2013 (SENTIKA 2013), 2013(Sentika), 163–168.
- Khan, I. A., & Choi, J. T. (2014). An application of educational data mining (EDM) technique for scholarship prediction. *International Journal of Software Engineering and Its Applications*, 8(12), 31–42.
- Wicaksono, A. E. (2016). Implementasi Data Mining Dalam Pengelompokan Peserta Didik di Sekolah untuk Memprediksi Calon Penerima Beasiswa Dengan Menggunakan Algoritma K-Means (Studi Kasus SMA N 6 Bekasi). *Jurusan Teknik Informatika, Universitas Gunadarma*, 21(3), 208.
- Kusnawi. (2012). Pengantar solusi data mining. *Seminar Nasional Teknologi 2012(SNT 2012)*, 2012(November), 1– 9
- Hijriana, N., Studi, P., Informatika, T., & Islam, U. (2017). Penerapan Metode Decision Tree Algoritma C4 . 5.
- MAURINA, D. (2015). Penerapan Data Mining Untuk Rekomendasi Beasiswa Pada Sma Muhammadiyah Gubug Menggunakan Algoritma C4.5.
- Utomo Pramudi. (2013). Analisis Kontribusi Pemberian Beasiswa Terhadap Peningkatan Prestasi Akademik Mahasiswa Fakultas Teknik Universitas Negeri Yogyakarta.
- Rismayanti. (2016). Implementasi algoritma c4.5 untuk menentukan penerima beasiswa di stt harapan medan, 12(2), 116–120.
- Achmad Basuki, I. S. (2009). *Decision Tree. Review Literature And Arts Of The Americas*.
- Sheila, P. (2015). Sistem Pendukung Keputusan Dengan Menggunakan Decision Tree Di Sekolah Menengah Pertama (Studi Kasus Di SMPN 2 Rembang) Sistem Pendukung Keputusan Dengan Menggunakan Decision Tree Di Sekolah Menengah Pertama (Studi Kasus di SMPN 2 Rembang)
-