

Analysis of The Theme Clustering Algorithm Using K-Means Method

Analisis Algoritma Clustering Tema Skripsi Menggunakan Metode K-Means

Erwin Dwika Putra¹⁾; Muhammad Husni Rifqo²⁾; Dwita Deslianti³⁾; Krismiyan⁴⁾

^{1,2,3,4)} Program Studi Teknik Informatika, Fakultas Teknik, Universitas Muhammadiyah Bengkulu

Email: ¹⁾ erwindikap@gmail.com; ²⁾ mhrifqo@umb.ac.id; ³⁾ dwitadeslianti@umb.ac.id;

⁴⁾ krismiyan739@gmail.com

How to Cite :

Putra, E.D.; Rifqo, M.H.; Deslianti, D.; Krismiyan. (2022). Analisis Algoritma Clustering Tema Skripsi Menggunakan Metode K-Means, Jurnal Komputer, Informasi dan Teknologi, 2 (2). DOI: <https://doi.org/10.53697/jkomitek.v2i2>

ARTICLE HISTORY

Received [10 September 2022]

Revised [21 Oktober 2022]

Accepted [09 November 2022]

Keywords :

Data Mining, Clustering,
K-Means

This is an open access article under the
[CC-BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license



ABSTRAK

Penelitian ini berjudul analisis algoritma clustering tema skripsi menggunakan metode k-means. Permasalahan utama adalah bagaimana kita dapat mengetahui mana tema yang paling banyak diminati oleh para mahasiswa skripsi di Fakultas Teknik Universitas Muhammadiyah Bengkulu. Pengklasteran ini menggunakan metode K-means. Metode K-Means dipilih karena metode ini merupakan salah satu metode pengklasteran data *nonhierarki* (sekatan) yang berusaha mempartisi data ke dalam bentuk dua atau lebih *cluster* yang berkarakteristik sama dimasukkan ke dalam satu *cluster* yang sama. Tujuan dari penelitian ini dapat membantu calon mahasiswa yang akan skripsi dalam mengetahui mana tema yang lebih banyak diminat.

ABSTRACT

The title of this research is the analysis of the thesis theme clustering algorithm using the k-means method. The main problem is how we can find out which theme is most in demand by thesis students at the Faculty of Engineering, University of Muhammadiyah Bengkulu. This clustering uses the K-means method. The K-Means method was chosen because this method is one of the non-hierarchical data clustering methods that seeks to partition data into two or more clusters with the same characteristics included in the same cluster. The purpose of this research is to help prospective students who will write their thesis in knowing which themes are more interested in them..

PENDAHULUAN

Setiap tahun penelitian skripsi mahasiswa semakin bertambah dan memungkinkan mahasiswa mengambil topik yang sama atau serupa, dikarenakan belum adanya sistem atau informasi yang melakukan pengklasteran judul berdasarkan tema skripsi. Di mana dokumen skripsi ini dapat diklasterkan berdasarkan pola sesuai dengan empat tema yang ada. Teks judul skripsi tersebut akan dilakukan pembobotan kata menggunakan metode *Text Mining* yang kemudian data vektor hasil pembobotan ini digunakan untuk melakukan pengklasteran dokumen skripsi berdasarkan empat tema skripsi yang ada. Metode pengklasteran yang digunakan yaitu metode *K-Means* dengan pemilihan *centroid* awal klaster yang telah dikembangkan menggunakan *Improved K-Means*. Untuk mengatasi permasalahan diatas, maka penelitian ini dilakukan untuk menguji metode

yang sudah ada dan diterapkan pada kasus *Clustering* tema skripsi menggunakan metode *K-Means*. [1]

Berdasarkan hasil pengumpulan data judul skripsi mahasiswa Program Studi Teknik Informatika Fakultas Teknik Universitas Muhammadiyah Bengkulu menampilkan berbagai macam judul berdasarkan empat perbedaan tema yang ada. Untuk itu diperlukan pengklasteran judul berdasarkan tema agar mengetahui tema skripsi apa yang paling banyak diambil oleh mahasiswa Program Studi Teknik Informatika Fakultas Teknik Universitas Muhammadiyah Bengkulu. Pembagian ini dilakukan berdasarkan empat tema yang sudah ditetapkan. Dengan pendekatan pengklasteran *K-Means*, tema skripsi dilakukan berdasarkan data judul skripsi yang ada dan tema skripsi yang sudah ditetapkan. Pada penelitian ini dilakukan pengklasteran tema skripsi menggunakan algoritma *K-Means*.

Salah satu teknik analisis *Data Mining* adalah *cluster analysis* yang lebih dikenal dengan *Clustering*. *Clustering* merupakan metode analisis data yang tujuannya mengklaster data dengan karakteristik yang sama ke suatu wilayah yang sama. Salah satu pendekatan yang digunakan dalam mengembangkan metode *clustering* yaitu metode *K-Means*, dimana metode ini merupakan salah satu metode pengklasteran data *nonhierarki* (sekatan) yang berusaha mempartisi data ke dalam bentuk dua atau lebih *cluster* yang berkarakteristik sama dimasukkan ke dalam satu *cluster* yang sama. [2]

Data Mining sudah banyak diterapkan dalam berbagai bidang, salah satunya metode *clustering* khususnya metode *K-Means*, antara lain oleh Andika (2008) menjelaskan tentang algoritma *K-Means clustering* dapat diimplementasikan dengan model perangkat lunak untuk verifikasi citra sidik jari poin *minutiae*. Suprihatin (2011) menyatakan bahwa teknik *clustering K-Means* digunakan untuk menentukan nilai huruf ujian akhir pada Universitas Ahmad Dahlan (UAD). [2]

Metode *K-Means* pertama kali diperkenalkan oleh MacQueen JB pada tahun 1976. Metode ini adalah salah satu metode *non hierarchi* yang umum digunakan. Metode *K-Means* sendiri memiliki kelebihan, yakni mampu mengklasterkan data besar dengan sangat cepat (Teknomo : 2007). Metode ini termasuk dalam teknik penyekatan (*partition*) yang membagi atau memisahkan objek ke *k* daerah bagian yang terpisah. Pada *K-Means*, setiap objek harus masuk dalam gelombang tertentu, tetapi dalam satu tahapan proses tertentu, objek yang sudah masuk dalam satu gelombang, pada satu tahapan berikutnya objek akan berpindah ke gerombol lain, yakni (1) jumlah *cluster* yang diinginkan; (2) Hanya memiliki atribut bertipe numerik. Metode *K-Means* merupakan metode *clustering* yang paling sederhana dan umum. *K-Means* merupakan algoritma *clustering* yang berulang-ulang. Algoritma *K-Means* dimulai dengan pemilihan secara acak *K*, *K* disini merupakan banyaknya *Cluster* yang ingin dibentuk. Kemudian tetapkan nilai-nilai *K* secara random, untuk sementara nilai tersebut menjadi pusat dari *cluster* atau biasa disebut dengan *centroid*, mean atau "means". [3]

LANDASAN TEORI

Penelitian Terkait

Penelitian yang dilakukan oleh Benri Melpa Metisen, Herlina Latipa sari pada tahun 2015 yang berjudul Analisis *clustering* menggunakan metode *k-means* dalam pengelompokkan penjualan produk pada swalayan Fadhila. Penelitian tersebut menghasilkan kesimpulan tentang aplikasi perangkat lunak Tanagra pada data mining. Tanagra adalah *software* data mining yang dapat digunakan untuk mengakses beberapa metode data mining yang ada. Aplikasi ini menggunakan dataset input. Dalam melaksanakan pengujian algoritma ini data yang dipakai adalah data barang di swalayan Fadhilla Bengkulu. Dalam penerapan ini, digunakan penerapan *clustering* dengan menggunakan algoritma *K-means*. Dari data yang diolah dengan sampel data yang diambil di swalayan Fadhilla Bengkulu, maka menghasilkan dua jenis kelompok data. Yaitu data penjualan rendah dan data penjualan tinggi. Sehingga dengan adanya pengelompokkan data ini pihak

swalayan Fadhillah Bengkulu dapat mengetahui jenis barang yang laris terjual dan tidak. Sehingga barang yang ada di gudang tidak menumpuk. [4]

Penelitian dengan judul penerapan datamining menentukan minat mahasiswa di perpustakaan Universitas Bina Darma Palembang menggunakan metode *clustering* oleh Lemi Iryani pada tahun 2020 menghasilkan bahwa Penelitian ini dilakukan untuk menentukan minat baca mahasiswa Bina Darma dengan metode *clustering*. Data yang digunakan dalam penelitian ini adalah database perpustakaan Bina Darma yang terdiri dari tiga tabel yaitu tabel *collection*, tabel *patron* dan tabel *visitor* yang akan diintegrasikan dan diseleksi untuk digabungkan sebagai data *warehouse*. Penerapan data mining untuk menentukan judul buku yang paling banyak diminati mahasiswa dan program studi apa yang paling banyak meminjam buku dan membaca buku di perpustakaan bina darma. Informasi yang didapat adalah bahwa Program Studi Sistem Informasi lebih berminat meminjam dan membaca buku dibandingkan dengan program studi teknik industri. Judul buku yang paling banyak diminati adalah judul buku dasar-dasar pemrograman sedangkan buku yang paling rendah diminati mahasiswa adalah judul buku panduan praktis menggunakan *microsoft word* dan *cluster* kategori buku yang paling banyak diminati mahasiswa adalah kategori buku ilmu komputer sedangkan kategori buku bahasa Indonesia lebih rendah. [5]

Penelitian yang dilakukan oleh Wahyu Fuadi, Ar Razi, Dedi Farida pada tahun 2022 yang berjudul "automasi penentuan tren topik skripsi menggunakan algoritma *k-means clustering*". Penelitian tersebut menghasilkan kesimpulan Skripsi merupakan salah satu mata kuliah wajib yang harus diselesaikan oleh setiap mahasiswa untuk mendapatkan status sarjana (S1). Kebutuhan mahasiswa terhadap informasi dalam bentuk skripsi semakin meningkat, sehingga pemberian label tren topik diharapkan dapat membantu mahasiswa dalam mengetahui topik apa yang sedang tren ditahun sebelumnya tanpa harus membaca secara keseluruhan skripsi yang ada di perpustakaan. Pada penelitian ini, pengambilan data skripsi berdasarkan data 5 tahun sebelumnya yaitu tahun 2015 - 2020 pada Jurusan Teknik Informatika Universitas Malikussaleh (Unimal). Tren topik skripsi digolongkan ke dalam 5 kategori, yaitu Data Mining, Kecerdasan Buatan, Pengolahan Citra, SPK, dan GIS. Klasifikasi topik skripsi berdasarkan judul dan abstrak. Bahasa pemrograman yang digunakan yaitu PHP dan database MySQL. Adapun metode yang digunakan dalam automasi penentuan tren topik skripsi adalah *text mining* dan algoritma *K-Means clustering*. Adapun keutamaan dari penelitian ini adalah untuk menghasilkan sebuah aplikasi yang dapat mengatasi kendala mahasiswa yaitu kesulitan dalam menentukan topik judul skripsi sehingga lama dalam proses pembuatan proposal skripsi. Dengan adanya aplikasi ini diharapkan dapat mengatasi kendala tersebut, sehingga mempermudah mahasiswa dalam mengetahui tren topik skripsi di tahun sebelumnya dan memudahkan dalam menentukan topik judul proposal skripsi. Pada penelitian ini persentase keakurasian tren topik skripsi dengan menggunakan algoritma *K-means clustering* adalah 84 % dari 70 data uji. [6]

Penelitian yang dilakukan oleh Muhammad Faishal Riyadhi pada tahun 2019 dengan judul aplikasi *text mining* untuk automasi penentuan tren topik skripsi dengan metode *k-means clustering*. Dalam penelitian ini dapat disimpulkan dengan banyaknya mahasiswa yang akan mengerjakan tugas akhir, maka diperlukan suatu sistem yang dapat memberikan informasi tentang tren topik skripsi apa saja yang sedang populer pada tahun-tahun tertentu. Oleh karena itu melalui penelitian ini dikembangkan suatu aplikasi yang dapat bekerja secara semiotomatis dengan memanfaatkan teknologi *Text Mining* dan Algoritma *K-Means Clustering*. Dari hasil penelitian yang telah dilakukan maka didapatkan hasil bahwa, sistem yang dibuat dapat membantu para mahasiswa untuk mengetahui informasi tren topik skripsi apa saja yang sedang tren di program studi sistem komputer. Untuk proses analisis menggunakan metode *k-means clustering*, tingkat keberhasilan yang didapat sebesar 66.66% untuk proses matematis. Dan untuk proses sistemnya sebesar 33.33% untuk data yang sama dengan proses matematis. [7]

Penelitian yang dilakukan oleh Rima Dias Ramadhani, pada tahun 2013 yang berjudul data mining menggunakan algoritma *K-means clustering* untuk menentukan strategi promosi Universitas Dian Nuswantoro. Proses penerimaan mahasiswa baru Universitas Dian Nuswantoro menghasilkan

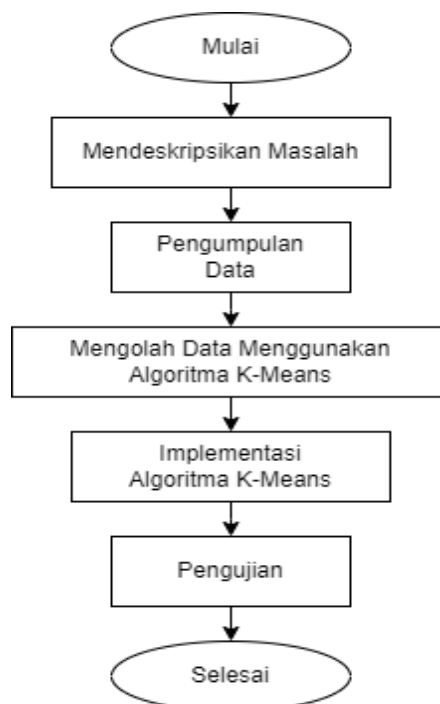
data mahasiswa yang sangat berlimpah berupa data profil mahasiswa dan data kegiatan belajar mengajar. Hal tersebut terjadi secara berulang dan menimbulkan penumpukan terhadap data mahasiswa, sehingga mempengaruhi pencarian informasi terhadap data tersebut. Penelitian ini bertujuan untuk melakukan pengelompokan terhadap data mahasiswa Universitas Dian Nuswantoro dengan memanfaatkan proses data mining dengan menggunakan teknik *Clustering*. Metode yang digunakan adalah CRISP-DM dengan melalui proses business understanding, data understanding, data *preparation*, *modeling*, *evaluation* dan *deployment*. Algoritma yang digunakan untuk pembentukan *cluster* adalah algoritma *K-Means*. *K-Means* merupakan salah satu metode data *non-hierarchical clustering* yang dapat mengelompokkan data mahasiswa ke dalam beberapa *cluster* berdasarkan kemiripan dari data tersebut, sehingga data mahasiswa yang memiliki karakteristik yang sama dikelompokkan dalam satu *cluster* dan yang memiliki karakteristik yang berbeda dikelompokkan dalam *cluster* yang lain. Implementasi menggunakan *RapidMiner 5.3* digunakan untuk membantu menemukan nilai yang akurat. Atribut yang digunakan adalah kota asal, program studi dan IPK mahasiswa. *Cluster* mahasiswa yang terbentuk adalah tiga *cluster*, dengan *cluster* pertama 804 mahasiswa, *cluster* kedua 2792 mahasiswa dan *cluster* ketiga sejumlah 223 mahasiswa. Hasil dari penelitian ini digunakan sebagai salah satu dasar pengambilan keputusan untuk menentukan strategi promosi berdasarkan *cluster* yang terbentuk oleh pihak admisi UDINUS. [8]

Dari penelitian di atas, dapat penulis simpulkan bahwa adanya analisis *algoritma clustering* tema skripsi menggunakan metode *k-means* akan membantu para mahasiswa untuk mengetahui informasi tema skripsi yang paling banyak diminati.

METODE PENELITIAN

Tahap Penelitian

Tahap penelitian merupakan gambaran dari tahap - tahap penelitian yang dilakukan dalam analisis algoritma *clustering* tema skripsi menggunakan metode *k-means*. Tahap-tahap penelitian tersebut dapat dilihat pada gambar 1.



Gambar 1 Tahap - Tahap Penelitian

Metode Pengumpulan Data

Penulis mendapatkan data yang akan diolah dari Program Studi Teknik Informatika Fakultas Teknik Universitas Muhammadiyah Bengkulu. Proses mengumpulkan data ini merupakan proses untuk pengumpulan data mentah berupa judul-judul skripsi dari tahun 2018-2020 sebelum dilakukan proses lainnya. Dalam proses ini diperlukan inisialisasi untuk mempermudah proses *Clustering* pada algoritma *K-means*. Data yang digunakan ini bersifat numerik, maka data nominal seperti pemilihan harus diubah dalam bentuk nilai numerik.

Data yang digunakan disini data skripsi dari tahun ajaran 2018 sampai dengan 2020. Pada penelitian ini akan diinputkan sebanyak 200 data berupa judul-judul skripsi dari mahasiswa Program Studi Teknik Informatika Fakultas Teknik Universitas Muhammadiyah Bengkulu.

Penulis melakukan proses pengolahan data dengan menggunakan algoritma *k-means* dengan melakukan pengklasteran. Dalam *clustering*, pada tahap ini dimulai dengan menentukan jumlah *cluster*, setelah itu ditentukan titik pusat *cluster* (*centroid*) dan menghitung jarak antara data dengan titik pusat *cluster* dengan persamaan jarak *Euclidean*. Kemudian data tersebut di alokasikan ke *cluster* terdekat sehingga diketahui data tersebut berada di *cluster* mana. Tahap selanjutnya adalah pengambilan keputusan, dimana data yang telah diklasterkan akan diranking. Dari data tersebut dibuat sebuah matriks keputusan berdasarkan kriteria dan alternatif yang dimiliki yang kemudian dinormalisasi. Selanjutnya menentukan solusi ideal positif dan negatif diikuti dengan menghitung ukuran utilitas.

HASIL DAN PEMBAHASAN

Hasil

Implementasi K-Means Pada Rapidminer

Di bawah ini merupakan langkah-langkah kerja pengimplementasian pada Rapidminer, yaitu sebagai berikut:

Langkah pertama pembuatan format tabular pada Ms. Excel

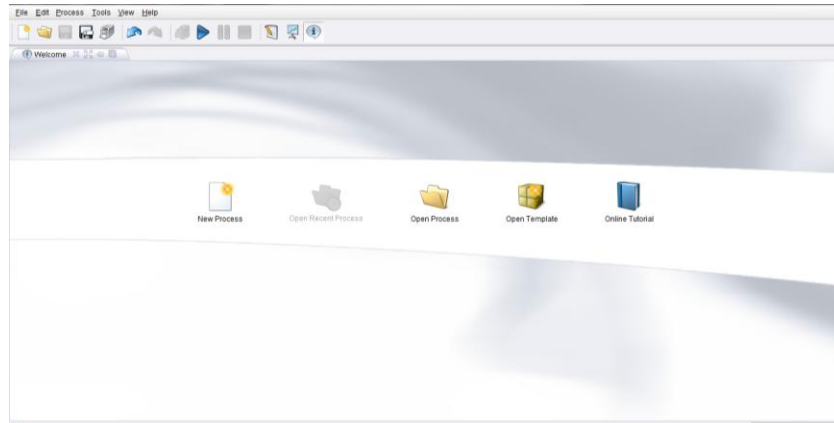
Tabel 1. Format Tabular Data Inialisasi

No	PRODI	TEMA	MINAT
1	1	1	158
2	1	2	21
3	1	3	14
4	1	4	4
5	1	5	5
6	2	1	12
7	2	2	8

Format tabular tersebut disimpan pada lembar kerja Ms.Excel yang menjadi *database* penyimpanan data tabular, dengan *Save as type* menjadi Excel 97/2003 Workbook. Ms. Excel tersebut akan dikoneksikan ke Rapidminer.

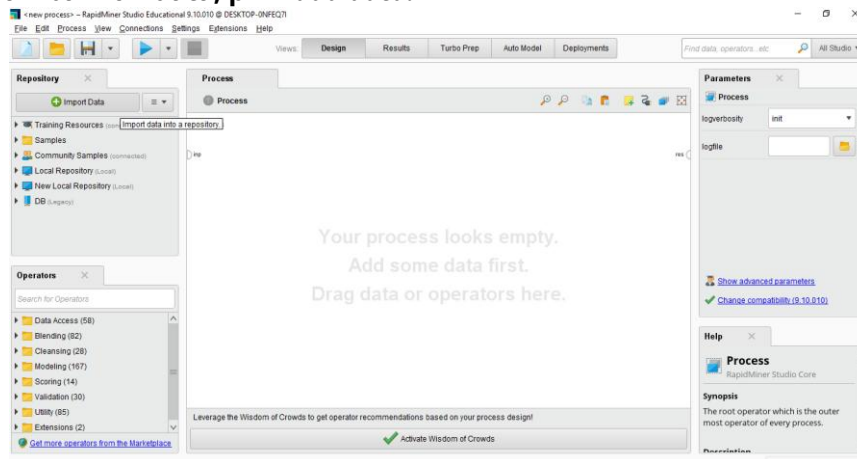
Buka aplikasi Rapidminer, lalu klik *New Process - Blank*.

Gambar 1. merupakan tampilan awal dari Rapidminer, ketika ingin memulai pilih menu *New Process* lalu *Blank* untuk memulai mengkoneksikan *database* ke Rapidminer.



Gambar 1. Tampilan Awal Rapidminer

1. Untuk menambahkan data, pilih *add data*.

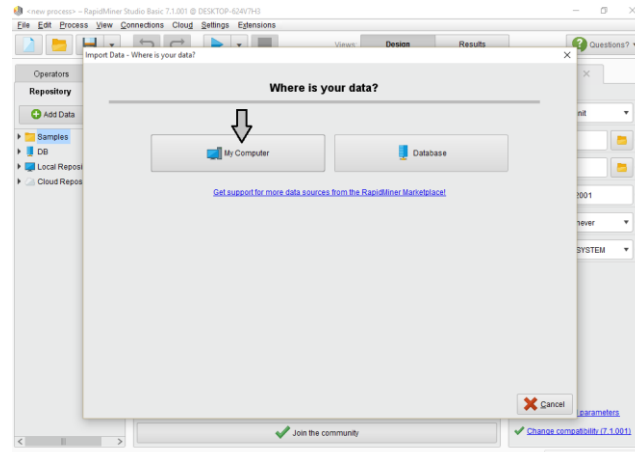


Gambar 2. Menambahkan data

Gambar 2. pada *Tab Menu Repository* pilih *Add Data* untuk menambahkan data yang baru dibuat, atau yang belum tersimpan pada Rapidminer.

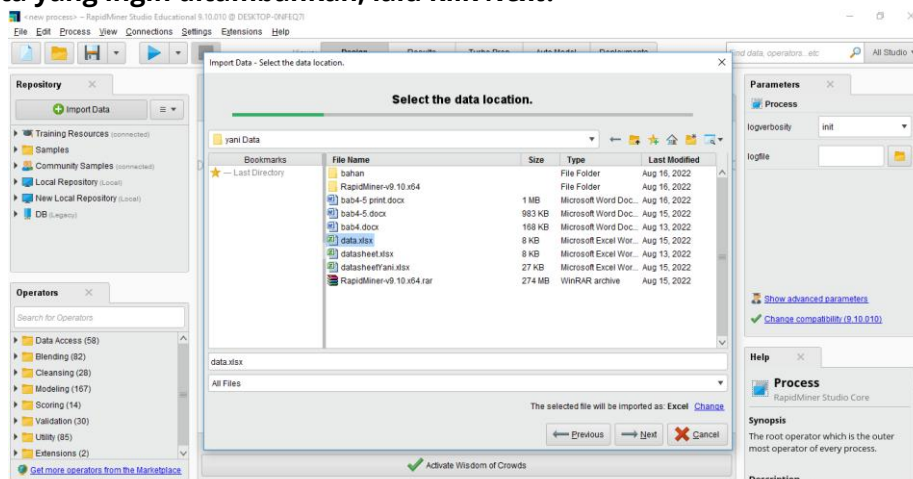
2. Pilih *My Computer*

Gambar 3. Pilih *My Computer* untuk membuka folder, yaitu tempat dimana kita menyimpan database yang ingin kita tambah.



Gambar 3. Pilih My Computer

3. Pilih data yang ingin ditambahkan, lalu klik Next.

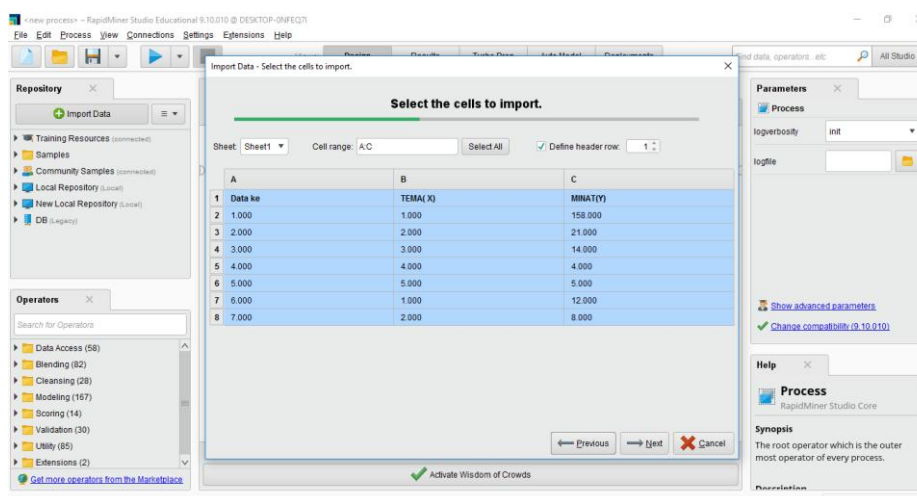


Gambar 4. Mencari Database yang Telah Disimpan

Pada Gambar 4. pilih sebuah file Excel yang ingin ditambahkan ke Rapidminer, lalu dilanjutkan dengan mengklik Next.

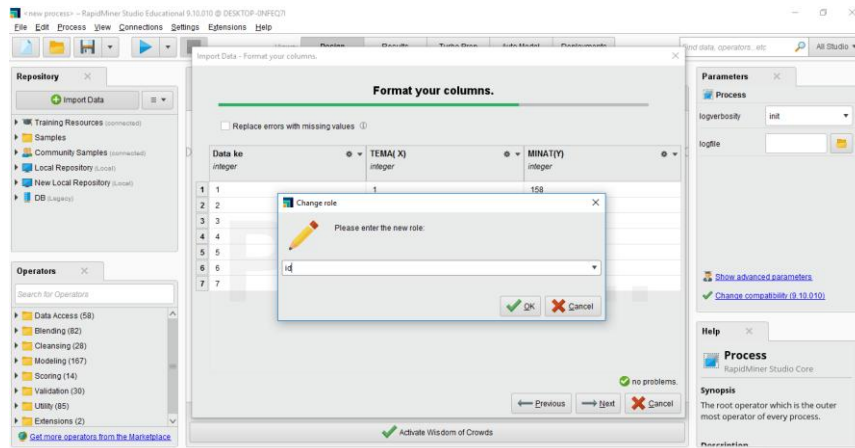
4. Pilih atribut yang ingin di proses.

Pada Gambar 5. pilih atribut yang ingin kita proses pada K-Means yaitu, tema dan minat. Kemudian klik Next.



Gambar 5. Memilih atribut yang ingin di proses

5. Mengubah format kolom.

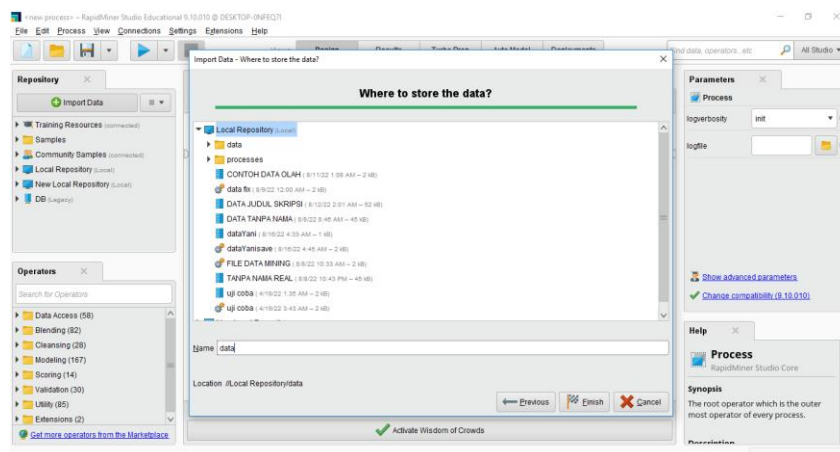


Gambar 6. Format kolom

Tahap selanjutnya yaitu mengubah format kolom atribut menjadi Id, kemudian klik *Next*.

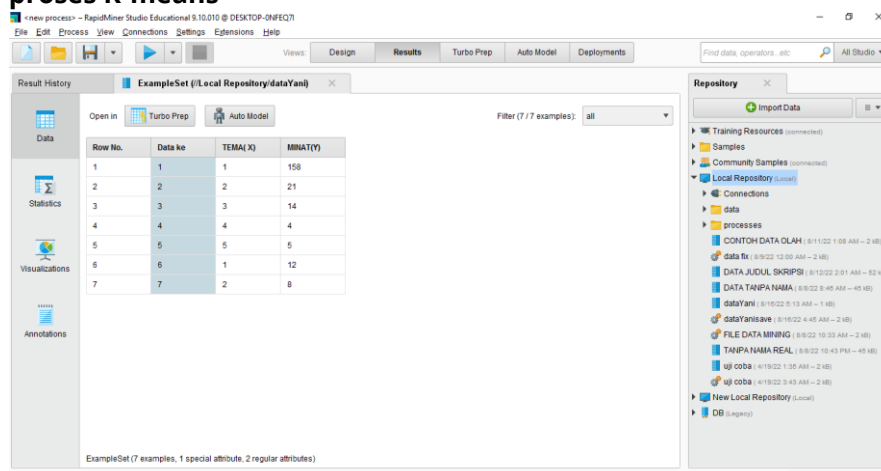
6. Memberi nama database

Pada Gambar 7. menampilkan lokasi penyimpanan pada Rapidminer. Pilih Local Repository/data, lalu kita simpan dengan nama Data Transaksi Burger. Kemudian klik *Finish*.



Gambar 7. Memberi nama data yang baru ditambahkan

7. Memulai proses K-means

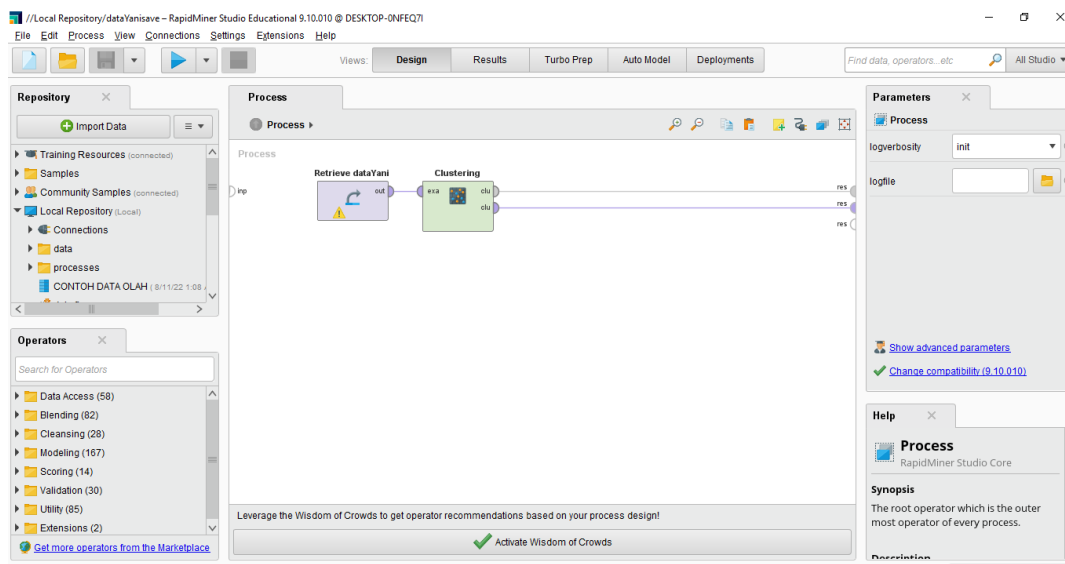


Gambar 8. Memulai proses K-Means

Pada Gambar 9. menampilkan hasil isi *database* yang telah tersimpan pada Rapidminer. Untuk memulai proses K-Means, pilih menu pada tab *Repository* lalu buka *amples/processes/07_Clustering/01_KMeans*.

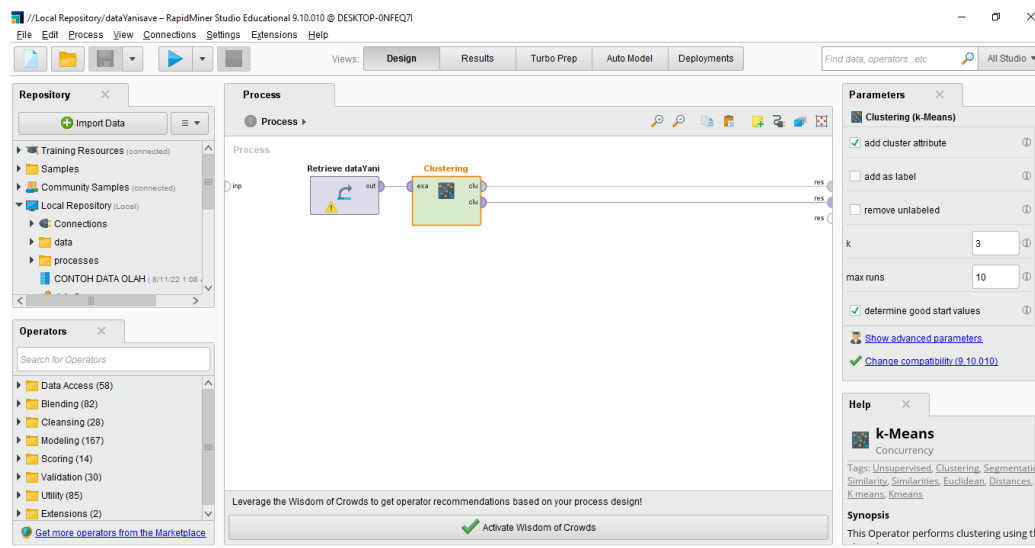
8. Koneksikan database ke dalam proses K-Means

Setelah muncul tampilan seperti pada Gambar 10, kita ubah parameter *Retrieve Data* dengan *database* yang telah di tambahkan. Yaitu dengan menggeser *database* pada *Local Repository/data/Data Transaksi Burger* ke dalam *Retrive Data*.



Gambar 10. Koneksikan database ke dalam proses K-Means

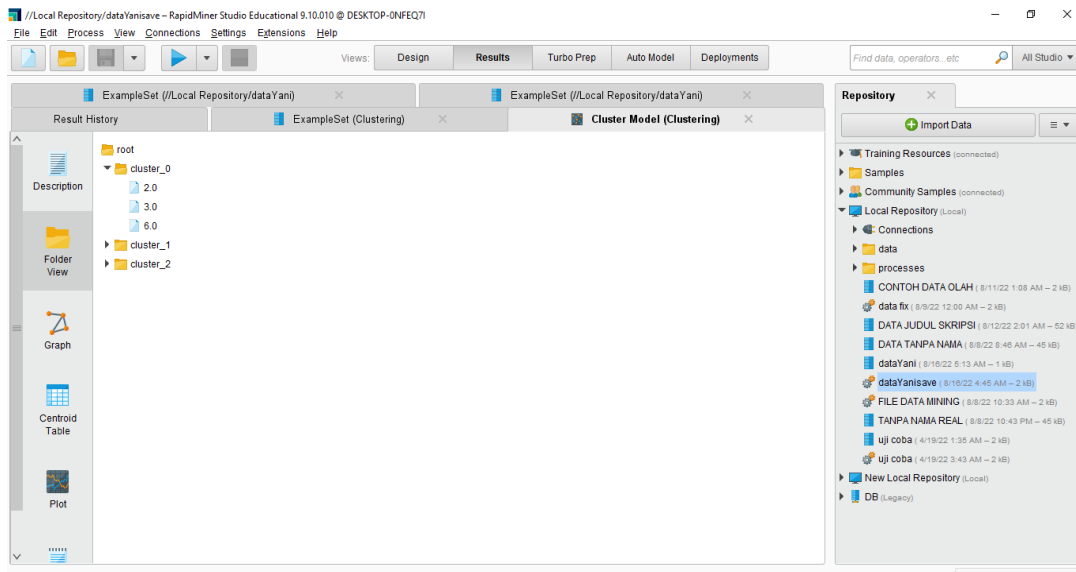
9. Ubah Parameter K-Means



Gambar 11. Mengubah parameter K-Means

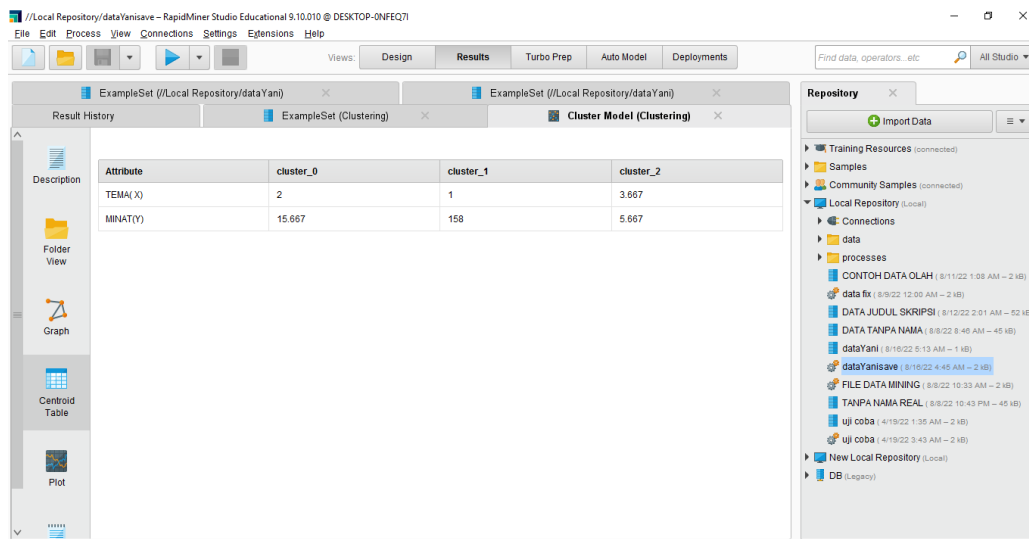
Selanjutnya pada Gambar 11. ubah parameter *k* menjadi 3 (jumlah *cluster* yg di inginkan), *measure type* menjadi *NumericalMeasures* dan *numerical measure* menjadi *EulideanDistance* (proses perhitungan yang diinginkan). Kemudian klik tombol *Play* untuk memulai proses K-Means clustering.

10. Hasil Clustering

Gambar 12. Hasil *clustering*

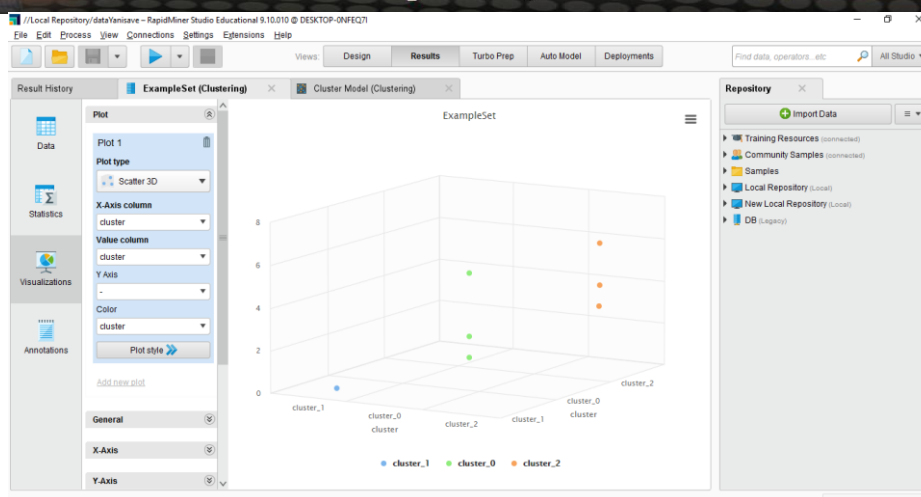
Setelah di proses maka akan muncul tampilan seperti pada Gambar 12. ketika di klik pada tampilan *Folder View* maka akan muncul jumlah kode yang masuk pada setiap *cluster*.

11. Hasil Centroid Akhir



Gambar 13. Hasil centroid akhir

Ketika di klik pada tampilan *Centroid table* seperti pada Gambar 13, maka akan muncul jumlah nilai *centroid* akhir pada setiap *cluster*.



Gambar 14. Hasil Pengolahan Data Pada Rapidminer

KESIMPULAN DAN SARAN

Kesimpulan

1. Penerapan *clustering* menggunakan algoritma *K-Means* bisa untuk mengelompokkan data judul skripsi sebanyak 7 iterasi, yaitu iterasi 1 belum convergen (belum dianggap optimal), iterasi 2 belum convergen (belum dianggap optimal), iterasi 3 belum convergen (belum dianggap optimal), iterasi 4 belum convergen (belum dianggap optimal), iterasi 5 belum convergen (belum dianggap optimal), iterasi 6 belum convergen (belum dianggap optimal) dan iterasi 7 convergen (sudah dianggap optimal) .
2. Hasil penghitungan *K-Means clustering* dapat diterapkan dengan melakukan pendekatan jarak terdekat masing-masing data terhadap *centroid cluster* lain.

Saran

1. Penelitian agar menggunakan data sample jumlah agar menghasilkan cluster model yang lebih bervariasi yang data skripsi memungkinkan muncul kasus baru dalam proses tahapan pembentukan cluster baru. Agar dapat mengembangkan aplikasi yang dapat mengolah data hasil clustering untuk menentukan jarak cluster terhadap centroid cluster lain.
2. Perlu dilakukan penelitian relevan menggunakan metode lain yang lebih baik sebagai pembandingan data yang dihasilkan pada masing-masing metode penelitian..

DAFTAR PUSTAKA

- [1] Pramesti, H. B. KLASERISASI DATA UNSUPERVISED MENGGUNAKAN METODE K-MEANS.
- [2] Darmi, Y. D., & Setiawan, A. (2016). Penerapan metode clustering k-means dalam pengelompokan penjualan produk. *Jurnal Media Infotama*, 12(2).
- [3] Saputra, A., Mulyawan, B., & Sutrisno, T. (2019). Rekomendasi Lokasi Wisata Kuliner di Jakarta Menggunakan Metode K-Means Clustering dan Simple Additive Weighting. *Jurnal Ilmu Komputer dan Sistem Informasi*, 7(1), 14-21.
- [4] Metisen, BM, & Sari, HL (2015). Analisis clustering menggunakan metode K-Means dalam pengelompokan penjualan produk pada Swalayan Fadhila. *Jurnal media infotama* , 11 (2).
- [5] Iryani, L. (2020). Penerapan Datamining Menentukan Minat Baca Mahasiswa Di Perpustakaan Universitas Bina Darma Palembang Menggunakan Metode Clustering. *INTECOMS: Journal of Information Technology and Computer Science*, 3(1), 82-89.

- [6] Fuadi, W., Razi, A., & Fariadi, D. (2022). Automasi Penentuan Tren Topik Skripsi Menggunakan Algoritma K-Means Clustering. *Jurnal Serambi Engineering*, 7(2).
- [7] Riyadhi, M. F. (2019). *Aplikasi Text Mining Untuk Automasi Penentuan Tren Topik Skripsi Dengan Metode K-Means Clustering (Studi Kasus: Prodi Sistem Komputer)* (Doctoral dissertation, Universitas Komputer Indonesia).
- [8] Ramadhani, R. D. (2013). Data Mining Menggunakan Algoritma K-Means Clustering Untuk Menentukan Strategi Promosi Universitas Dian Nuswantoro. *vol, 1*, 1-9.
- [9] Chasanah, T. T., & Widiyono, W. (2017). Penentuan Strategi Promosi Penerimaan Mahasiswa Baru dengan Algoritma Clustering K-Means. *IC-Tech*, 12(1).
- [10] Syahputra, S., Ramadani, S., & Pardede, A. M. H. (2020). Menentukan Strategi Promosi Menggunakan Algoritma Clustering K-Means. *JOISIE (Journal Of Information Systems And Informatics Engineering)*, 4(1), 7-14.
- [11] Nofitri, R., & Irawati, N. (2019). Analisis Data Hasil Keuntungan Menggunakan Software Rapidminer. *JURTEKSI (Jurnal Teknologi dan Sistem Informasi)*, 5 (2), 199-204.