

Analisis Pola Kinerja Anak dalam Tes Membaca untuk Mengidentifikasi Anak yang Membutuhkan Pendampingan Dini Menggunakan Algoritma *K-Means Clustering* di PAUD Seroja

Frencis Matheos Sarimole ¹, Muhamad Aqil Septiansyah ^{2*}

^{1,2*} Program Studi Teknik Informatika, Sekolah Tinggi Ilmu Komputer Cipta Karya Informatika, Kota Jakarta Timur, Daerah Khusus Ibukota Jakarta, Indonesia.

Email: matheosfrancis.s@gmail.com ¹, aqilseptiansyah@gmail.com ^{2*}

Histori Artikel:

Dikirim 27 Juli 2024; *Diterima dalam bentuk revisi* 10 Agustus 2024; *Diterima* 20 Agustus 2024; *Diterbitkan* 20 September 2024. Semua hak dilindungi oleh Lembaga Penelitian dan Pengabdian Masyarakat (LPPM) STMIK Indonesia Banda Aceh.

Abstrak

Penelitian ini bertujuan untuk menganalisis pola kinerja anak dalam tes membaca guna mengidentifikasi anak yang membutuhkan pendampingan dini di PAUD (Pendidikan Anak Usia Dini) Seroja. Identifikasi dini sangat penting untuk memberikan bantuan tepat waktu kepada anak-anak yang mengalami kesulitan membaca, sehingga dapat meningkatkan kemampuan membaca mereka sejak usia dini. Dalam penelitian ini, digunakan algoritma K-Means Clustering untuk mengelompokkan anak-anak berdasarkan hasil tes membaca mereka. Data yang digunakan dalam penelitian ini terdiri dari hasil tes membaca yang diambil dari sejumlah anak di PAUD Seroja. Algoritma K-Means Clustering diterapkan untuk mengelompokkan anak-anak ke dalam beberapa kelompok berdasarkan kinerja mereka. Hasil dari pengelompokan ini kemudian dianalisis untuk mengidentifikasi pola-pola kinerja yang signifikan dan untuk mengenali anak-anak yang memerlukan pendampingan dini.

Kata Kunci: Kinerja Membaca; Pendampingan Dini; Algoritma K-Means Clustering; PAUD; Pendidikan Anak Usia Dini.

Abstract

This research aims to analyze children's performance patterns in reading tests in order to identify children who need early assistance at PAUD (Early Childhood Education) Seroja. Early identification is very important to provide timely assistance to children who have reading difficulties, so as to improve their reading abilities from an early age. In this research, the K-Means Clustering algorithm was used to group children based on their reading test results. The data used in this research consisted of reading test results taken from a number of children at PAUD Seroja. K-Means Clustering algorithm is applied to Cluster children into groups based on their performance. The results of this grouping are then analyzed to identify significant performance patterns and to identify children who need early assistance.

Keyword: Reading Performance; Early Mentoring; K-Means Clustering Algorithm; PAUD; Early Childhood Education.

1. Pendahuluan

Pendidikan Anak Usia Dini (PAUD) memiliki peran penting dalam mengembangkan kemampuan dasar anak, terutama dalam hal keterampilan membaca. Membaca merupakan fondasi utama dalam proses pembelajaran di berbagai tingkatan. Keterampilan ini tidak hanya mempengaruhi aspek akademis, tetapi juga berdampak signifikan pada perkembangan sosial dan emosional anak. Anak-anak yang memiliki kemampuan membaca yang baik sejak dini akan lebih mampu mengikuti pembelajaran di tingkat yang lebih tinggi. Namun, tidak semua anak dapat menguasai keterampilan ini dengan cepat dan mudah, sehingga beberapa anak memerlukan bantuan tambahan untuk mencapai kemampuan membaca yang optimal. Identifikasi dini terhadap anak-anak yang mengalami kesulitan membaca sangatlah penting agar intervensi yang diperlukan dapat diberikan secara tepat waktu. Deteksi awal terhadap anak yang mengalami kesulitan ini membantu mencegah masalah akademis di masa depan dan mendukung perkembangan kemampuan literasi mereka. Styaningtyas *et al.* (2022) menyatakan bahwa penerapan metode *data mining* seperti algoritma *k-means Clustering* dapat digunakan untuk mengelompokkan individu berdasarkan pola yang ditemukan pada data yang tersedia. Algoritma ini sangat efektif dalam mengidentifikasi individu yang membutuhkan bantuan berdasarkan analisis data.

Algoritma *k-means Clustering* merupakan salah satu metode analisis data yang umum digunakan untuk mengelompokkan data ke dalam beberapa kluster berdasarkan kemiripan atribut tertentu. Situmorang *et al.* (2022) menunjukkan bahwa algoritma ini dapat diaplikasikan dalam berbagai bidang, termasuk mengelompokkan rekomendasi kategori baru pada *IMDb*. Dalam pendidikan, penerapan *k-means Clustering* dapat membantu menganalisis kinerja anak dalam membaca dan memberikan gambaran yang lebih jelas mengenai kebutuhan mereka akan bantuan tambahan. Pada PAUD, algoritma ini dapat digunakan untuk mengidentifikasi anak-anak yang memiliki kinerja membaca yang kurang memadai. Rizki *et al.* (2021) mengemukakan bahwa *k-means Clustering* mampu mengelompokkan individu berdasarkan kemampuan tertentu, seperti kemampuan membaca, yang memudahkan pendidik dalam menentukan anak-anak yang memerlukan perhatian lebih. Dengan menggunakan metode ini, data kinerja membaca anak dapat dipetakan secara objektif, memungkinkan pemberian intervensi yang lebih tepat.

Penelitian di PAUD Seroja bertujuan untuk memperbaiki proses identifikasi anak yang membutuhkan pendampingan lebih lanjut dalam hal membaca. Penggunaan algoritma *k-means Clustering* dalam penelitian ini diharapkan dapat membantu guru dalam mengelompokkan anak-anak berdasarkan kinerja membaca mereka. Sebagai tambahan, penelitian oleh Fahmi *et al.* (2022) menunjukkan bahwa klasterisasi menggunakan *k-means* dapat memberikan hasil yang bermanfaat dalam memahami kebiasaan membaca. Oleh karena itu, penelitian ini dapat diterapkan pada PAUD untuk membantu guru dalam memetakan kemampuan membaca anak-anak dan memberikan intervensi yang sesuai.

Hasil penelitian yang dilakukan oleh Priyatman *et al.* (2021) yang menggunakan *k-means Clustering* untuk memprediksi waktu kelulusan mahasiswa menunjukkan bahwa algoritma ini efektif dalam memetakan individu berdasarkan berbagai indikator akademis. Berdasarkan hal tersebut, penerapan algoritma *k-means Clustering* dalam menganalisis kemampuan membaca anak di PAUD Seroja dapat membantu mengidentifikasi anak-anak yang memerlukan pendampingan khusus agar proses pembelajaran mereka lebih optimal. Penelitian ini akan dilakukan di PAUD Seroja untuk menganalisis pola kinerja anak-anak dalam tes membaca, serta mengidentifikasi mereka yang memerlukan pendampingan lebih lanjut. Dengan penerapan algoritma *k-means Clustering*, diharapkan hasil penelitian ini dapat memberikan kontribusi penting bagi pengembangan strategi pendidikan yang lebih efektif di PAUD Seroja maupun di institusi pendidikan anak usia dini lainnya.

2. Metode Penelitian

Metode penelitian merupakan bagian yang sangat penting dalam pelaksanaan riset karena metode inilah yang akan menentukan bagaimana data dikumpulkan, dianalisis, dan diinterpretasikan. Dalam penelitian ini, berbagai pendekatan digunakan untuk memastikan data yang diperoleh relevan dan sesuai dengan tujuan penelitian, yaitu mengidentifikasi anak-anak yang membutuhkan pendampingan dini dalam pembelajaran membaca di PAUD Seroja.

2.1 Data Penelitian

Data penelitian yang digunakan dalam studi ini terdiri dari data primer dan data sekunder. Menurut *Kamus Besar Bahasa Indonesia (KBBI)*, data adalah keterangan atau bahan nyata yang bisa dijadikan dasar kajian atau analisis dalam suatu penelitian. Data dalam penelitian ini merujuk pada informasi yang dikumpulkan untuk mendukung proses analisis, yang dapat berupa fakta, angka, statistik, pernyataan, observasi, atau informasi relevan lainnya terkait topik yang diteliti. Data dapat diperoleh melalui berbagai metode, seperti survei, eksperimen, wawancara, observasi, analisis literatur, serta pengumpulan data sekunder.

2.2 Data Primer

Data primer dalam penelitian ini diperoleh langsung dari lapangan melalui pengamatan dan wawancara yang dilakukan secara langsung di PAUD Seroja. Data ini diambil dari sumber asli yang merupakan responden atau objek penelitian, yaitu siswa PAUD yang menjadi subjek dalam penelitian ini. Metode pengumpulan data primer yang digunakan meliputi wawancara, observasi langsung, dan pengisian kuesioner oleh responden terkait. Wawancara dilakukan dengan guru dan staf sekolah guna mendapatkan informasi rinci tentang perilaku dan kemajuan siswa dalam tes membaca. Selain itu, observasi langsung di kelas juga dilakukan untuk melihat secara langsung aktivitas siswa selama proses pembelajaran membaca berlangsung. Data yang diperoleh dari pengamatan ini kemudian dianalisis untuk melihat pola-pola yang bisa digunakan sebagai dasar dalam mengidentifikasi anak-anak yang membutuhkan pendampingan.

2.3 Data Sekunder

Selain data primer, data sekunder juga digunakan untuk memperkaya informasi yang diperoleh. Data sekunder adalah data yang sudah ada sebelumnya, biasanya diperoleh dari berbagai sumber lain, seperti laporan penelitian sebelumnya, artikel ilmiah, atau basis data publik yang relevan dengan penelitian ini. Pengumpulan data sekunder dilakukan melalui studi pustaka, yang melibatkan penelaahan literatur dari buku-buku dan jurnal yang membahas tentang algoritma *k-means Clustering* dan penerapannya dalam dunia pendidikan, khususnya dalam mendeteksi pola kinerja anak dalam pembelajaran membaca. Data sekunder ini berperan penting dalam memberikan landasan teoritis dan panduan bagi penelitian, sehingga hasil penelitian ini dapat dikaitkan dengan studi-studi sebelumnya dan dapat dipertanggungjawabkan secara akademis.

2.4 Teknik Pengumpulan Data

1) Observasi

Observasi merupakan salah satu metode utama yang digunakan dalam penelitian ini. Pengamatan langsung dilakukan di PAUD Seroja untuk memperoleh informasi empiris tentang aktivitas siswa selama proses pembelajaran berlangsung. Observasi ini dilakukan secara teliti dan sistematis untuk mencatat kinerja siswa dalam tes membaca. Aktivitas siswa diamati dan dicatat berdasarkan berbagai indikator yang telah ditetapkan, seperti kemampuan mengenali huruf dan kata, kecepatan membaca, ketepatan membaca, serta kemampuan memahami teks. Dengan menggunakan observasi langsung, peneliti dapat memperoleh gambaran yang lebih jelas tentang kinerja setiap siswa dalam tes membaca dan mengidentifikasi siswa yang membutuhkan intervensi lebih lanjut.

2) Wawancara

Selain observasi, wawancara juga dilakukan dengan guru dan staf sekolah di PAUD Seroja. Wawancara ini bertujuan untuk mendapatkan informasi yang lebih rinci mengenai kondisi dan perkembangan siswa, serta praktik pembelajaran yang diterapkan di PAUD Seroja. Dengan melakukan wawancara, peneliti bisa menggali informasi tambahan terkait perilaku dan kemajuan siswa dalam belajar membaca. Guru juga memberikan pandangan mereka tentang siswa yang membutuhkan pendampingan lebih lanjut, yang selanjutnya digunakan sebagai bahan pertimbangan dalam analisis data.

2.5 Analisis Data

Setelah data terkumpul, proses analisis dilakukan dengan menggunakan metode statistik yang sesuai. Data yang diperoleh dianalisis menggunakan algoritma *k-means Clustering*, yang digunakan untuk mengelompokkan siswa berdasarkan hasil tes membaca mereka. Algoritma ini membantu mengidentifikasi siswa yang membutuhkan pendampingan lebih lanjut dalam belajar membaca. Analisis statistik deskriptif juga digunakan untuk menggambarkan distribusi frekuensi variabel-variabel seperti kecepatan membaca, ketepatan membaca, kemampuan mengenali huruf dan kata, pemahaman teks, serta kemampuan memori visual.

3. Hasil dan Pembahasan

3.1 Hasil

3.1.1 Alat Penelitian

Dalam penelitian, membutuhkan alat untuk mendukung berjalannya penelitian. Alat penelitian yang digunakan yaitu berupa perangkat lunak (*software*) dan perangkat keras (*hardware*). Aplikasi Rapidminer untuk Pengolahan data dan menggunakan aplikasi excel dengan metode elbow method untuk menentukan nilai K berapa *Cluster* yang paling bagus untuk data yang akan diolah.

Tabel 1. Spesifikasi Perangkat lunak (*Software*)

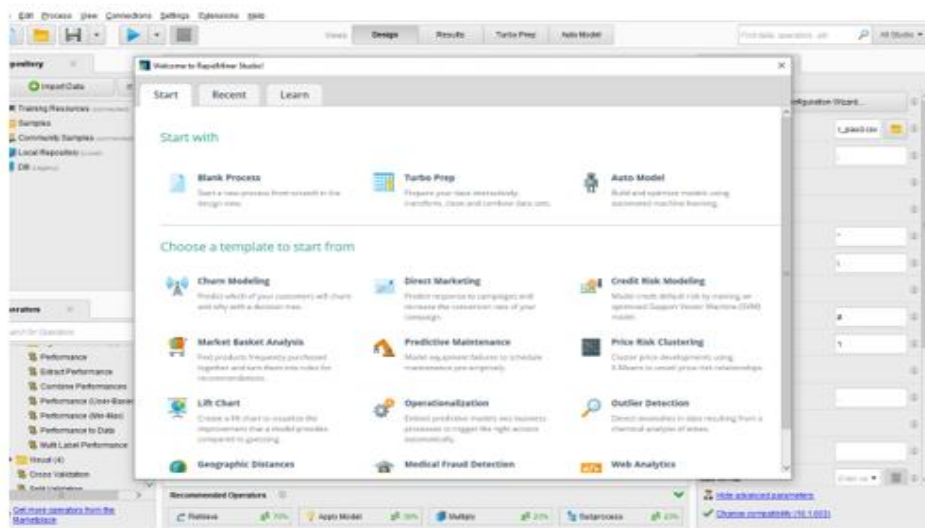
No	<i>Software</i>	Versi	Kegunaan
1	Rapid Miner	10.1	Untuk Melakukan Analisis Data Mining
2	Excel 2019	2406	Untuk Memfilter Data Dan Menyimpan Data

Tabel 2. Spesifikasi Perangkat Keras (*Hardware*)

No	Hardware	Spesifikasi
1	Processor	Intel Core i3-2100 3. 10Ghz
2	RAM	16 GB
3	Operations System	Windows 10 Pro 64 Bit
4	VGA	Radeon R7 200 Series 2 Gb
5	Monitor	Spc 24 Inc
6	Motherboard	Asrock H61M

3.1.2 Pengujian

Pengujian Menggunakan *Software Rapidminer*, pada Implementasi dan Pengujian disini kita menggunakan sebuah *Software RapidMiner Studio Versi10.1*, dengan pengujian data menggunakan *software RapidMiner*.

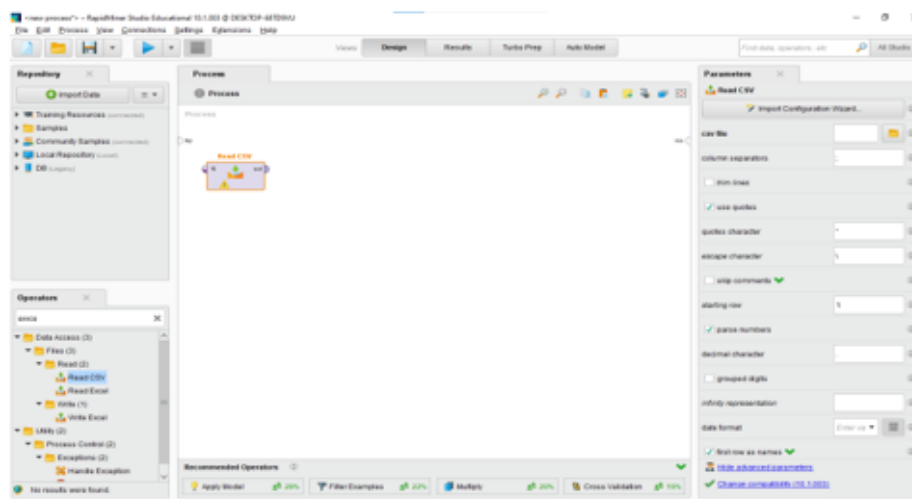


Gambar 1. Pengujian Menggunakan Software Rapidminer

Pada tampilan halaman utama ada lima menu yang akan di gunakan:

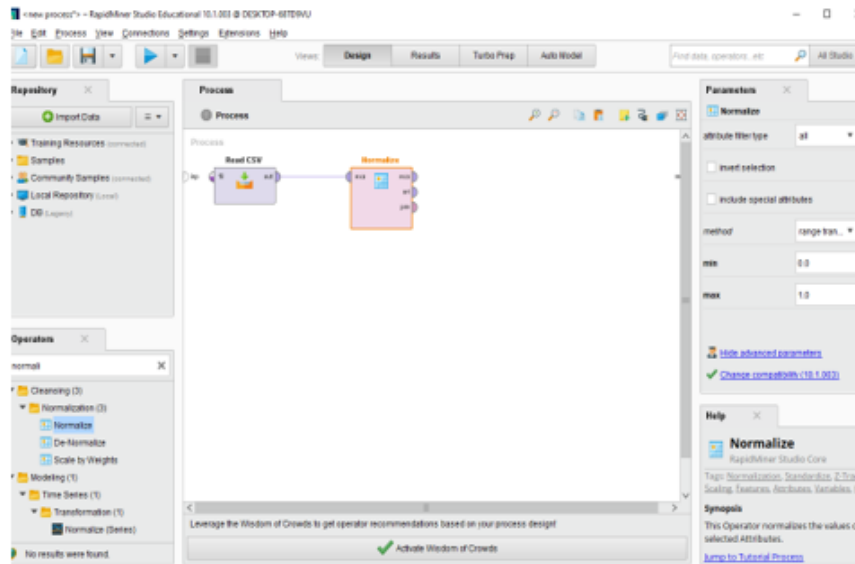
- Start berfungsi sebagai untuk membuat halaman proses kerja pada Data Mining yang baru
- Recent berfungsi untuk membuka proses yang sudah ada direpository sebelumnya.
- Learn berisikan petunjuk-petunjuk menggunakan Rapidminer studio Versi 10.1

Tampilan *new process* adalah untuk membuat halaman kerja pada *RapidMiner Studio v.10.1*. Import data dengan operator Read Csv drag dan drop pada kanvas hal ini dilakukan untuk memasukkan data yang akan diuji dalam bentuk format .xls atau .csv. Berikut adalah cara untuk melakukan *import file* Microsoft Excel yang sudah di ubah dalam bentuk format .csv. Untuk membuat mengimport data yang akan diproses, maka dilakukan *New process*.

Gambar 2. Tampilan *New process*

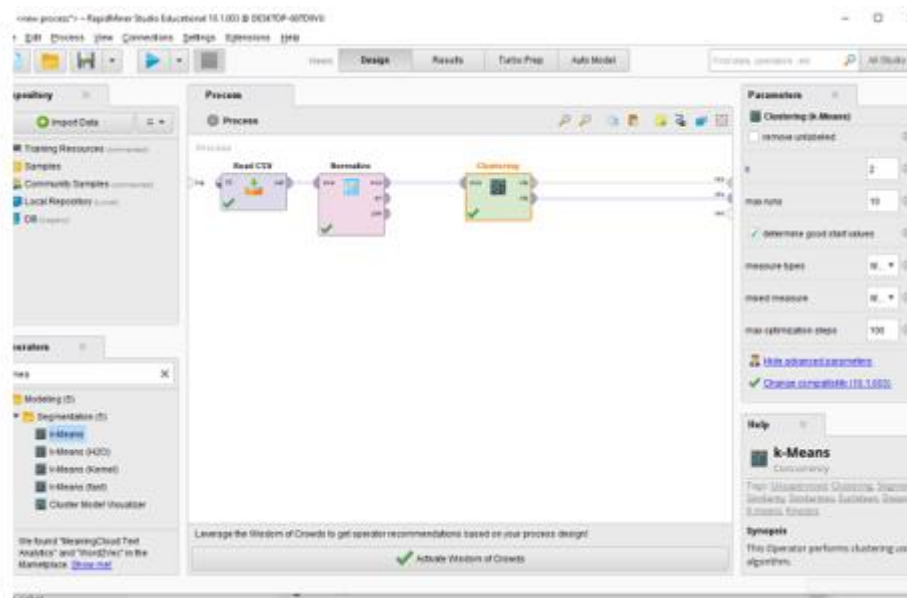
Setelah itu cari file data yang ingin kita masukan pada Read Csv. Setelah file sudah di import maka kita harus menyeleksi terlebih dahulu kolom mana saja yang ingin di tampilkan. Setelah selesai diseleksi, maka selanjutnya yaitu menormalize data tujuannya untuk mengubah skala data sehingga berada dalam rentang tertentu. Langkah awal untuk normalized yaitu mencari operator

normalize pada menu operator lalu di *drag* ke kanvas dan setelah itu klik dua kali pada operator *normalize* dan pada method ubah menjadi *range transformation* dan tentukan rentang minimum dan maximum 0 dan 1, sambungkan Read Csv *output* ke *input* *Normalize*.



Gambar 3. Import Data

Dan berikut masuk ke pengklusteran data menggunakan operator *K- Means*, cari operator *K-means* terlebih dahulu kemudian di drag dan drop pada kanvas dan *output normalized* di sambungkan ke operator *input K-Means*, lalu klik dua kali pada operator *k-Means* untuk mengubah nilai *K* karena hanya ingin menggunakan dua kluster maka nilai *K* nya di ubah menjadi dua kluster.



Gambar 4. Pengklusteran Data

Datas dilakukan pengaturan nilai k , di mana k merupakan nilai yang digunakan untuk menentukan jumlah *Cluster* yang akan dibentuk. Di sini jumlah *Cluster* yang akan dibentuk adalah sebanyak 2 *Cluster* sesuai tingkatan yaitu rendah, dan tinggi.

id	cluster	keceptian	ketepatan m.	kesempurnaan	respon terk.	kesempurnaan	nama siswa
1	cluster_0	0.007	1	1	0.333	0.007	Harini Rina
2	cluster_0	0.333	1	1	0.007	0.007	Huraidi
3	cluster_0	0.007	1	1	0.333	0.007	Zamar Mudo
4	cluster_1	0.333	0	1	0.333	0.007	Muhammad Ra
5	cluster_0	0.007	1	1	0.007	0.007	Danu Umbara
6	cluster_1	0.007	0	1	0.333	0.007	Rully Chandra
7	cluster_0	0.007	1	1	0.333	0.007	Harani
8	cluster_1	0.007	0	1	0.007	0.007	Mahmudi
9	cluster_0	0.007	1	1	0.333	0.007	Faham
10	cluster_1	0.007	0	1	0.007	0.007	Itan
11	cluster_0	0.007	1	1	0.333	0.007	Fito Triani
12	cluster_1	0.007	1	1	0.333	0.007	Agha Faniha
13	cluster_1	0.007	0	1	0.007	0.007	Agung Dharma
14	cluster_0	1	1	1	0.007	0.007	San Sumia
15	cluster_0	0.007	1	1	0.007	0.007	Berka Kama
16	cluster_1	0.007	0	1	0.007	0.007	Kusala
17	cluster_0	0.333	1	1	0.007	0.007	Wahyu Harko
18	cluster_0	1	1	1	0.007	0.007	Sangar Hadi
19	cluster_1	1	1	1	0.333	0.007	Egi Misa Pua
20	cluster_1	0.007	0	1	0.333	0.007	Inda Harahap

Gambar 5. Example set

Pada bagian *Example Set* terdapat beberapa bentuk hasil *Cluster* yaitu: *Data View*, *Statistics View*, *visualizations*, dan *Annotations*. Untuk *Cluster model* terdapat *Description*, *Folder View*, *Graph View*, *Centroid Table*, *Centroid Plot View* dan *Annotations*.

a) *Example Set*

Pada *Example Set* dapat dilihat beberapa tampilan hasil *Cluster*, yaitu *Data View*. *Data View* merupakan tampilan hasil *Cluster* data secara keseluruhan sesuai dengan data yang telah diinputkan

b) *Visualizations*

Merupakan tampilan grafik hasil pengelompokan atau *Cluster* sampel data Paud Seroja dengan 2 *Cluster*.

Cluster Model	
Cluster 0:	326 items
Cluster 1:	45 items
Total number of items:	371

Gambar 6. Cluster Model

Pada penelitian ini diperoleh hasil 2 *Cluster*. *Cluster* (0) merupakan *Cluster* yang tinggi dan *Cluster* (1) merupakan *Cluster* yang rendah. *Cluster* (0) memiliki data 327 dan *Cluster* (1) memiliki data 45. *Cluster* (0) memiliki angka yang tinggi dimana angka itu menunjukan anak-anak yang tidak Membutuhkan pendampingan lebih dimana anak-anak pada *Cluster* 0 ini adalah anak-anak yang memiliki Tingkat kemampuan yang sama rata dengan indikator yang di nilai oleh guru, dan pada *Cluster* 1 dimana *Cluster* ini dimiliki oleh anak-anak yang Membutuhkan pendampingan lebih dari guru dimana anak-anak pada *Cluster* 1 ini adalah anak-anak yang memiliki nilai di salah satu indikator ada yang kurang dari anak yang berada di *Cluster* 0.

3.2 Pembahasan

Penelitian ini menggunakan algoritma *k-means Clustering* untuk mengelompokkan siswa PAUD Seroja berdasarkan kinerja membaca mereka, dengan tujuan utama mengidentifikasi siswa yang membutuhkan pendampingan lebih lanjut. Hasil analisis menunjukkan dua kluster utama, yaitu siswa dengan kinerja membaca yang baik dan siswa yang membutuhkan pendampingan dalam membaca. Hasil ini memberikan gambaran yang jelas mengenai perbedaan kinerja siswa dan menjadi dasar untuk intervensi pendidikan yang lebih tepat sasaran. Penggunaan algoritma *k-means Clustering* dalam penelitian ini didukung oleh literatur sebelumnya. Menurut Styaningtyas *et al.* (2022), *k-means Clustering* adalah salah satu teknik *data mining* yang populer dan efektif dalam mengelompokkan data ke dalam kluster berdasarkan kemiripan atribut. Dalam penelitian ini, atribut yang digunakan meliputi kecepatan membaca, ketepatan dalam mengenali huruf dan kata, serta pemahaman teks. Penggunaan algoritma ini memungkinkan pendidik untuk mengelompokkan siswa dengan cara yang obyektif, sehingga dapat memberikan intervensi yang lebih terarah berdasarkan kebutuhan individu.

Selain itu, penerapan *k-means Clustering* dalam pendidikan juga telah dibahas oleh beberapa studi. Situmorang *et al.* (2022) menunjukkan bahwa algoritma ini efektif dalam mengelompokkan data yang memiliki kesamaan karakteristik, seperti yang diterapkan dalam analisis rekomendasi kategori baru di *IMDb*. Dalam konteks penelitian ini, algoritma *k-means Clustering* mampu membedakan siswa yang menunjukkan kinerja membaca yang baik dari mereka yang membutuhkan pendampingan tambahan, memberikan informasi yang sangat berguna bagi guru untuk menyesuaikan strategi pengajaran mereka. Temuan ini juga sejalan dengan penelitian yang dilakukan oleh Rizki *et al.* (2021), di mana *k-means Clustering* digunakan untuk mengelompokkan minat membaca penduduk. Dalam penelitian ini, algoritma tersebut berhasil mengelompokkan individu berdasarkan kemampuan literasi, yang kemudian dapat dijadikan dasar untuk intervensi yang lebih efektif. Pada PAUD Seroja, algoritma ini terbukti mampu mengidentifikasi siswa yang membutuhkan dukungan tambahan dalam pembelajaran membaca, memungkinkan guru untuk memberikan perhatian yang lebih personal kepada siswa-siswa tersebut.

Penelitian Fahmi *et al.* (2022) juga memperkuat hasil penelitian ini. Fahmi *et al.* menggunakan *k-means Clustering* untuk mengelompokkan data kegemaran membaca siswa di SMA Al-Islam Cirebon. Hasil penelitian mereka menunjukkan bahwa algoritma ini dapat dengan jelas memisahkan kelompok siswa berdasarkan kegemaran membaca mereka. Dalam konteks PAUD, penggunaan *k-means Clustering* serupa dapat membantu guru dalam mengidentifikasi siswa yang memiliki kesulitan membaca sehingga intervensi dapat dirancang dengan lebih akurat. Priyatman *et al.* (2021) juga menemukan bahwa *k-means Clustering* efektif dalam memprediksi waktu kelulusan mahasiswa berdasarkan berbagai indikator akademik. Temuan mereka relevan dengan hasil penelitian ini, yang menggunakan algoritma yang sama untuk memprediksi siswa mana yang memerlukan pendampingan dini dalam pembelajaran membaca. Penggunaan algoritma ini memberikan pandangan yang lebih obyektif tentang kemampuan siswa dan membantu dalam menentukan langkah intervensi yang tepat untuk meningkatkan kemampuan literasi mereka.

Hasil penelitian ini memiliki beberapa implikasi penting dalam pembelajaran di PAUD Seroja. Pertama, hasil pengelompokan ini memberikan pedoman bagi guru untuk lebih memahami perbedaan kinerja siswa dalam membaca. Siswa yang tergolong dalam kluster dengan kinerja rendah dapat segera mendapatkan intervensi yang lebih spesifik, seperti pelatihan tambahan atau pendekatan belajar yang

lebih intensif, yang disesuaikan dengan kelemahan individu masing-masing siswa. Hal ini dapat meningkatkan efektivitas pembelajaran di kelas dan memastikan bahwa setiap siswa mendapatkan bantuan yang sesuai dengan kebutuhannya. Kedua, penggunaan teknologi dalam proses analisis kinerja siswa membuka peluang baru dalam pendidikan. Algoritma *k-means Clustering* terbukti dapat memberikan analisis obyektif terhadap kinerja siswa dan menghasilkan data yang dapat diandalkan untuk mendukung pengambilan keputusan di tingkat kelas dan sekolah. Oleh karena itu, penerapan algoritma ini tidak hanya relevan bagi PAUD Seroja, tetapi juga memiliki potensi untuk diimplementasikan di lembaga pendidikan lain yang menghadapi tantangan serupa dalam mengidentifikasi siswa yang membutuhkan pendampingan dini. Selain itu, hasil penelitian ini juga menekankan pentingnya identifikasi dini bagi siswa yang mengalami kesulitan dalam membaca. Dengan memanfaatkan hasil pengelompokan yang dihasilkan oleh algoritma *k-means Clustering*, guru dapat lebih cepat dan tepat dalam memberikan bantuan kepada siswa yang membutuhkan. Hal ini sesuai dengan temuan Priyatman *et al.* (2021), yang menekankan bahwa identifikasi dini dengan algoritma *k-means Clustering* dapat meningkatkan efektivitas intervensi pendidikan, khususnya dalam prediksi dan perbaikan performa akademik.

Penelitian ini berhasil menunjukkan efektivitas algoritma *k-means Clustering* dalam mengelompokkan siswa PAUD berdasarkan kinerja membaca mereka. Pengelompokan ini membantu guru dalam merancang strategi intervensi yang lebih sesuai dengan kebutuhan siswa, sehingga dapat meningkatkan kemampuan literasi mereka secara lebih efektif. Hasil ini didukung oleh penelitian sebelumnya, seperti yang dilakukan oleh Styaningtyas *et al.* (2022), Situmorang *et al.* (2022), Rizki *et al.* (2021), Fahmi *et al.* (2022), dan Priyatman *et al.* (2021), yang menunjukkan bahwa *k-means Clustering* adalah metode yang andal untuk mengidentifikasi pola kinerja dalam berbagai konteks pendidikan.

4. Kesimpulan

Berdasarkan hasil analisis dan pengolahan data yang telah dilakukan, beberapa kesimpulan dapat diambil. Penelitian ini menunjukkan bahwa penerapan teknik *data mining* dengan menggunakan metode *k-means Clustering* dapat memberikan informasi yang berguna dalam mengidentifikasi anak-anak yang membutuhkan pendampingan dalam tes membaca. Metode ini memungkinkan pengelompokan siswa secara lebih obyektif, sehingga pendidik dapat dengan mudah mengenali siswa yang memerlukan intervensi dini. Algoritma *k-means Clustering* berhasil mengelompokkan anak-anak menjadi dua kelompok utama. Kelompok pertama terdiri dari siswa dengan kinerja membaca yang sangat baik, di mana mereka tidak memerlukan pendampingan tambahan. Sementara itu, kelompok kedua terdiri dari siswa dengan kinerja membaca yang rendah, sehingga mereka memerlukan pendampingan lebih lanjut untuk meningkatkan kemampuan mereka dalam membaca.

Hasil lebih rinci dari penelitian ini mengungkapkan bahwa dari 372 data siswa yang dianalisis, dua kluster berhasil diidentifikasi. Kluster pertama, yang dikenal sebagai kluster 0, terdiri dari 327 siswa yang tidak membutuhkan pendampingan tambahan dalam membaca. Sementara itu, kluster kedua, atau kluster 1, terdiri dari 45 siswa yang membutuhkan intervensi lebih lanjut untuk memperbaiki kemampuan membaca mereka. Temuan ini menunjukkan efektivitas algoritma *k-means Clustering* dalam mengelompokkan siswa berdasarkan kinerja mereka, sehingga pendidik dapat memberikan perhatian yang lebih tepat dan efektif kepada siswa yang memerlukan bantuan tambahan.

5. Ucapan Terima Kasih

Dengan penuh rasa syukur dan terima kasih kepada keluarga yang selalu memberikan support dan saya mahasiswa dari Sekolah Tinggi Ilmu Komputer Cipta Karya Informatika, ingin menyampaikan ucapan terima kasih yang tulus kepada Bapak Francis Matheos Sarimole atas

bimbingan, arahan, dan inspirasi yang luar biasa selama pelaksanaan skripsi ini. Bapak telah memberikan panduan yang sangat berharga, serta kesabaran dan dukungan yang tak terhingga selama kami menjalankan tugas ini. Saya juga ingin menyampaikan terimakasih kepada Sekolah Tinggi Ilmu Komputer Cipta Karya Informatika atas fasilitas dan dukungan yang diberikan selama saya menjalankan skripsi ini. Semua bantuan dan dukungan ini sungguh berarti bagi kelancaran skripsi saya. Dengan rendah hati, kami menyadari bahwa terselesaikannya skripsi ini tidak mungkin terwujud tanpa kontribusi dari semua.

6. Daftar Pustaka

- Al Masykur, A. (2023). Penerapan Metode K-Means Clustering untuk Pemetaan Pengelompokan Lahan Produksi Tandan Buah Segar. *Penerapan Metode K-Means Clustering untuk Pemetaan Pengelompokan Lahan Produksi Tandan Buah Segar*, 10(1), 92-100.
- Anggraeni, D., Rizaldi, R., & Putra, G. M. (2021). Penerapan K-Means Clustering Untuk Pengelompokan Kelas Pada Taman Kanak-Kanak. *Building of Informatics, Technology and Science (BITS)*, 3(3), 400-404. DOI: <https://doi.org/10.47065/bits.v3i3.1125>.
- Bhatnagar, S., & Saxena, P. S. (2018). ANALYSIS OF FACULTY PERFORMANCE EVALUATION USING CLASSIFICATION. *International Journal of Advanced Research in Computer Science*, 9(1).
- Hossain, M. Z., Akhtar, M. N., Ahmad, R. B., & Rahman, M. (2019). A dynamic K-means clustering for data mining. *Indonesian Journal of Electrical engineering and computer science*, 13(2), 521-526.
- Messakh, G. C. (2023). COMPARISON K-MEANS AND FUZZY C-MEANS IN REGENCIES/CITIES GROUPING BASED ON EDUCATIONAL INDICATORS. *Jurnal Varian*, 7(1).
- Nugraha, G. S., & Hairani, H. (2018). Aplikasi pemetaan kualitas pendidikan di Indonesia menggunakan metode k-means. *MATRIK: Jurnal Manajemen, Teknik Informatika dan Rekayasa Komputer*, 17(2), 13-23. DOI: <https://doi.org/10.30812/matrik.v17i2.84>.
- Pradnyana, G. A., & Permana, A. A. J. (2018). Sistem Pembagian Kelas Kuliah Mahasiswa Dengan Metode K-Means Dan K-Nearest Neighbors Untuk Meningkatkan Kualitas Pembelajaran. *Jurnal Ilmiah Teknologi Informasi*, 16(1), 59-68. DOI: <http://dx.doi.org/10.12962/j24068535.v16i1.a696>.
- Priyatman, H., Sajid, F., & Haldivany, D. (2019). Klasterisasi Menggunakan Algoritma K-Means Clustering untuk Memprediksi Waktu Kelulusan Mahasiswa. *Jurnal Edukasi Dan Penelitian Informatika (JEPIN)*, 5(1), 62.
- Qoiriah, A., Harimurti, R., & Nurhidayat, A. I. (2020, November). Application of k-means algorithm for clustering student's computer programming performance in automatic programming assessment tool. In *International Joint Conference on Science and Engineering (IJCSE 2020)* (pp. 421-425). Atlantis Press. DOI: <https://doi.org/10.2991/aer.k.201124.075>.
- Rauthan, A., Singh, A. S., & Singh, N. (2021, December). Impact on higher education in pandemic: analysis k-means clustering using urban & rural areas. In *2021 3rd International Conference on*

Advances in Computing, Communication Control and Networking (ICAC3N) (pp. 1974-1980). IEEE.
DOI: <https://doi.org/10.1109/ICAC3N53548.2021.9725709>.

Rizki, M. Y., Maysaroh, S., & Windarto, A. P. (2021). Implementasi K-Means Clustering dalam Mengelompokkan Minat Membaca Penduduk Menurut Wilayah. *Just IT: Jurnal Sistem Informasi, Teknologi Informasi dan Komputer*, 11(2), 41-49. DOI: <https://doi.org/10.24853/justit.11.2.41-49>.

Setyaningtyas, S., Nugroho, B. I., & Arif, Z. (2022). TINJAUAN PUSTAKA SISTEMATIS PADA DATA MINING: STUDI KASUS ALGORITMA K-MEANS CLUSTERING. *J Teknoif Teknik Informatika*, 10, 52-61.

Situmorang, A., Arifin, A., Rusilpan, I., & Juliane, C. (2022). Analisa dan Penerapan Metode Algoritma K-Means Clustering Untuk Mengidentifikasi Rekomendasi Kategori Baru Pada List Movie IMDb. *JURNAL MEDIA INFORMATIKA BUDIDARMA*, 6(4), 2171-2179. DOI: <http://dx.doi.org/10.30865/mib.v6i4.4729>.

Stephen, K. W. (2016). Data Mining Model for Predicting Student Enrolment in STEM Courses in Higher Education Institutions.

Zahroh, L. F. A., Rahaningsih, N., & Dana, R. D. (2024). KLASTERISASI DATA KECEMUKURAN MEMBACA MENGGUNAKAN ALGORITMA K-MEANS DI SMA AL-ISLAM CIREBON. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 8(3), 2692-2698. DOI: <https://doi.org/10.36040/jati.v8i3.9543>.