

Terbit online pada laman : <http://teknosi.fti.unand.ac.id/>

Jurnal Nasional Teknologi dan Sistem Informasi

| ISSN (Print) 2460-3465 | ISSN (Online) 2476-8812 |



Artikel Penelitian

Analisis Metode SAW Pada Perangkingan Hasil *Clustering* K-Means Data Evaluasi Kemampuan Dasar Pemrograman Mahasiswa

Made Pasek Agus Ariawan^{a,*}, Ida Bagus Adisimakrisna Peling^b, Gde Brahupadhy Subiksa^c

^{a,b,c}Program Studi Teknologi Rekayasa Perangkat Lunak, Jurusan Teknologi Informasi, Politeknik Negeri Bali, Badung, Bali

INFORMASI ARTIKEL

Sejarah Artikel:

Diterima Redaksi: 21 September 2024

Revisi Akhir: 31 Agustus 2025

Diterbitkan Online: 13 September 2025

KATA KUNCI

K-means,
Clustering, Analisis,
SAW

KORESPONDENSI

E-mail: pasekagus@pnb.ac.id*

A B S T R A C T

Di era digital, volume data yang dihasilkan terus meningkat dengan cepat, sehingga kebutuhan akan metode efektif untuk menganalisis dan menemukan pola dalam data menjadi sangat penting. Salah satu teknik yang sering digunakan untuk mengelompokkan data adalah *clustering*, khususnya metode *K-Means*. Namun, *K-Means* memiliki kelemahan dalam pelabelan *cluster* karena sifatnya yang unsupervised, yang membuat interpretasi hasil menjadi lebih sulit. Untuk mengatasi kelemahan ini, penelitian ini menggabungkan metode Simple Additive Weighting (SAW) dengan *K-Means*, di mana SAW digunakan untuk menilai hasil pengelompokan yang dihasilkan *K-Means* dan memberikan bobot pada berbagai kriteria evaluasi. Penelitian ini dilakukan dalam konteks Program Studi D4 Teknologi Rekayasa Perangkat Lunak guna mengevaluasi kemampuan dasar pemrograman mahasiswa. Dengan menggunakan kombinasi kedua metode ini, diharapkan pengelompokan data menjadi lebih akurat dan optimal, serta dapat membantu proses penilaian hasil *clustering* secara lebih informatif. Hasil penelitian ini menunjukkan bahwa integrasi SAW dapat meningkatkan kualitas pengelompokan dengan membantu menentukan centroid yang lebih tepat dan menilai hasil pengelompokan secara lebih mendalam.

1. PENDAHULUAN

Di era digital, jumlah data yang dihasilkan semakin meningkat dengan cepat, sehingga muncul kebutuhan akan cara yang efektif untuk menganalisis dan menemukan pola dalam data tersebut. Salah satu metode yang sering digunakan adalah pengelompokan data (*clustering*), yang berfungsi untuk mengelompokkan data berdasarkan kesamaan tertentu. Teknik ini sangat berguna di berbagai sektor, mulai dari bisnis untuk membedakan segmen pelanggan, di bidang kesehatan untuk menganalisis data medis, hingga dalam ilmu pengetahuan guna menemukan hubungan tersembunyi dalam kumpulan data yang kompleks. Dengan demikian, pengelompokan data membantu pengambilan keputusan yang lebih cerdas dan berbasis informasi.

Perkembangan ilmu pengetahuan dan teknologi mendorong perguruan tinggi, sebagai lembaga pendidikan formal, untuk mencetak lulusan yang kompeten dan berdaya saing tinggi. Oleh

karena itu, metode pengajaran di perguruan tinggi harus lebih kreatif, inovatif, dan responsif terhadap kebutuhan pasar kerja[1]. Program Studi Sarjana Terapan Teknologi Rekayasa Perangkat Lunak (D4 TRPL) adalah bagian dari Jurusan Teknologi Informasi. Program ini bertujuan untuk menghasilkan lulusan berkualitas di bidang rekayasa perangkat lunak atau *software development*, dengan gelar Sarjana Terapan (S.Tr.Kom.). Oleh karena itu, penguasaan keterampilan pemrograman menjadi syarat utama bagi para mahasiswa dalam program ini..

Penelitian yang dilakukan oleh Yudistira menyebutkan bahwa salah satu faktor yang menentukan kualitas pendidikan dari tinggi dan rendahnya tingkat keberhasilan siswa dalam proses penilaian pembelajaran. Dunia pendidikan pada era digital ini dituntut untuk memiliki kemampuan dalam bersaing dengan memanfaatkan sumber daya yang dimiliki, baik sumber daya sarana, sumber daya prasarana, sampai sumber daya manusia. Dengan memiliki sumber daya yang baik akan menunjang seluruh kegiatan

operasional pembelajaran serta akan menunjang semua kegiatan dalam proses pengambilan keputusan[2].

Kemampuan dasar pemrograman merupakan salah satu kemampuan penting yang harus dimiliki oleh mahasiswa, khususnya mahasiswa yang mengambil jurusan atau program studi yang berkaitan dengan teknologi informasi. Kemampuan dasar pemrograman ini meliputi pemahaman tentang konsep-konsep pemrograman, seperti variabel, tipe data, operator, kontrol aliran, dan fungsi.

Evaluasi kemampuan dasar pemrograman mahasiswa dapat dilakukan dengan berbagai cara, salah satunya adalah dengan menggunakan metode klusterisasi. Metode klusterisasi adalah metode yang digunakan untuk mengelompokkan objek-objek berdasarkan kesamaan karakteristiknya[3]. Dalam konteks evaluasi kemampuan dasar pemrograman mahasiswa, metode klusterisasi dapat digunakan untuk mengelompokkan mahasiswa berdasarkan tingkat kemampuan pemrograman mereka. Metode *k-means clustering* adalah teknik pengelompokan data yang tidak memerlukan data pelatihan (training data) untuk melakukan pengelompokan atau klasifikasi objek data. Oleh karena itu, metode ini dikategorikan sebagai bagian dari *unsupervised machine learning*[4].

Clustering merupakan teknik yang penting dalam ilmu data, digunakan untuk mengidentifikasi struktur kelompok dalam suatu kumpulan data. Metode ini bekerja dengan cara mengelompokkan data berdasarkan tingkat kesamaan di dalam satu kelompok dan memaksimalkan perbedaan antara kelompok-kelompok yang berbeda[5]. *K-means clustering* adalah algoritma *unsupervised learning* yang dipakai untuk mengelompokkan dataset yang belum dilabel ke dalam *cluster* yang berbeda[6]. Analisis deskriptif dalam metode *unsupervised learning* memiliki beberapa kelemahan. Salah satu kelemahan utamanya adalah kesulitan dalam menginterpretasi hasilnya karena tidak adanya label yang mengarahkan model dalam mengenali pola data[7]. Metode *k-means* sangat tergantung pada jumlah *cluster* yang dibentuk dan tidak menentukan label pada *cluster* data yang dibentuk sehingga menyulitkan dalam melakukan analisis secara deskriptif[8][9]. Untuk mengatasi kelemahan dalam *K-means*, diperlukan metode lain yang dapat membantu proses penilaian hasil pengelompokan atau memandu pemilihan centroid yang lebih optimal.

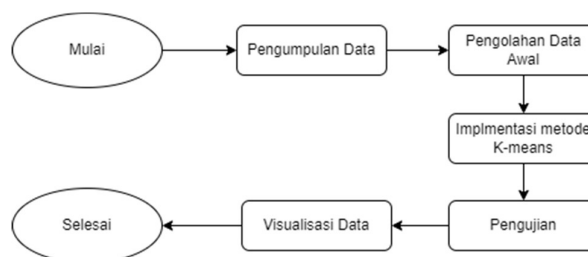
Metode SAW merupakan metode pengambilan keputusan (SPK) [10], Metode SAW adalah satu teknik dalam pengambilan

keputusan multi-kriteria yang dapat digunakan untuk menilai dan memberikan bobot pada hasil pengelompokan data. SAW memungkinkan evaluasi berbagai kriteria yang berbeda untuk menentukan mana solusi terbaik[11]. Dengan menggabungkan metode SAW dan *K-means*, diharapkan pengelompokan data bisa menjadi lebih akurat dan optimal. SAW dapat digunakan untuk menilai hasil pengelompokan yang dilakukan oleh *K-means*.

Penelitian serupa yang dilakukan oleh Ngaeni N dan tim mengombinasikan metode K-means dengan TOPSIS. Penelitian ini bertujuan menggunakan algoritma K-means Clustering untuk mengelompokkan data pendaftar berdasarkan latar belakang yang digambarkan melalui atribut seperti asal sekolah, pekerjaan orangtua, dan daerah asal. Setelah data dikelompokkan, proses dilanjutkan dengan penerapan metode DSS menggunakan TOPSIS untuk menentukan prioritas strategi pemasaran yang sesuai untuk setiap cluster[12]. Saputra E menggunakan metode SAW dalam menentukan top ranking dari masing – masing kelompok data siswa berprestasi yang terlebih dahulu dikelompokkan dengan metode *k-means*[13].

Penelitian ini berfokus pada mengisi celah penting dalam evaluasi kemampuan dasar pemrograman mahasiswa, yakni keterbatasan metode *K-means* standar yang belum mampu memberikan penilaian kuantitatif atas kualitas tiap kluster. Meskipun *K-means* mampu mengelompokkan mahasiswa berdasarkan kemiripan skor pemrograman, terdapat tiga permasalahan utama yaitu tidak adanya mekanisme penentuan centroid yang optimal, sehingga hasil clustering sangat dipengaruhi oleh pilihan nilai *k* dan inisialisasi acak, hasil clustering tidak dilengkapi label atau bobot yang menilai “seberapa baik” tiap kelompok, sehingga dosen kesulitan mengidentifikasi grup yang memerlukan intervensi atau penguatan materi, interpretasi hasil bersifat deskriptif dan subjektif, yang membuat keputusan.

Berdasarkan latarbelakang yang sudah dijelaskan oleh peneliti maka tujuan dari penelitian ini adalah Menganalisis penggabungan metode SAW dalam pengelompokan data menggunakan *K-means* untuk mengatasi kelemahan metode *k-means* dalam pelabelan *cluster* yang dibentuk. *K-Means* bersifat *unsupervised* dan tidak menghasilkan label yang memudahkan interpretasi. SAW berperan sebagai mekanisme evaluasi terarah, sehingga hasil *clustering* menjadi lebih interpretabel dan dapat diprioritaskan sesuai tujuan evaluasi (misalnya mengidentifikasi mahasiswa berpotensi tinggi).



Gambar 1. Tahapan Penelitian

2. METODE

2.1. Metodologi

Bagian ini akan menjelaskan mengenai tahapan dalam penelitian yang akan dilakukan. Tahapan yang dilakukan dalam penelitian ini seperti yang ditunjukkan pada Gambar 1.

2.1.1. Pengumpulan data

Teknik pengumpulan data dengan memberikan tes pada responden merupakan salah satu metode yang umum digunakan untuk mendapatkan informasi tentang pengetahuan, keterampilan, atau kemampuan responden. Tes dapat dilakukan dalam berbagai format, seperti:

Tes tertulis: Tes ini terdiri dari pertanyaan pilihan ganda, benar-salah, atau uraian singkat. Tes tertulis dibuat untuk mengetahui nilai teori dari mahasiswa.

Tes kinerja: Tes ini meminta responden untuk melakukan tugas berupa membuat program console. Tes kinerja dapat digunakan untuk mengukur keterampilan dan kecepatan mahasiswa dalam membuat suatu program..

2.1.2. Pengolahan data awal

Pengolahan data awal merupakan tahap krusial dalam analisis data. Tujuan dari langkah ini adalah untuk membersihkan, mentransformasi, serta merangkul data agar siap untuk dianalisis lebih lanjut. Pada penelitian ini, akan digunakan metode Minmax normalization untuk melakukan transformasi data. Proses normalisasi ini dilakukan dengan mengurangi setiap nilai data pada fitur dengan nilai minimum fitur tersebut, kemudian hasilnya dibagi dengan selisih antara nilai maksimum dan nilai minimum dari fitur yang sama.[14]

2.1.3. Implementasi metode K-means

Pada tahapan ini akan dilakukan proses pembuatan prototype sistem dengan menggunakan metode *K-means* dalam proses *clustering*. Tool yang digunakan pada tahapan ini adalah Matlab

2.1.4. Pengujian

Pengujian performa dilakukan dengan metode Dunn dan Silhouette untuk mengevaluasi nilai k yang digunakan. Metode Dunn berfokus pada pencarian nilai tertinggi, yang menandakan bahwa cluster yang dihasilkan semakin berbeda satu sama lain[15]. *Silhouette* mencari nilai tertinggi karena hal ini menunjukkan bahwa tingkat keyakinan terhadap penempatan data pada kluster semakin baik. [16].

2.1.5. Analisis SAW dalam Pelabelan metode k-means

Pada tahapan ini metode SAW akan melakukan parangkingan pada kelompok *cluster* yang dibentuk oleh metode *k-means* berdasarkan rata – rata nilai fitur pada tiap- tiap kelompok *cluster* data mahasiswa. Tahapan ini metode SAW melakukan pelabelan berdasarkan hasil dari implementasi metode K-Means. Metode SAW melakukan pelabelan pada hasil *cluster* yang terbentuk.

2.1.6. Visualisasi data

Hasil deteksi *clustering* kemudian ditampilkan dalam bentuk grafik untuk mempermudah pembacaan data. Analisis data dalam bentuk visualisasi dapat berupa diagram pie, histogram, scatter dan lainnya

3. HASIL

3.1. Pengumpulan data

Pengumpulan data dilakukan berdasarkan penilaian terhadap mahasiswa dalam aspek teori, praktek, serta ketepatan waktu dalam penyelesaian tugas teori dan praktek. Data ini digunakan untuk menganalisis performa masing-masing mahasiswa dalam dua aspek utama, yaitu nilai akademis (teori dan praktek) dan kedisiplinan dalam menyelesaikan tugas. Berikut adalah tabel 1 hasil pengumpulan data dari beberapa mahasiswa.

Tabel 1. Hasil Pengumpulan data

Nim	Nilai Teori	Nilai Praktek	Waktu Teori	Waktu Praktek
23xx01	Sangat Baik	Baik	Tepat Waktu	Tepat Waktu
23xx04	Baik	Sangat Baik	Tepat Waktu	Tepat Waktu
23xx05	Cukup	Baik	Tepat Waktu	Tepat Waktu
23xx09	Baik	Baik	Tepat Waktu	Tepat Waktu
23xx13	Baik	Sangat Kurang	Tepat Waktu	Tepat Waktu
23xx16	Baik	Cukup	Tepat Waktu	Tepat Waktu
23xx17	Baik	Baik	Tepat Waktu	Tepat Waktu
23xx21	Cukup	Sangat Kurang	Tepat Waktu	Tepat Waktu
23xx25	Cukup	Cukup	Tepat Waktu	Tidak Tepat Waktu
23xx28	Baik	Sangat Baik	Tepat Waktu	Tepat Waktu
23xx29	Kurang	Cukup	Tidak Tepat Waktu	Tidak Tepat Waktu
23xx33	Cukup	Cukup	Tepat Waktu	Tidak Tepat Waktu
23xx37	Baik	Baik	Tepat Waktu	Tepat Waktu
23xx40	Sangat Baik	Cukup	Tepat Waktu	Tidak Tepat Waktu
23xx41	Baik	Baik	Tepat Waktu	Tepat Waktu
23xx45	Sangat Kurang	Baik	Tepat Waktu	Tepat Waktu

3.2. Pengolahan data awal

Proses **pengolahan data awal** (prapemrosesan) data mentah seperti yang terlihat pada tabel 1 menjadi bentuk yang lebih terstruktur seperti tabel 2 yang ditunjukkan biasanya melibatkan beberapa langkah. Berikut adalah penjelasan tahapan prapemrosesan yang dilakukan untuk menormalisasi dan membersihkan data

3.2.1. Entri Data & Normalisasi

Nilai-nilai mentah, seperti "Sangat Baik," "Baik," "Cukup," "Kurang," dan "Sangat Kurang," pertama-tama harus dikonversi

menjadi representasi numerik. Proses ini melibatkan pemetaan label kualitatif ke skor numerik. Contohnya:

- "Sangat Baik" → 1
- "Baik" → 0.8
- "Cukup" → 0.6
- "Kurang" → 0.4
- "Sangat Kurang" → 0.2

Langkah ini penting untuk analisis kuantitatif dan komputasi, karena nilai kualitatif tidak bisa diproses langsung untuk operasi numerik seperti rata-rata atau perbandingan.

Kinerja waktu untuk tugas teori dan praktek dikategorikan menjadi "Tepat Waktu" dan "Tidak Tepat Waktu." Untuk konsistensi, kategori-kategori ini juga dipetakan ke nilai numerik:

- "Tepat Waktu" → 1
- "Tidak Tepat Waktu" → 0

Untuk kolom "RATA WAKTU," rata-rata dari kedua nilai ini (teori dan praktek) dihitung untuk mendapatkan satu nilai per baris. Misalnya:

- Jika teori dan praktek keduanya tepat waktu (1 dan 1), rata-rata adalah 1.
- Jika teori tepat waktu tetapi praktek tidak (1 dan 0), rata-ratanya adalah 0.5.
- Jika salah satunya terlambat dan yang lainnya tepat waktu, rata-ratanya adalah 0.75.

Tabel 2. Data ternormalisasi

Nilai Teori	Nilai Praktek	Rata Waktu
1	0.8	1
0.8	1	1
0.6	0.8	1
0.8	0.8	1
0.8	0.2	1
0.8	0.6	1
0.8	0.8	1
0.6	0.2	1
0.6	0.6	0.75
0.8	1	1
0.4	0.6	0.5
0.6	0.6	0.75
0.8	0.8	1
1	0.6	0.75
0.8	0.8	1
0.2	0.8	1
0.8	0.6	1
0.6	0.2	1
0.6	0.2	1

3.3. Implementasi metode K-means

Metode *K-means* adalah salah satu algoritma *clustering* (pengelompokan) yang sering digunakan dalam analisis data. Algoritma ini berfungsi untuk membagi sekumpulan data ke dalam sejumlah *cluster* (*K cluster*) berdasarkan kesamaan antar data tersebut. Pada penelitian ini dilakukan percobaan dalam menguji nilai *k* dari *k* = 2 sampai *k*=10 kemudian dilakukan

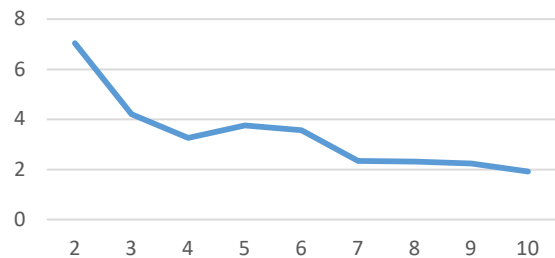
<https://doi.org/10.25077/TEKNOSI.v11i2.2025.178-184>

perhitungan menggunakan metode elbow untuk mencari nilai *k* terbaik. Tabel 3. Merupakan nilai SSE dari setiap nilai *k* yang diuji

Tabel 3. Nilai SSE untuk setiap nilai K

Nilai K	Iterasi	SSE
2	1	7.037133
3	2	4.211711
4	3	3.254313
5	3	3.766824
6	2	3.560687
7	4	2.334985
8	3	2.311544
9	2	2.231505
10	3	1.910044

Grafik Elbow



Gambar 2. Grafik metode elbow

Berdasarkan **grafik Elbow** pada gambar 2, dapat dilihat penurunan yang tajam dari *K*=2 hingga *K*=3, kemudian penurunan mulai melambat setelah itu. Pada titik ini (sekitar *K*=3), penurunan dalam inertia mulai lebih stabil, sehingga menambah *cluster* lebih banyak tidak akan memberikan banyak keuntungan dalam hal penurunan inertia, dengan demikian, berdasarkan grafik ini, jumlah *cluster* optimal adalah **3**, karena setelah titik tersebut (*K*=3), penurunan inertia mulai mendatar dan tidak seefisien penurunan sebelumnya.

Grafik Sebaran Cluster



Gambar 3. Distribusi anggota cluster

Pada gambar 3 merupakan distribusi anggota pada tiap-tiap *cluster* yang dibentuk. *Cluster* 3 memiliki jumlah anggota terbesar dengan **31 anggota**, diikuti oleh *Cluster* 2 dengan **28 anggota**, dan *Cluster* 1 dengan jumlah anggota paling sedikit, yaitu **24 anggota**. **Distribusi anggota yang cukup merata** di

antara ketiga *cluster* menunjukkan bahwa data tidak terkonsentrasi dalam satu *cluster* saja. Ini mengindikasikan bahwa data berhasil dikelompokkan dengan baik, meskipun terdapat sedikit variasi dalam ukuran *cluster*.

3.4. Pengujian

Tabel 4 menunjukkan nilai Dunn Index dan Silhouette Index (SI) untuk berbagai nilai K (jumlah *cluster*)

Tabel 4. Hasil pengujian *cluster*

Nilai K	Dunn	SI
2	0.215353	0.358233
3	0.298142	0.585965
4	0.298142	0.488347
5	0.298142	0.428232
6	0.298142	0.326209
7	0.298142	0.482994
8	0.312348	0.381258
9	0.312348	0.456119
10	0.312348	0.468782

3.4.1. Pengujian Dunn Index

Dunn Index mengukur kualitas *clustering* dengan melihat rasio antara jarak terdekat antar *cluster* dengan jarak terjauh antar anggota dalam *cluster*. Semakin tinggi nilai Dunn Index, semakin baik pemisahan antar *cluster* dan semakin kompak *cluster* tersebut.

- Pada tabel di atas, nilai Dunn Index mulai meningkat dari 0.215 pada K=2 menjadi 0.298 pada K=3 dan seterusnya.
- K=3 hingga K=7 mempertahankan nilai yang sama (0.298), menandakan bahwa setelah K=3, pemisahan *cluster* mungkin tidak signifikan membaik.
- Pada K=8 hingga K=10, nilai Dunn Index sedikit meningkat menjadi 0.312, menunjukkan sedikit perbaikan dalam kualitas *clustering*, tetapi peningkatan ini relatif kecil.

3.4.2. Silhouette Index

Silhouette Index (SI) mengukur seberapa baik sebuah data berada dalam *cluster*nya dan seberapa jauh data tersebut dari *cluster* lain. Nilai SI berkisar antara -1 hingga 1, di mana nilai mendekati 1 menunjukkan *clustering* yang baik, sementara nilai mendekati -1 menunjukkan *clustering* yang buruk.

- Nilai SI tertinggi ditemukan pada K=3 dengan 0.585, menunjukkan bahwa pada K=3, data di *cluster* memiliki pengelompokan yang paling baik dan jarak antar *cluster* yang signifikan.
- Setelah K=3, nilai SI menurun untuk K=4 hingga K=6, dengan penurunan yang cukup signifikan pada K=6 menjadi 0.326. Ini menunjukkan bahwa *clustering* setelah K=3 menjadi kurang optimal, dengan data yang semakin sulit dikelompokkan secara baik.
- Pada K=7 hingga K=10, SI sedikit meningkat kembali tetapi tidak mencapai nilai setinggi K=3, dengan nilai SI tertinggi sebesar 0.482 pada K=7.

3.5. Analisis SAW dalam Pelabelan metode k-means

Metode Simple Additive Weighting (SAW) adalah salah satu metode yang sering digunakan untuk pengambilan keputusan berbasis multi-kriteria. Dalam kasus ini, SAW digunakan untuk menentukan ranking pada *cluster* hasil *K-means* berdasarkan beberapa atribut, yaitu AVG teori, AVG praktek, dan AVG waktu

Tabel 5. Hasil Pelabelan *cluster* dengan Metode SAW

<i>Cluster</i>	AVG teori	AVG Praktek	AVG waktu	SAW	Rank
<i>Cluster 1</i>	0.500	0.525	0.854	0.606	3
<i>Cluster 2</i>	0.893	0.464	0.955	0.859	2
<i>Cluster 3</i>	0.819	0.858	0.992	0.951	1

Cluster 1 memiliki nilai SAW yang relatif rendah dibandingkan dengan *cluster* lainnya, yang menunjukkan bahwa performanya berada di bawah *cluster* lain dalam hal gabungan dari ketiga kriteria. *Cluster 2* memiliki nilai teori yang sangat tinggi (0.893), tetapi performa pada praktek (0.464) lebih rendah dibandingkan *cluster* lainnya. Nilai SAW 0.859 menempatkan *cluster* ini di posisi kedua. *Cluster 3* memiliki nilai tinggi pada semua kriteria. Ini menyebabkan nilai SAW-nya paling tinggi, yaitu 0.951, sehingga ditempatkan pada ranking pertama

4. PEMBAHASAN

Pada gambar 3. Garfik menunjukan perbedaan ukuran *cluster* yang relatif kecil (selisih 3 hingga 7 anggota antara *cluster*) menyiratkan bahwa data dalam dataset mungkin memiliki kemiripan yang cukup besar di antara beberapa kelompok, sehingga sulit memisahkannya ke dalam *cluster* yang sangat berbeda dalam hal ukuran. Ukuran *cluster* yang hampir seimbang ini juga dapat mengindikasikan bahwa tidak ada outlier atau kelompok data yang sangat dominan.

Berdasarkan hasil pengujian, dapat menyimpulkan bahwa K=3 memberikan hasil *clustering* yang paling optimal. Hal ini didukung oleh Dunn Index yang naik signifikan dari K=2 ke K=3, menunjukkan pemisahan *cluster* yang lebih baik pada K=3.

Silhouette Index tertinggi pada K=3, yang menandakan *cluster* yang lebih kompak dan terpisah dengan baik dibandingkan nilai K lainnya. Untuk nilai K di atas 3, baik Dunn Index maupun Silhouette Index menunjukkan penurunan, yang berarti menambah jumlah *cluster* tidak memberikan peningkatan signifikan dalam kualitas *clustering*.

Nilai K=3 adalah jumlah *cluster* yang paling optimal untuk dataset ini. Pada titik tersebut, *clustering* menghasilkan pemisahan yang baik dan anggota *cluster* yang kompak sesuai dengan hasil Dunn dan Silhouette Index.

Pelabelan menggunakan metode SAW dapat mempermudah analisis deskriptif pada metode *k-means* dengan hasil *Cluster 3* merupakan *cluster* terbaik karena memiliki nilai agregat (SAW) tertinggi, yang berarti kinerja teorinya, praktek, dan waktu secara keseluruhan lebih baik dibandingkan *cluster* lain.

Cluster 2 berada di peringkat kedua, meskipun memiliki nilai teori tinggi, namun kinerja praktiknya lebih rendah.

Cluster 1 berada di posisi terakhir karena memiliki nilai agregat SAW terendah, terutama karena nilai teorinya yang lebih rendah dibandingkan cluster lainnya.

Peringkat yang dihasilkan oleh metode SAW ini mencerminkan performa umum dari setiap cluster dengan mempertimbangkan semua kriteria. kombinasi evaluasi kualitas clustering (*Dunn & Silhouette*) dengan metode penilaian multi-kriteria SAW memberikan gambaran yang komprehensif tentang performa tiap kelompok serta arah perbaikan yang tepat.

5. KESIMPULAN

Berdasarkan hasil dan pembahasan dapat disimpulkan bahwa centroid awal sangat mempengaruhi kinerja cluster pada metode k-means. Metode perangkingan SAW dapat membantu dalam melakukan analisis deskriptif pada metode k-means clustering. Peringkat yang dihasilkan oleh metode SAW mencerminkan performa umum dari setiap cluster dengan mempertimbangkan semua kriteria

DAFTAR PUSTAKA

- [1] Y. E. Fadrial, "Klasterisasi Hasil Evaluasi Akademik Menggunakan Metode K-Means (Studi Kasus Fakultas Ilmu Komputer Unilak)," in *Prosiding-Seminar Nasional Teknologi Informasi & Ilmu Komputer (SEMASTER)*, 2020, pp. 53–65.
- [2] A. Yudhistira and R. Andika, "Pengelompokan Data Nilai Siswa Menggunakan Metode K-means Clustering," *Journal of Artificial Intelligence and Technology Information (JAITI)*, vol. 1, no. 1, pp. 20–28, Feb. 2023, doi: [10.58602/jaiti.v1i1.22](https://doi.org/10.58602/jaiti.v1i1.22).
- [3] F. Handayani, "Aplikasi Data Mining Menggunakan Algoritma K-means Clustering untuk Mengelompokkan Mahasiswa Berdasarkan Gaya Belajar," *Jurnal Teknologi dan Informasi*, doi: [10.34010/jati.v12i1](https://doi.org/10.34010/jati.v12i1).
- [4] A. Abdurrahman and F. Ismawan, "Model Machine Learning Klasifikasi Data Sekolah Tk Berdasarkan Status Dan Kabupaten/Kota Administrasi Provinsi Dki Jakarta," vol. 15, no. 2, pp. 1979–276, 2022, doi: [10.30998/faktorexacta.v15i2.13211](https://doi.org/10.30998/faktorexacta.v15i2.13211).
- [5] K. P. Sinaga and M. S. Yang, "Unsupervised K-means clustering algorithm," *IEEE Access*, vol. 8, 2020, doi: [10.1109/ACCESS.2020.2988796](https://doi.org/10.1109/ACCESS.2020.2988796).
- [6] M. Rival, M. Misriani, and L. O. Bakrim, "Penerapan Metode Cluster Dalam Data Mining Mengelompokkan Kenakalan Remaja (Studi Kasus Polda Sultra)," *SIMKOM*, vol. 9, no. 1, 2024, doi: [10.51717/simkom.v9i1.375](https://doi.org/10.51717/simkom.v9i1.375).
- [7] N. Husnaningtyas and T. Dewayanto, "Financial Fraud Detection And Machine Learning Algorithm (Unsupervised Learning): Systematic Literature Review," *Jurnal Riset Akuntansi dan Bisnis Airlangga*, vol. 8, no. 2, p. 2023, [Online]. Available: <https://e-journal.unair.ac.id/jraba>
- [8] M. M. Sebatubun and A. Prayitna, "Implementasi Metode Elbow Dan K-means Clustering Untuk," *Jurnal Informasi Interaktif*, vol. 8, no. 2, 2023.
- [9] M. P. A. Ariawan, I. B. A. Peling, and G. B. Subiksa, "Prediksi Nilai Akhir Matakuliah Mahasiswa Menggunakan Metode K-means Clustering (Studi Kasus : Matakuliah Pemrograman Dasar)," *Jurnal Nasional Teknologi dan Sistem Informasi*, vol. 9, no. 2, 2023, doi: [10.25077/teknosi.v9i2.2023.122-131](https://doi.org/10.25077/teknosi.v9i2.2023.122-131).
- [10] Refiza, *Penerapan Metode Simple Additive Weighting (Saw) Untuk Seleksi Tenaga Kerja*.
- [11] D. A. Rismana, M. I. Hanafri, and M. Iqbal, "Sistem Penunjang Keputusan Evaluasi Kinerja Karyawan Kontrak menggunakan Metode SAW Berbasis Web pada PT EKSTRINDO LAMINASI," *Jurnal Teknologi, Pendidikan dan Manajemen Global*, vol. 1, no. 1, pp. 1–6, 2022.
- [12] N. S. Ngaeni, K. Kusriani, and K. Kusnawi, "Analisis Kombinasi Algoritma K-means Clustering dan TOPSIS Untuk Menentukan Pendekatan Strategi Marketing Berdasarkan Background Target Audiens," *Journal of Computer System and Informatics (JoSYC)*, vol. 5, no. 2, pp. 393–403, Feb. 2024, doi: [10.47065/josyc.v5i2.4948](https://doi.org/10.47065/josyc.v5i2.4948).
- [13] E. A. Saputra and Y. Nataliani, "Analisis Pengelompokan Data Nilai Siswa untuk Menentukan Siswa Berprestasi Menggunakan Metode Clustering K-means," *Journal of Information Systems and Informatics*, vol. 3, no. 3, 2021, doi: [10.51519/journalisi.v3i3.164](https://doi.org/10.51519/journalisi.v3i3.164).
- [14] D. Azzahra Nasution, H. H. Khotimah, and N. Chamidah, "Perbandingan Normalisasi Data Untuk Klasifikasi Wine Menggunakan Algoritma K-Nn," *Journal of Computer Engineering System and Science*, vol. 4, no. 1, pp. 2502–7131, 2019.
- [15] S. Monalisa and F. Kurnia, "Analysis of DBSCAN and K-means algorithm for evaluating outlier on RFM model of customer behaviour," *Telkomnika (Telecommunication Computing Electronics and Control)*, vol. 17, no. 1, 2019, doi: [10.12928/TELKOMNIKA.v17i1.9394](https://doi.org/10.12928/TELKOMNIKA.v17i1.9394).
- [16] S. Monalisa, "Clusterisasi Customer Lifetime Value dengan Model LRFM menggunakan Algoritma K-means," *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 5, no. 2, 2018, doi: [10.25126/jtiik.201852690](https://doi.org/10.25126/jtiik.201852690).

LAMPIRAN

Lampiran 1. Hasil Cluster 1

NIM	Teori	Praktek	Rata-rata waktu	Cluster
23xx021	0.6	0.2	1	1
23xx025	0.6	0.6	0.75	1
23xx029	0.4	0.6	0.5	1
23xx033	0.6	0.6	0.75	1
23xx045	0.2	0.8	1	1
23xx052	0.6	0.2	1	1
23xx053	0.6	0.2	1	1
23xx057	0.6	0.6	0.5	1
23xx061	0.6	0.6	1	1
23xx068	0.6	0.6	0.5	1

23xx077	0.2	0.6	0.75	1
23xx010	0.2	0.6	0.75	1
23xx020	0.6	0.6	1	1
23xx034	0.6	0.4	1	1
23xx044	0.2	0.8	0.75	1
23xx050	0.6	0.6	1	1
23xx056	0.2	0.6	0.75	1
23xx066	0.6	0.2	0.5	1
23xx024	0.6	0.6	1	1
23xx027	0.6	0.2	1	1
23xx043	0.4	0.6	1	1
23xx047	0.6	0.6	1	1
23xx059	0.6	0.6	1	1
23xx080	0.6	0.6	1	1

Lampiran 2. Hasil Cluster 2

NIM	Teori	Praktek	Rata-rata waktu	Cluster
23xx013	0.8	0.2	1	2
23xx016	0.8	0.6	1	2
23xx040	1	0.6	0.75	2
23xx049	0.8	0.6	1	2
23xx064	0.8	0.6	1	2
23xx069	1	0.6	1	2
23xx073	0.8	0.6	1	2
23xx002	1	0.6	1	2
23xx014	0.8	0.2	0.75	2
23xx023	0.8	0.2	0.75	2
23xx026	1	0.2	0.75	2
23xx030	1	0.6	1	2
23xx032	1	0.6	1	2
23xx038	0.8	0.6	1	2
23xx054	1	0.6	1	2
23xx062	0.8	0.6	1	2
23xx078	0.8	0.6	1	2
23xx003	0.8	0.2	1	2
23xx011	1	0.6	1	2
23xx019	0.8	0.2	1	2
23xx031	0.8	0.6	1	2
23xx035	1	0.2	0.75	2
23xx036	1	0.2	1	2
23xx051	0.8	0.6	1	2
23xx060	1	0.6	1	2
23xx072	0.8	0.4	1	2
23xx079	1	0.6	1	2
23xx084	1	0.2	1	2

Lampiran 3. Hasil Cluster 3

NIM	Teori	Praktek	Rata-rata waktu	Cluster
23xx001	1	0.8	1	3
23xx004	0.8	1	1	3
23xx005	0.6	0.8	1	3
23xx009	0.8	0.8	1	3
23xx017	0.8	0.8	1	3
23xx028	0.8	1	1	3
23xx037	0.8	0.8	1	3
23xx041	0.8	0.8	1	3
23xx065	1	0.8	1	3
23xx081	0.8	1	1	3
23xx006	0.6	1	1	3
23xx018	1	0.8	1	3
23xx022	1	0.8	1	3
23xx042	0.8	0.8	1	3
23xx046	1	0.8	1	3
23xx058	0.6	1	1	3
23xx070	0.8	0.8	1	3
23xx074	1	0.8	1	3
23xx076	0.6	1	1	3
23xx082	0.6	0.8	1	3
23xx007	1	0.8	1	3
23xx012	0.8	0.8	1	3
23xx015	0.8	0.8	0.75	3
23xx039	0.6	0.8	1	3
23xx048	1	1	1	3
23xx055	1	1	1	3
23xx063	1	0.8	1	3
23xx067	0.6	0.8	1	3
23xx071	1	0.8	1	3
23xx075	0.8	1	1	3
23xx083	0.6	0.8	1	3