
Pengukuran Kinerja Algoritma K-Means dan Hierarchical Agglomerative Clustering Dalam Pengelompokan Perkara Dispensasi Kawin di Wilayah Pengadilan Tinggi Agama Makassar

Slamet^{*1}, Kusri²

^{1,2}) Magister Teknik Informatika, Universitas Amikom Yogyakarta, Yogyakarta, Indonesia

Email: ^{*1} slametmieno@students.amikom.ac.id, ² kusri@amikom.ac.id,

(Naskah masuk: 8 Februari 2025, diterima untuk diterbitkan: 20 Oktober 2025)

Abstrak: Dispensasi kawin atau dispensasi nikah merupakan sebuah upaya bagi masyarakat yang ingin menikah, namun belum memenuhi persyaratan batas usia untuk menikah yang ditetapkan oleh pemerintah, sehingga perlu mengajukan proses dispensasi kawin ke Pengadilan Agama melalui proses persidangan. Melalui algoritma K-Means Clustering dan Hierarchical Agglomerative Clustering dilakukan klusterisasi terhadap data dispensasi kawin di wilayah Pengadilan Tinggi Agama Makassar. Hasil Klusterisasi di evaluasi menggunakan metode Davies-Bouldin Index. Proses klusterisasi yang dilakukan pada data dispensasi kawin Pengadilan Tinggi Agama Makassar, dihasilkan jumlah kluster yang sama dari dua algoritma yang digunakan yakni 4 kluster. Nilai evaluasi terhadap kluster yang dihasilkan oleh dua algoritma yang digunakan, algoritma K-Means menghasilkan nilai Davies-Bouldin Index 1,898 dan algoritma Hierarchical Agglomerative Clustering dihasilkan nilai Davies-Bouldin Index 1,906. Mengacu pada nilai evaluasi Davies-Bouldin Index terhadap dua algoritma tersebut, dikatakan proses klusterisasi yang dihasilkan cukup baik.

Kata Kunci - Dispensasi Kawin; K-Means Clustering; Hierarchical Agglomerative Clustering; Pengadilan Agama; Davies-Bouldin Index

Performance Measurement of K-Means Algorithm and Hierarchical Agglomerative Clustering in Grouping Marriage Dispensation Cases at Makassar High Religious Court

Abstract: Marriage dispensation or marriage dispensation is an effort for people who want to get married, but do not yet meet the age limit requirements for marriage set by the government, so they need to apply for the marriage dispensation process to the Religious Court through a trial process. Using the K-Means Clustering and Hierarchical Agglomerative Clustering algorithms, clustering of marriage dispensation data in the Makassar High Religious Court area was carried out. Clustering results were evaluated using the Davies-Bouldin Index method. The clustering process carried out on marriage dispensation data from the Makassar High Religious Court resulted in the same number of clusters from the two algorithms used, namely 4 clusters. The evaluation value of the clusters produced by the two algorithms used, the K-Means algorithm produced a Davies-Bouldin Index value of 1.898 and the Hierarchical Agglomerative Clustering algorithm produced a Davies-Bouldin Index value of 1.906. Referring to the Davies-Bouldin Index evaluation value for the two algorithms, it is said that the resulting clustering process is quite good.

Keywords - Marriage Dispensation; K-Means Clustering; Hierarchical Agglomerative Clustering; Religious Courts; Davies-Bouldin Index

1. PENDAHULUAN

Dispensasi kawin atau dispensasi nikah merupakan sebuah upaya bagi masyarakat yang ingin menikah, namun belum memenuhi persyaratan batas usia untuk menikah yang ditetapkan oleh pemerintah, sehingga perlu mengajukan proses dispensasi kawin ke Pengadilan Agama melalui proses persidangan. Kasus dispensasi kawin yang terjadi di Indonesia, khususnya di wilayah Pengadilan Tinggi Agama Makassar, Sulawesi Selatan tergolong tinggi. Hal ini menjadikan perhatian serius bagi pemerintah provinsi Sulawesi Selatan dan jajarannya, Kementerian Agama

dan Mahkamah Agung dalam upaya pencegahan melonjaknya angka kasus dispensasi kawin di wilayah Sulawesi Selatan. Sehingga perlu adanya pemetaan atau pengelompokan perkara atau kasus dispensasi kawin berdasarkan wilayah yang ada di Sulawesi Selatan. Upaya-upaya yang mungkin bisa dilakukan adalah menjalin kerja sama dengan mengadakan Memorandum of Understanding (MoU) diantara Pemerintah Provinsi, Kementerian Agama dan Mahkamah Agung. Dengan adanya MoU ini diharapkan akan muncul gagasan-gagasan yang dapat mengetahui sebaran kasus dispensasi kawin, penyebab adanya kasus dispensasi kawin dan juga tindakan yang dapat mencegah terjadinya kasus dispensasi kawin di wilayah hukum Pengadilan Tinggi Agama Makassar, Sulawesi Selatan. Salah satu upaya yang dilakukan oleh Mahkamah Agung, melalui Pengadilan Tinggi Agama Makassar selaku pemilik data kasus dispensasi kawin di wilayah Sulawesi Selatan, menerapkan metode pada data mining untuk mengetahui tingkat sebaran kasus dispensasi kawin di wilayah Sulawesi Selatan dan mengelompokkan kasus dispensasi kawin berdasarkan wilayah yang ada di Sulawesi Selatan.

Data mining merupakan sebuah proses pengumpulan data yang bertujuan untuk mengekstraksi informasi penting yang ada pada data kasus dispensasi kawin[1]. Data mining yang digunakan pada penelitian ini bertujuan untuk mengetahui tingkat akurasi algoritma K-Means Clustering dalam kasus dispensasi kawin di Sulawesi Selatan, sehingga didapatkan pola atau informasi yang dapat digunakan sebagai acuan dalam menentukan kebijakan dalam mengatasi atau menurunkan tingkat kasus dispensasi kawin di wilayah hukum Pengadilan Tinggi Agama Makassar, Sulawesi Selatan.

Penelitian di dalamnya meneliti tentang Penerapan Algoritma K-Means Clustering dan Hierarchical Clustering dalam Mengelompokkan Data Pengangguran di Kawarang [2], menjadi salah satu alasan mengapa peneliti mengukur kinerja algoritma serupa, yakni K-Means dan Hierarchical Agglomerative Clustering pada Pengelompokan Perkara Dispensasi Kawin di Wilayah Pengadilan Tinggi Agama Makassar dan dievaluasi menggunakan Davies Boulden Index (BDI)[3], sehingga dapat ditentukan model yang terbaik dalam pengelompokan kasus dispensasi kawin di wilayah Pengadilan Tinggi Agama Makassar.

Analisis terhadap data dengan menggunakan pembelajaran mesin (machine learning) dengan menggunakan algoritma K-Means Clustering ini sudah digunakan pada pengelompokan data perkara perceraian berdasarkan kelurahan di kota Jambi[4]. Tujuan dari penelitian ini adalah bagaimana algoritma K-Means Clustering dapat digunakan untuk peningkatan kasus perceraian berdasarkan kelurahan yang berada di Kota Jambi.

Pengamatan lain juga dilakukan pada kasus yang berbeda oleh Dhimas Alif Fajar Fadhilah dan rekan, metode algoritma K-Means Clustering digunakan untuk mengetahui pemetaan lahan kopi di Kabupaten Malang. Penelitian ini menghasilkan sebuah kesimpulan bahwa metode algoritma K-Means Clustering mampu mengelompokkan data lahan pertanian kopi menjadi tiga kluster, yakni kluster rendah (<300), sedang ($300-1000$) dan kluster tinggi (>1000)[5].

Penelitian dengan metode K-Means juga dilakukan oleh Yosinta Pratiwi dengan fokus pada pengelompokan klaster harapan hidup berdasarkan tingkat provinsi. Penelitian ini menghasilkan sebuah kesimpulan bahwa algoritma K-Means Clustering mampu menghasilkan 2 klaster harapan hidup pada tingkat provinsi dengan nilai evaluasi Davies Boulden Index 0,222 pada klaster 2[6].

Penelitian internasional yang dilakukan oleh Erwin Lanceta Chapin dkk pada tahun 2023, dengan judul Clustering of students admission data using k-means, hierarchical and DBSCAN algorithm, menghasilkan 2 buah klaster yang memungkinkan bagi admisi untuk mengarahkan siswa pada kelas yang sesuai dengan minat dan bakat yang dimiliki oleh siswa[7].

Dengan memanfaatkan metode algoritma K-Means Clustering, peneliti dapat menentukan pemetaan kasus dispensasi kawin di wilayah hukum Pengadilan Tinggi Agama Makassar berdasarkan alasan terjadinya dispensasi kawin, usia pemohon, latar belakang keluarga dan pendidikan, sehingga dapat dilakukan tindakan-tindakan yang dapat mengurangi kasus dispensasi kawin di wilayah hukum Pengadilan Tinggi Agama Makassar.

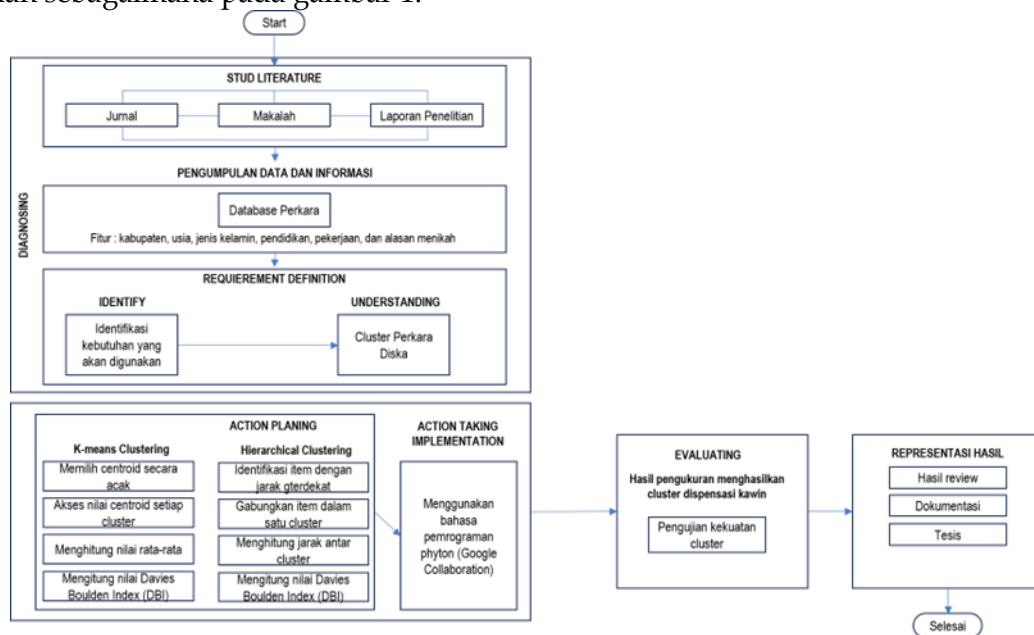
2. METODE PENELITIAN

2.1. Pengumpulan Data

Dataset penelitian merupakan data kasus dispensasi kawin di wilayah Pengadilan Tinggi Agama Makassar, Sulawesi Selatan yang didapatkan dengan cara observasi data, studi dokumen dan wawancara melalui bagian kepaniteraan Pengadilan Tinggi Agama Makassar.

2.2. Tahapan Penelitian

Tahapan penelitian diawali dengan melakukan study literature untuk menentukan kriteria data penelitian. Tahapan selanjutnya adalah melakukan analisis data dengan menggunakan Algoritma K-Means Clustering dan Hierarchical Agglomerative Clustering untuk mendapatkan klaster pada perkara dispensasi kawin. Selanjutnya hasil klaster yang terbentuk dievaluasi menggunakan Davies-Bouldin Index. Tahapan selanjutnya adalah menyusun rekomendasi dan kesimpulan sebagaimana pada gambar 1.



Gambar 1. Tahapan Penelitian

2.3. Penentuan Masalah dan Fitur Data

Tujuan yang akan dicapai pada penelitian ini adalah mengukur kinerja algoritma K-Means dan Hierarchical Agglomerative Clustering dalam dimensi data perkara dispensasi kawin di Pengadilan Tinggi Agama makassar dengan metode evaluasi Davies-Bouldin Index (DBI). Data yang diperoleh berjumlah 20.216 data dan fitur data pada penelitian ini berjumlah 6 fitur meliputi kabupaten, usia, jenis kelamin, pendidikan, pekerjaan, dan alasan nikah. Fitur data dalam dataset dispensasi kawin terdiri dari data numeric dan katgorikal, untuk memudahkan proses klasterisasi, maka fitur data yang berbentuk kategorikal harus dikonversi menjadi data numerik. Proses konversi (Preprocessing) ini dilakukan dengan menggunakan metode One-HotEncoding[8]. Data sebelum dilakukan preprocessing ditunjukkan pada tabel 1.

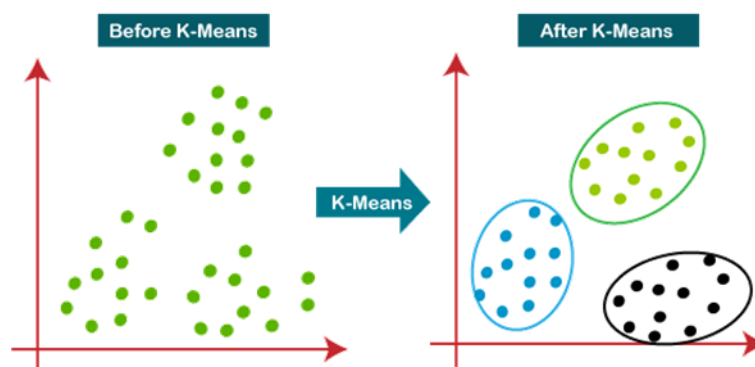
Tabel 1. Data Dispensasi Kawin

kabupaten	usia	jenis_kelamin	pendidikan	pekerjaan	alasan_nikah
Kab. Pangkajene	18	Laki-Laki	Sekolah Dasar	Lainnya	Menghindari Zina
Kab. Pangkajene	17	Perempuan	Sekolah Dasar	Tidak diketahui	Menghindari Zina

Kab. Pangkajene	17	Laki-Laki	Sekolah Dasar	Lainnya	Pergaulan Bebas
Kab. Pangkajene	15	Perempuan	Sekolah Dasar	Tidak diketahui	Pergaulan Bebas
Kab. Pangkajene	15	Perempuan	Sekolah Dasar	Tidak diketahui	Pergaulan Bebas
Kab. Pangkajene	17	Laki-Laki	Sekolah Dasar	Pedagang	Pergaulan Bebas
.....					

2.4. Pengelompokan Algoritma K-Means Clustering

Algoritma K-Means Clustering digunakan untuk membuat kluster-kluster yang saling berdekatan dan memiliki kemiripan satu sama lain[9]. Cara kerja dari algoritma K-Means ini ditunjukkan pada gambar 1.



Gambar 2. Cara Kerja K-Means Clustering

Sumber : <https://www.javatpoint.com/k-means-clustering-algorithm-in-machine-learning>

K-Means Clustering merupakan salah satu algoritma yang dapat digunakan untuk menemukan pusat kluster yang mewakili wilayah tertentu dari data obyek. Algoritma secara bergantian menetapkan setiap titik data ke pusat kluster terdekat dan kemudian menetapkan setiap pusat kluster sebagai rata-rata dari titik data yang ditugaskan padanya. Algoritma selesai ketika penugasan instan ke kluster tidak lagi berubah[10].

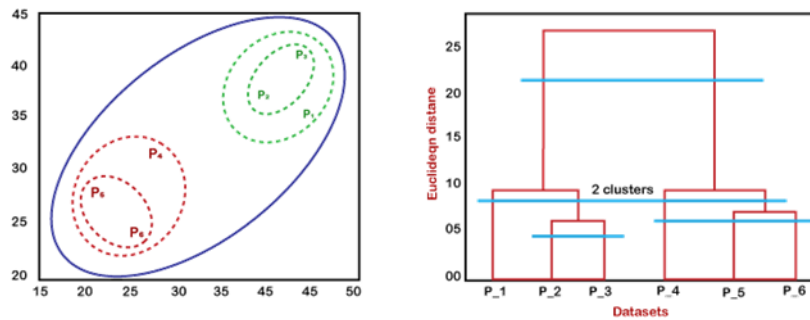
Tahapan untuk melakukan klustering dengan menggunakan algoritma K-Means Clustering adalah sebagai berikut[11] :

1. Peneliti menentukan jumlah kluster yang nantinya akan dibentuk
2. Menentukan centroid secara acak
3. Menghitung alokasi data berdasarkan jumlah kluster yang dibentuk. Kedekatan antara satu obyek dengan obyek lainnya akan ditentukan oleh jarak. Artinya, obyek mana yang paling dekat akan menentukan suatu data masuk ke dalam kluster tertentu.'
4. Menghitung nilai pusat kluster (centroid) yang merupakan nilai rata-rata atas semua obyek data dalam kluster.
5. Lakukan kembali alokasi masing-masing obyek menggunakan centroid baru sampai nilai pusat kluster tidak berubah.
6. Apabila nilai pusat kluster masih berubah atau ada obyek yang berpindah kelompok, maka proses akan dikembalikan pada jarak antara obyek ke centroid.

2.5. Pengelompokan Algoritma Hierarchical Agglomerative Clustering (HAC)

Hierarchical Agglomerative Clustering merupakan salah satu metode klustering yang mengelompokkan sejumlah data berdasarkan kemiripan yang membentuk pohon hirarki dari bawah ke atas, atau dikenal dengan istilah dendogram. Metode ini menggabungkan obyek atau

kluster secara bertahap, dimulai dengan setiap obyek sebagai kluster individual yang kemudian digabungkan sehingga membentuk kluster yang besar[12].



Gambar 2. Cara Kerja Hierarchical Agglomerative Clustering

Sumber : <https://www.javatpoint.com/hierarchical-clustering-in-machine-learning>

Prinsip dari Agglomerative Hierarchical Clustering ini menginisialisasi setiap obyek atau data individu sebagai kluster terpisah, yang kemudian digabungkan dengan kluster terdekat berdasarkan metrik jarak tertentu secara berulang sehingga membentuk kluster besar[13]. Proses ini akan berlanjut hingga mencapai kluster atau kriteria yang diinginkan. Proses penggabungan kluster pada metode agglomerative hierarchical clustering ini dapat divisualisasikan dalam bentuk dendrogram yang menunjukkan langkah-langkah penggabungan kluster secara hierarki. Visualisasi dalam bentuk dendrogram ini mampu menggambarkan secara jelas tentang jarak atau kemiripan antar kluster selama proses penggabungan[14].

Agglomerative Hierarchical Clustering memiliki beberapa keunggulan dalam melakukan metode clustering, beberapa keunggulan dari metode Agglomerative Hierarchical Clustering adalah [15]:

1. Tidak memerlukan jumlah kluster (nilai k) sebelumnya
2. Memberikan hasil yang informatif dalam bentuk dendrogram
3. Memiliki fleksibilitas dalam memilih metrik jarak dan metode penggabungan

2.6. Davies-Bouldin Index (DBI)

Davies-Bouldin Index adalah salah satu alat ukur atau metrik yang digunakan untuk mengevaluasi hasil klusterisasi dalam analisis data. Indeks ini mengukur sejauh mana kluster yang terbentuk dari algoritma klustering terpisah satu sama lain dan seberapa ketat atau dekat jarak antar kluster yang terbentuk. Tujuan dari metrik Davies-Bouldin Index adalah meminimalkan jarak antar kluster, semakin kecil angka indeks yang dihasilkan, maka kluster yang terbentuk akan semakin baik, lebih terpisah dan lebih padat[16].

Indeks yang dihasilkan pada Davies-Bouldin Index merepresentasikan rata-rata rasio antara ukuran kluster dengan jarak antar pusat kluster. Semakin kecil nilai indeks, maka semakin baik kluster yang terbentuk[17].

3. HASIL DAN PEMBAHASAN

Penelitian ini menggunakan data perkara dispensasi kawin di wilayah Pengadilan Tinggi Agama Makassar dari tahun 2015 sampai dengan 2024, yang berjumlah 20.216 data dan fitur atau atribut sebanyak 6.

3.1. Pemrosesan Data Awal

Pengolahan data awal dalam proses klustering sangat penting dalam menentukan kluster yang dihasilkan. Proses ini dilakukan untuk memastikan tidak adanya missing data, data tidak valid karena missing value sehingga data dapat proses klusterisasi dengan baik dan tidak merusak penelitian. Tahapan yang dilakukan dalam pemrosesan data awal sebagai berikut [18]:

1. Pemilihan data awal

Proses pemilihan data awal dilakukan pada dataset dengan menentukan fitur yang akan digunakan pada proses pemodelan. Data yang akan digunakan pada proses klusterisasi penelitian ini adalah 20.216 data dengan fitur yang telah ditentukan adalah 6 fitur meliputi kabupaten, usia, jenis kelamin, pendidikan, pekerjaan, dan alasan nikah.

2. Pembersihan Data

Proses pembersihan data dilakukan untuk menghindari data yang tidak memiliki value atau missing value. Tujuan lain dari proses pembersihan data ini adalah untuk memastikan tidak ada data yang ganda atau redundan, serta menghindari data yang memiliki tipe data yang tidak sesuai.

3. Inisiasi Data

Inisiasi data atau proses konversi dari tipe data kategorikal menjadi data numerik sehingga dapat dilakukan proses klusterisasi. Proses inisiasi data pada penelitian ini menggunakan metode One-HotEncoding, hal ini dilakukan karena data yang ada memiliki kategori yang cukup variatif. Proses inisiasi data dilakukan menggunakan bahasa python. Script untuk inisiasi data atau preprocessing dapat dilihat pada gambar 3.

```
fitur_numerik = ['usia']
fitur_kategorikal = ['kabupaten', 'jenis_kelamin', 'pendidikan', 'pekerjaan', 'alasan_nikah']

preprocessor = ColumnTransformer(
    transformers=[
        ('num', StandardScaler(), fitur_numerik),
        ('cat', OneHotEncoder(handle_unknown='ignore', sparse_output=False), fitur_kategorikal)
    ]
)
```

Gambar 3. Proses Preprocessing dengan One-HotEncoding

Hasil dari inisiasi data atau preprocessing ditunjukkan pada gambar 4.

```
[[ 0.49532952  0.         ...  0.         0.
   0.         ]
 [-0.02485654  0.         ...  0.         0.
   0.         ]
 [-0.02485654  0.         ...  1.         0.
   0.         ]
 ...
 [ 1.01551558  0.         ...  0.         0.
   0.         ]
 [ 0.49532952  0.         ...  0.         0.
   0.         ]
 [ 1.01551558  0.         ...  0.         0.
   0.         ]]
```

Gambar 4. Hasil Proses Preprocessing dengan One-HotEncoding

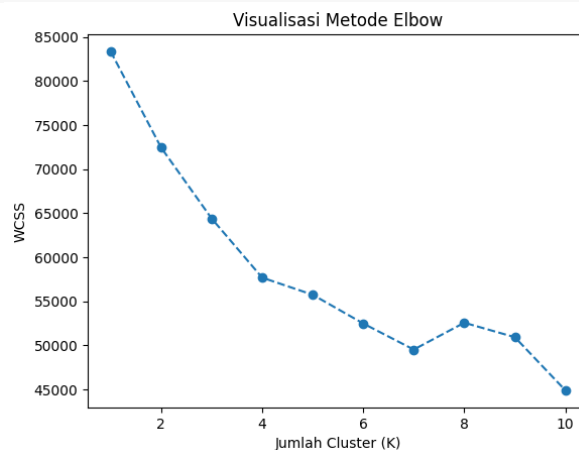
3.2. Perhitungan K-Means Clustering

Perhitungan K-Means Clustering diawali dengan menentukan jumlah klaster yang akan dibentuk. Proses dalam menentukan jumlah klaster dalam penelitian ini digunakan metode elbow[19]. Metode elbow digunakan karena dapat membantu menentukan jumlah klaster secara optimal. Metode elbow memberikan visualisasi dalam bentuk grafik yang membantu dalam menganalisis nilai wcss. Script untuk menentukan jumlah klaster dengan menggunakan metode elbow dapat ditunjukkan pada gambar 5.

```
wcss = []
k_range = range(1, min(11, len(data)))

for k in k_range:
    kmeans = KMeans(n_clusters=k, init='k-means++', random_state=42)
    kmeans.fit(X)
    wcss.append(kmeans.inertia_)

plt.plot(k_range, wcss, marker='o', linestyle='--')
plt.xlabel('Jumlah Cluster (K)')
plt.ylabel('WCSS')
plt.title('Visualisasi Metode Elbow')
plt.show()
```



Gambar 5. Menentukan jumlah kluster metode elbow

Setelah menentukan jumlah kluster, langkah selanjutnya adalah melakukan proses klasterisasi menggunakan k-means. Proses klasterisasi dengan k-means dapat ditunjukkan pada gambar 6.

```
# Klasterisasi menggunakan K-Means
kmeans = KMeans(n_clusters=k, init='k-means++', random_state=42)
kmeans_clusters = kmeans.fit_predict(X)
data['Klaster_K-Means'] = kmeans_clusters
```

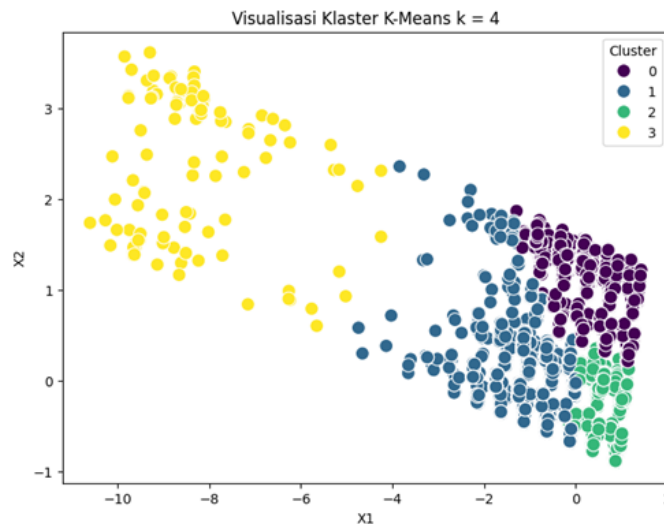
Gambar 6. Script klasterisasi k-means

Proses klasterisasi yang telah dilakukan menghasilkan nilai centroid dari masing-masing cluster. Nilai centroid dari hasil klasterisasi ditunjukkan pada gambar 7.

```
[ [ 4.13916709e-01  9.34928945e-03  3.08526552e-02  3.79581152e-02
    6.43231114e-02  3.88930441e-02  3.14136126e-02  2.50560957e-02
    3.73971578e-04  2.37471952e-02  4.02019447e-02  3.25355273e-02
    3.01047120e-02  3.42183994e-02  8.11518325e-02  1.75766642e-01
    5.29169783e-02  7.51682872e-02  4.30067315e-03  3.92670157e-03
    1.24532536e-01  4.22587883e-02  2.33732236e-02  1.75766642e-02
    1.00000000e+00 -3.81916720e-14  1.86347248e-18  1.86985789e-04
    2.79543755e-01  2.51495886e-01  3.87434555e-01  5.60957367e-04
    1.86985789e-04  2.24382947e-03  7.83470456e-02  1.96709050e-01
    1.86985789e-04  5.14210920e-02  1.12191473e-03  9.34928945e-04
    5.79655946e-03  3.73971578e-04  7.47943156e-04  3.73971578e-04
    1.86347248e-18  1.86347248e-18  7.47943156e-04  1.49588631e-03
    7.29244577e-02  2.43081526e-03  6.43231114e-02  1.00972326e-02
    1.86985789e-04  3.29094989e-02  1.86985789e-04  4.03889304e-02
```

Gambar 7. Centorid hasil klasterisasi

Proses klasterisasi yang dihasil sesuai dengan jumlah kluster yang telah ditentukan, dapat divisualiasikan dalam bentuk plotting pada gambar 8.



Gambar 8. Visualisasi hasil klasterisasi k-means

Langkah selanjutnya setelah dilakukan klasterisasi adalah mengevaluasi hasil pemodelan atau hasil klasterisasi untuk menentukan kualitas hasil klasterisasi. Metode evaluasi yang digunakan pada penelitian ini adalah Davies-Bouldin Index. Nilai indeks yang dihasilkan pada metode evaluasi Davies-Bouldin Index dapat dijadikan sebagai metrik untuk mengukur kinerja algoritma k-means clustering. Hasil evaluasi klasterisasi k-means dapat ditunjukkan pada gambar 9.

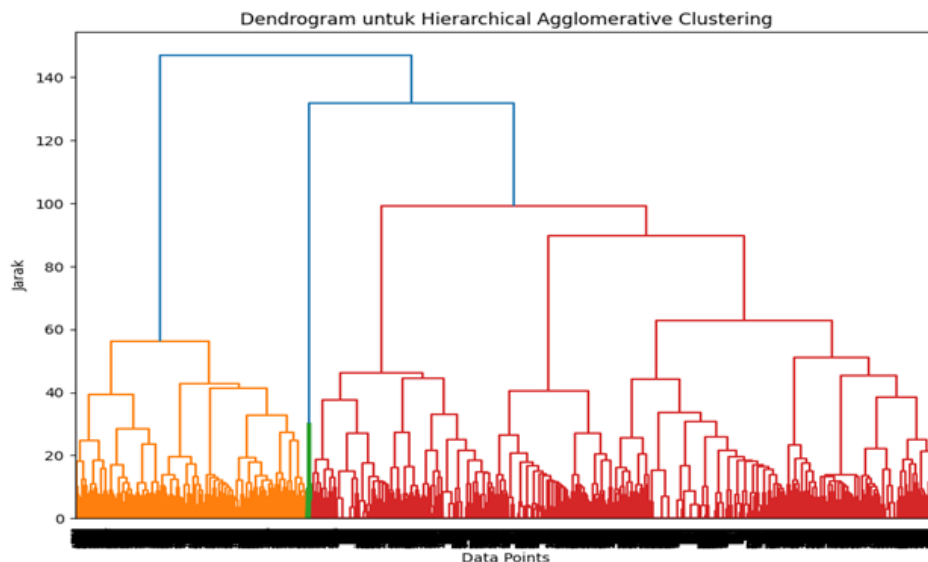
```
kmeans_dbi_index = davies_bouldin_score(X, kmeans_clusters) if k > 1 else None  
  
print("Nilai Davies-Bouldin Index K-Means Clustering : ", kmeans_dbi_index)  
  
Nilai Davies-Bouldin Index K-Means Clustering : 1.898387081608798
```

Gambar 8. Visualisasi hasil klasterisasi k-means

Hasil Davies-Bouldin Index yang diperoleh dari hasil klasterisasi k-means dalam dimensi data perkara dispensasi kawin, diperoleh nilai 1,898. Nilai indeks ini menunjukkan bahwa hasil klaster yang terbentuk cukup baik, dimana dalam Davies-Bouldin Index nilai yang mendekati nilai 0 menunjukkan bahwa hasil klaster yang terbentuk cukup baik.

3.3. Perhitungan Hierarchical Agglomerative Clustering

Perhitungan Hierarchical Agglomerative Clustering tidak memerlukan jumlah klaster diawal, sehingga tahapan yang dilakukan setelah proses preprocessing, mengukur jarak dan linkage antar klaster. Pengukuran jarak dan linkage antar klaster ini menggunakan metode ward's linkage[20]. Metode ward's linkage ini digunakan karena dapat menghasilkan kualitas klaster yang lebih baik dengan meminimalkan varian antar klaster. Pengukuran jarak antar klaster ini sangat menentukan kualitas klaster yang terbentuk pada saat proses klasterisasi. Langkah selanjutnya setelah dilakukan pengukuran jarak antar klaster dengan menggunakan metode ward's linkage, langkah selanjutnya adalah memvisualisasikan klasterisasi melalui dendogram. Dendogram berguna untuk menganalisis proses klaster yang terbentuk. Dendogram pada proses klasterisasi dapat ditunjukkan pada gambar 9.



Gambar 9. Dendrogram Hierarchical Agglomerative Clustering

Melalui hasil dendrogram, dapat ditentukan jumlah kluster yang terbentuk adalah 4 kluster. Setelah jumlah kluster terbentuk, selanjutnya adalah melakukan proses klusterisasi HAC. Proses klusterisasi dengan HAC digunakan script yang ditunjukkan gambar 10.

```
hac = AgglomerativeClustering(n_clusters=k, linkage='ward')  
hac_clusters = hac.fit_predict(X)  
data['HAC_Cluster'] = hac_clusters
```

Gambar 10. Script klusterisasi HAC

Hasil dari proses klusterisasi yang telah dilakukan, kemudian dievaluasi dengan menggunakan Davies-Bouldin Index. Hasil dari evaluasi klusterisasi HAC dapat ditunjukkan pada gambar 11.

```
print("Nilai Davies-Bouldin Index HAC : ", hac_dbi_index)  
  
Nilai Davies-Bouldin Index HAC : 1.9060205394598015
```

Gambar 11. Davies-Bouldin Index HAC

Nilai Davies-Bouldin Index yang dihasilkan dari klusterisasi HAC adalah 1,906. Nilai ini menunjukkan bahwa hasil klusterisasi HAC cukup baik.

4. KESIMPULAN

Kesimpulan harus mengindikasikan secara jelas hasil-hasil yang diperoleh, kelebihan dan kekurangannya, serta kemungkinan pengembangan selanjutnya.

Kesimpulan dapat berupa paragraf, namun sebaiknya berbentuk point-point dengan menggunakan numbering.

UCAPAN TERIMA KASIH

Pembahasan pada bagian hasil menunjukkan bahwa hasil klusterisasi menggunakan algoritma K-Means Clustering dapat menghasilkan jumlah kluster sebanyak 4 kluster. Perhitungan nilai indeks dengan metode Davies-Bouldin Index diperoleh nilai 1,898. Sedangkan pada proses klusterisasi menggunakan Hierarchical Agglomerative Clustering dihasilkan jumlah kluster sebanyak 4, dengan nilai Davies-Bouldin Index 1,906. Melalui proses klusterisasi dengan menggunakan dua algoritma yang berbeda pada dimensi data dispensasi kawin di wilayah Pengadilan Tinggi Agama Makassar, dihasilkan jumlah kluster yang sama, yakni 4 kluster dengan nilai Davies-Bouldin Index yang memiliki kemiripan yakni 1,898 dan 1,906. Sehingga, proses klusterisasi yang dihasilkan dari dua algoritma tersebut cukup baik.

DAFTAR PUSTAKA

- [1] Bahar, A., Pramono, B., & Sagala, L. H. S. (2016). Penentuan strategi penjualan alat-alat tattoo di studio sonyxattoo menggunakan metode K-Means Clustering. *SemanTIK*, 2(2), 75–86.
- [2] A. A. R. Mulyana, A. R. Juwita, A. M. Siregar, dan A. Fauzi, “Penerapan Algoritma K-means Clustering dan Hierarchical Clustering dalam Mengelompokkan Data Pengangguran di Karawang”, *Jurnal Algoritma*, vol. 21, no. 2, hlm. 209–220, Nov 2024. DOI: 10.33364/algoritma/v.21-2.2155.
- [3] D. A. Tarigan, “Optimization of the K-Means Clustering Algorithm Using Davies Bouldin Index in Iris Data Classification,” *Media Online*, vol. 4, no. 1, pp. 545–552, 2023, doi: 10.30865/klik.v4i1.964.
- [4] Elvi Yanti. (2021). Analisis Algoritma K-Means Dalam Pengelompokan Perkara Perceraian Berdasarkan Kelurahan di Kota Jambi. *Processor Vol 16 No1*
- [5] Dimas Alif Fajar Fadhillah. (2022). Penerapan Metode K-Means Clustering pada Pemetaan Lahan Kopi di Kabupaten Malang. *JATI (Jurnal Mahasiswa Teknik Informatika)*, Vol 6 No 1.
- [6] Y. Pratiwi and M. Mulyawan, “Implementasi Algoritma K-Means untuk Menentukan Angka Harapan Hidup berdasarkan Tingkat Provinsi”, *Blend Sains J. Teknik*, vol. 1, no. 4, pp. 284–294, Mar. 2023.
- [7] E. L. Cahapin, B. A. Malabag, C. S. Santiago, J. L. Reyes, G. S. Legaspi, and K. L. Adrales, “Clustering of students admission data using k-means, hierarchical, and DBSCAN algorithms,” *Bulletin of Electrical Engineering and Informatics*, vol. 12, no. 6, pp. 3647–3656, Dec. 2023, doi: 10.11591/eei.v12i6.4849.
- [8] T. Smutek, J. Sikora, S. Bogacki, M. Rutkowski, and D. Woźniak, “Use of Autoencoder and One-Hot Encoding for Customer Segmentation,” 2024.
- [9] T. Hidayat, “Klasifikasi Data Jamaah Umroh Menggunakan Metode K-Means Clustering,” *J. Sistim Inf. dan Teknol.*, pp. 19–24, Feb. 2022, doi: 10.37034/jsisfotek.v4i1.115.
- [10] B. Raharjo, P Y YAYASAN PRIMA AGUS TEKNIK YAYASAN PRIMA AGUS TEKNIK YAYASAN PRIMA AGUS TEKNIK Pembelajaran Mesin (Machine Learning). 2021.
- [11] Y. Agusta, “K-Means-Penerapan, Permasalahan dan Metode Terkait,” 2007.
- [12] R. K. Dinata and N. Hasdyna, *Machine Learning*. UNIMAL PRESS, 2020.
- [13] B. Moseley and J. R. Wang, “Approximation Bounds for Hierarchical Clustering: Average Linkage, Bisecting K-means, and Local Search,” 2023. [Online]. Available: <http://jmlr.org/papers/v24/18-080.html>.
- [14] F. Bajaber, "Enhancing Subcluster Identification in IoT Sensor Networks with Hierarchical Clustering Algorithms and Dendrograms," 2023 IEEE 8th International Conference On Software Engineering and Computer Systems (ICSECS), Penang, Malaysia, 2023, pp. 459-464, doi: 10.1109/ICSECS58457.2023.10256413.
- [15] S. Wulandari, “Clustering Indonesian Provinces on Prevalence of Stunting Toddlers Using Agglomerative Hierarchical Clustering,” *Faktor Exacta*, vol. 16, no. 2, Jul. 2023, doi: 10.30998/faktorexacta.v16i2.17186.
- [16] N. Eliza, I. Husein, and S. Dur, “Determining Zoning of Areas Affected by Flood Disasters in Medan City Using Silhouette Coefficient and Davies Bouldin Index Analysis,” *Jurnal Pijar Mipa*, vol. 19, no. 3, pp. 558–563, May 2024, doi: 10.29303/jpm.v19i3.6707.
- [17] E. Purwaningsih and E. Nurelasari, “Implementasi Metode K-Means Clustering Dengan Davies Bouldin Index Pada Analisis Faktor Penyebab Perceraian,” *INFORMATION MANAGEMENT FOR EDUCATORS AND PROFESSIONALS*, vol. 7, no. 2, pp. 134–143, 2023, [Online]. Available: <https://databoks.katadata.co.id/>
- [18] A. A. Rismayadi, N. N. Fatonah, and E. Junianto, “ALGORITMA K-MEANS CLUSTERING UNTUK MENENTUKAN STRATEGI PEMASARAN DI CV. INTEGRREET KONSTRUKSI,”

JURNAL RESPONSIF, vol. 3, no. 1, pp. 30-36, 2021, [Online]. Available: <http://ejurnal.ars.ac.id/index.php/jti>

- [19] A. R. Lashiyanti, I. Rasyid Munthe, F. A. Nasution, and E. P. Korespondensi, "Optimisasi Klasterisasi Nilai Ujian Nasional dengan Pendekatan Algoritma K-Means, Elbow, dan Silhouette," Jurnal Ilmu Komputer dan Sistem Informasi (JIKOMSI, vol. 6, no. 1, pp. 14-20, 2023.
- [20] S. Indra Maiyanti, O. Dwipurwani, and N. Aisyah, "Pengelompokan Provinsi Di Indonesia Berdasarkan Konsumsi Kalori Per Kapita Sehari Menurut Kelompok Komoditas / Makanan Menggunakan Average Linkage Dan Ward's Method," Azizah, dan Nurul Aisyah, vol. 2, pp. 1698-1713, 2024, [Online]. Available: <https://journal.institercom-edu.org/index.php/multiple>