

OPTIMALISASI PENGIRIMAN BARANG DENGAN MENGELOMPOKKAN TITIK PENGIRIMAN MENGGUNAKAN METODE *K-MEANS CLUSTERING* DAN *HIERARCHICAL CLUSTERING*

Levina Olivia

Program Studi Teknik Informatika, Fakultas Teknologi Informasi, Universitas Tarumanagara,
Jln. Letjen S. Parman No. 1, Jakarta, 11440, Indonesia
E-mail: levina.53210041@stu.untar.ac.id

ABSTRAK

Optimasi merupakan suatu proses yang dilakukan untuk mencapai hasil yang ideal atau nilai efektif yang dapat dicapai. Optimasi juga dapat diartikan sebagai suatu bentuk mengoptimalkan sesuatu yang sudah ada, atau merancang dan membuat sesuatu secara optimal. Pada penelitian ini akan dilakukan optimasi terkait dengan pengiriman barang yang sebelumnya mungkin dilakukan secara manual sehingga kurang efektif dari segi waktu dan biaya. Optimasi ini akan menerapkan suatu metode *machine learning* yaitu *K-Means Clustering* dan *Hierarchical Clustering*, dimana titik-titik pengiriman akan dikumpulkan berdasarkan pengangkut yang tersedia. Setiap pengangkut pastinya akan memiliki kapasitas berat dan *volume* masing-masing. Hasil dari penelitian ini berupa beberapa *cluster* yang berisi titik-titik mana saja yang akan dikirim beserta pengangkut apa yang akan digunakan. Jumlah *cluster* disini akan diinisialisasi secara otomatis. Algoritma *K-Means Clustering* dengan jumlah *cluster* sebanyak 12 dan nilai *silhouette* sebesar 0,756 serta nilai DBI sebesar 0,36 cocok untuk dataset pada penelitian ini.

Kata kunci: Optimasi, *Machine Learning*, *Clustering*, *Cluster*, *Silhouette Score*, DBI

ABSTRACT

Optimization is a process carried out to achieve ideal results or effective values that can be achieved. Optimization can also be interpreted as a form of optimizing something that already exists, or designing and making something optimally. In this study, optimization will be carried out related to the delivery of goods that were previously possibly done manually so that it was less effective in terms of time and cost. This optimization will apply a method of machine learning, namely K-Means Clustering and Hierarchical Clustering, where delivery points will be collected based on the available carriers. Each carrier will definitely have their own weight and volume capacity. The results of this study are in the form of several clusters containing which points will be sent along with what carrier will be used. The number of clusters here will be initialized automatically. The K-Means Clustering algorithm with a number of clusters of 12 and a silhouette value of 0.756 and a DBI value of 0.36 is suitable for the dataset in this study

Keywords: Optimization, Machine Learning, Clustering, Cluster, Silhouette Score, DBI

1. PENDAHULUAN

Di era digitalisasi saat ini, banyak peluang bisnis yang didapat dengan menggunakan media online. Banyak orang saat ini yang membeli sesuatu barang baik untuk pribadi maupun untuk dijual kembali tidak lagi langsung datang ke tempatnya, tetapi banyak yang menggunakan platform belanja secara online. Hal ini dikarenakan adanya keuntungan secara efisiensi waktu dan biaya yang diperlukan sehingga memudahkan pelanggan dalam bertransaksi.

Bagi perusahaan distribusi, pengiriman barang merupakan aspek yang sangat krusial karena mempengaruhi efisiensi operasional, dan daya saing pasar. Pengelolaan yang buruk dapat menyebabkan peningkatan biaya operasional, keterlambatan, dan kerusakan barang, merugikan reputasi perusahaan dan kehilangan pelanggan. Riset menyatakan 77% pembeli online menyatakan bahwa mereka mengharapkan pesanan mereka akan dikirim dalam waktu dua jam atau kurang.

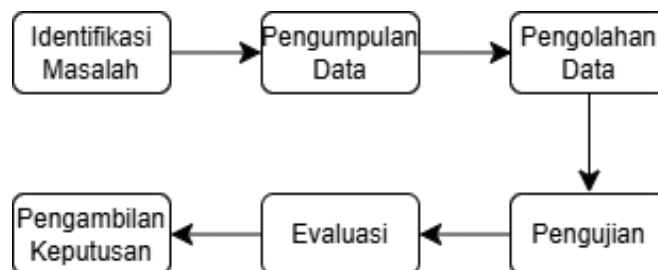
Pengiriman barang sangat berpengaruh dalam kepuasan pelanggan saat membeli suatu barang. Kepuasan pelanggan sendiri menjadi hal yang sangat diamati di dunia pemasaran. Upaya meningkatkan kepuasan pelanggan menjadi salah satu strategi membentuk kesetiaan pelanggan. Dengan citra baik yang diberikan oleh perusahaan akan memberikan dampak kepada niat beli dan membentuk kebiasaan membeli ulang. Sehingga sangat diperlukan optimalisasi pengiriman barang.

Optimisasi pengiriman barang adalah sebuah proses, dan cara (aktivitas/kegiatan) untuk mencari solusi terbaik dalam melakukan pengiriman barang sehingga tercipta nya suatu hasil yang efektif dari suatu proses yang telah dimaksimalkan [1]. Permasalahannya sekarang adalah sering kali penggolompokan barang biasanya didasarkan pada titik sesuai daerah dimana kurir harus mengantar barang tersebut. Tanpa mempertimbangkan kedekatan antara lokasi yang satu dengan yang lain. Yang sebenarnya dapat dimaksimalkan karena titik yang berdekatan. Hal ini biasa menyebabkan ketidakefisienan baik dari segi waktu ataupun biaya.

Maka dari itu, dibutuhkan sentuhan teknologi untuk mengoptimasi pengiriman barang. Penggunaan algoritma dan sistem berbasis machine learning dapat membantu perusahaan dalam merencanakan pengiriman yang lebih baik, mengelompokkan pengiriman berdasarkan kedekatan geografis. Tujuan dari penelitian ini, yaitu mengoptimalkan pengelompokkan pengiriman barang dengan mengimplementasi metode machine learning. Dengan demikian, penerapan strategi dan teknologi dalam optimisasi pengiriman barang menjadi sangat penting dalam menghadapi tantangan di era digital ini, dan dapat membawa perusahaan ke tingkat yang lebih tinggi dalam hal pelayanan dan kepuasan pelanggan

2. METODE PENELITIAN

Penelitian ini terdiri dari beberapa langkah atau tahapan yang dapat dilihat pada Gambar 1.



Gambar 1. Tahapan Penelitian

Langkah-langkah dalam penelitian ini, yaitu :

1. Identifikasi Masalah

Identifikasi masalah merupakan langkah awal dalam proses pemecahan sebuah masalah yang tujuannya untuk lebih mengenal dan mengartikan permasalahan yang perlu diatasi. Masalah yang saat ini terjadi yaitu tidak optimalnya penyebaran titik pengiriman.

2. Pengumpulan Data

Tahapan ini sangat penting yaitu mengumpulkan data dengan jelas dan informasinya harus sesuai dengan kebutuhan. Pada penelitian ini data didapatkan dari data pengiriman PT. Enseval Putera Megatrading, Tbk. Pada bulan Mei sampai Juli 2024 dengan jumlah kurang lebih 28.000 data. Pada Gambar 2, merupakan contoh dari data pengiriman yang akan digunakan.

SJ_LOADING_DATE	CARRIER_CODE	CARRIER_DESCRIPTION	WEIGHT_CAPACITY	VOLUME_CAPACITY	SHIP_TO_ORG_ID	LONGITUDE	LATITUDE	SHIPPED_QUANTITY	TOTAL_WEIGHT	TOTAL_VOLUME
05/20/2024 10:28:00	CB02	COLT L-300	1,160	5	3,367,684	-6.1357282	106.790511	24	7	0.0087
06/14/2024 10:27:21	CB02	COLT L-300	1,160	5	3,367,732	-6.1442151	106.7968015	3	1	0.0010
05/13/2024 15:59:08	CB01	MOTOR BOX	50	0	1,167,762	-6.1677488	106.8722069	1	0	0.0006
05/07/2024 16:26:40	CB01	MOTOR BOX	50	0	1,167,762	-6.1677488	106.8722069	15	1	0.0050
07/10/2024 15:36:31	CB01	MOTOR BOX	50	0	1,167,762	-6.1677488	106.8722069	80	2	0.0071
06/19/2024 12:16:00	CB02	COLT L-300	1,160	5	1,169,032	-6.2716835500000006	106.8603731	4	0	0.0007
07/10/2024 09:44:01	CB02	COLT L-300	1,160	5	1,168,686	-6.217996066666666	106.85379296666666	7	2	0.0032
07/23/2024 15:41:20	CB01	MOTOR BOX	50	0	1,170,328	-6.2333968	106.897259	52	4	0.0253
07/20/2024 13:02:53	CB02	COLT L-300	1,160	5	1,170,328	-6.2333968	106.897259	210	29	0.1409
07/26/2024 15:18:13	CB01	MOTOR BOX	50	0	1,170,328	-6.2333968	106.897259	54	2	0.0103
06/22/2024 13:05:18	CB02	COLT L-300	1,160	5	1,170,328	-6.2333968	106.897259	23	2	0.0081
06/26/2024 15:24:37	CB01	MOTOR BOX	50	0	1,170,328	-6.2333968	106.897259	12	0	0.0016
05/18/2024 12:21:07	CB01	MOTOR BOX	50	0	1,170,328	-6.2333968	106.897259	1	0	0.0007
07/22/2024 14:27:29	CB01	MOTOR BOX	50	0	3,323,788	-6.1616929	106.83421776666667	15	3	0.0067
06/18/2024 15:46:50	CB01	MOTOR BOX	50	0	3,323,788	-6.1616929	106.83421776666667	1	0	0.0011
06/20/2024 15:28:28	CB01	MOTOR BOX	50	0	3,323,788	-6.1616929	106.83421776666667	5	2	0.0031
05/03/2024 15:50:58	CB02	COLT L-300	1,160	5	3,326,734	-6.1950258	106.8908611	76	3	0.0197
05/31/2024 14:55:00	CB04	COLT 6 BAN	4,000	12	3,326,734	-6.1950258	106.8908611	409	19	0.1119
06/08/2024 12:16:48	CB02	COLT L-300	1,160	5	3,326,734	-6.1950258	106.8908611	10	2	0.0049

Gambar 2. Sample Data Pengiriman

3. Pengolahan Data

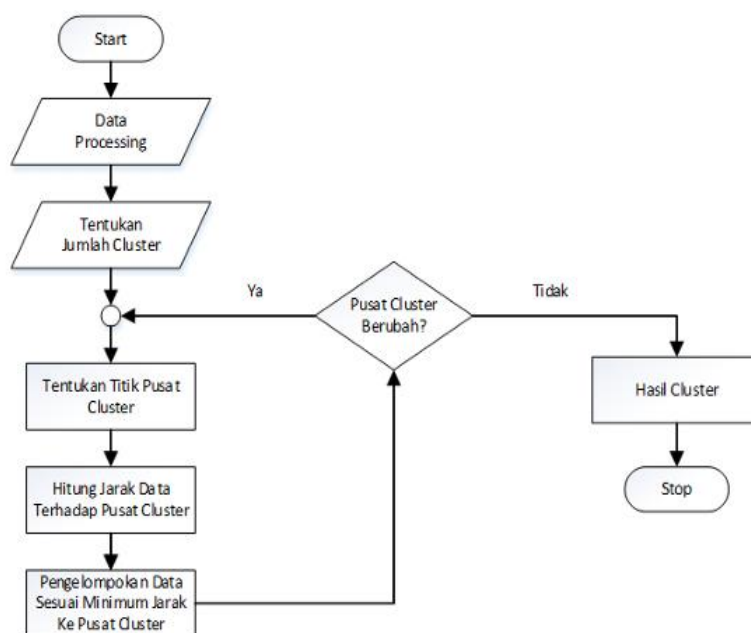
Data yang sudah dimiliki akan diolah pada tahap ini. Langkah pertama, kolom-kolom yang tidak digunakan akan dihapus dan untuk data yang memiliki *missing value* akan di hapus juga. Selanjutnya untuk kolom “SJ_LOADING_DATE” akan diubah menjadi datetime agar bisa mengakses data pada tanggal dan waktu tertentu. Lalu untuk kolom “LANGITUDE” dan “LATITUDE” akan di *scaling*. Kemudian data baru diolah menggunakan 2 algoritma yang berbeda, yaitu algoritma *K-Means Clustering*, dan *Hierarchical Clustering*. Penelitian ini menggunakan 2 algoritma yang berbeda dengan tujuan untuk membandingkan kinerja dari 2 algoritma yang sebelumnya sudah melalui tahap pemilihan.

4. Evaluasi

Pada tahap ini, hasil dari kedua algoritma yang digunakan berupa *sillhouette score* akan dibandingkan dan disimpulkan algoritma yang cocok atau sesuai dengan *case* dan data yang dimiliki.

2.1. Algoritma K-Means Clustering

K-Means Clustering merupakan salah satu algoritma yang masuk ke dalam *unsupervised learning*. Algoritma *K-Means Clustering* digunakan untuk mengelompokkan sejumlah data yang tidak memiliki label ke dalam sebuah *cluster* tanpa melalui proses *training* data terlebih dahulu [2]. Pengelompokkan data pada metode ini yaitu memilih data dengan karakteristik yang sama yang akan dikelompokkan ke dalam satu *cluster* yang sama (Purnamaningsih, Saptono, & Aziz, 2014).



Gambar 3. Flow Chart Algoritma K-Means Clustering [3].

Pada Gambar 3. Merupakan *flow chart* dari algoritma *K-Means Clustering*. Dalam penelitian ini, *K-Means Clustering* bekerja dengan baik jika jumlah *cluster* telah diketahui sebelumnya. Jadi, untuk jumlah *cluster* akan diinisialisasikan terlebih dahulu. Cara kerja dari algoritma ini dalam penelitian saat ini yaitu dengan memilih titik pusat (*centroid*) dari setiap *cluster* secara acak, kemudian dilakukan pengelompokkan data berdasarkan kedekatan jarak euclidean dari suatu titik ke *centroid* tersebut.

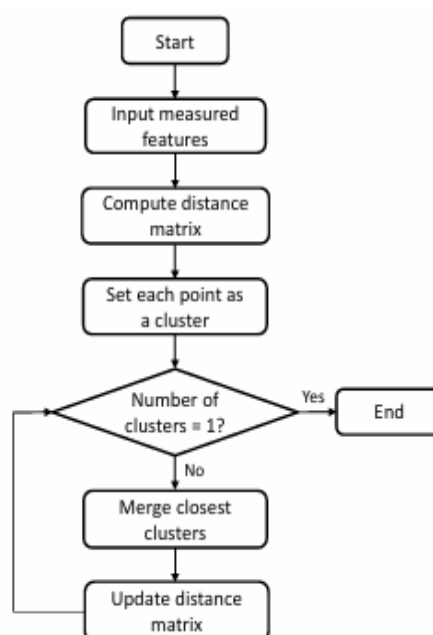
Beberapa keunggulan algoritma *K-Means Clustering* dibanding dengan algoritma *clustering* lainnya:

1. Efisiensi Waktu Komputasi
Proses iteratif sederhana dalam *K-Means* memungkinkan algoritma ini bekerja lebih cepat dibandingkan dengan metode *clustering* lainnya [4].
2. Kemudahan Implementasi dan Penyesuaian
K-Means adalah algoritma yang relatif mudah diimplementasikan dan disesuaikan, terutama untuk kasus-kasus seperti pengelompokkan titik pengiriman yang memiliki distribusi yang cukup seragam. Selain itu, *K-Means* dapat diperluas untuk *cluster* besar atau skenario distribusi dengan mudah, hanya dengan memilih nilai *K* yang sesuai [5].
3. Kemampuan Mengatasi Dataset Besar dengan Variasi Kecil
Dalam kasus pengiriman, data sering kali memiliki variasi kecil dalam hal lokasi geografis (jarak yang tidak terlalu jauh). *K-Means* lebih efisien dan efektif untuk menangani data seperti ini karena setiap titik dapat dikelompokkan dengan cepat tanpa perlu banyak iterasi [6].

2.2. Algoritma *Hierarchical Clustering*

Algoritma *Hierarchical Clustering* adalah metode analisis *clustering* yang bertujuan membangun struktur hirarki untuk data [7]. Pendekatan pengelompokannya terbagi menjadi dua strategi utama, yaitu *Agglomerative (Bottom-Up)* dan *Divisive (Top-Down)*. Pendekatan *Agglomerative*, di mana proses *clustering* dimulai dengan setiap data sebagai *cluster* individu yang kemudian secara bertahap digabungkan hingga membentuk satu kesatuan *cluster* besar.

Pendekatan *Divisive* dimulai dengan menganggap seluruh data sebagai satu *cluster* besar, kemudian secara bertahap membagi *cluster* tersebut menjadi sub-*cluster* yang lebih kecil hingga setiap data berdiri sebagai *cluster* tersendiri atau sesuai dengan jumlah *cluster* yang diinginkan. Proses ini dilakukan secara rekursif dengan memilih cara optimal untuk memisahkan data dalam setiap langkah pembagian.



Gambar 4. Flow Chart Algoritma *Hierarchical Clustering* [8]

Flow chart algoritma *Hierarchical Clustering* dapat dilihat pada Gambar 4. Dalam penelitian ini, algoritma *Hierarchical Clustering* dengan pendekatan *Agglomerative (Bottom-Up)* bekerja dengan memulai dari setiap titik pengiriman sebagai *cluster* terpisah, kemudian menggabungkan *cluster* yang paling berdekatan secara iteratif berdasarkan jarak seperti *Euclidean Distance* atau *Great Circle Distance*. Proses ini terus berlangsung hingga semua titik pengiriman tergabung dalam satu *cluster* besar. Struktur hierarki penggabungan ini divisualisasikan dalam bentuk dendrogram, yang menunjukkan hubungan bertingkat antar *cluster*. Dengan memotong dendrogram pada level tertentu, Anda dapat menentukan jumlah *cluster* yang diinginkan, memastikan bahwa setiap titik hanya berada dalam satu *cluster* tanpa tumpang tindih. Pendekatan ini sangat berguna untuk memahami pola geografis bertingkat pada titik pengiriman dan memungkinkan fleksibilitas dalam memilih jumlah *cluster* yang sesuai dengan kebutuhan.

Beberapa keunggulan algoritma *Hierarchical Clustering* dengan pendekatan *Agglomerative (Bottom-Up)* dibanding dengan algoritma *clusterring* lainnya:

1. Dapat Menangani Bentuk *Cluster* yang Tidak Beraturan
Agglomerative clustering tidak bergantung pada asumsi bentuk *cluster* tertentu. Oleh karena itu, metode ini cocok untuk data dengan distribusi yang kompleks atau berbentuk tidak beraturan, seperti data geografis [9].
2. Menghasilkan *Cluster* yang Stabil
Agglomerative clustering umumnya memberikan hasil yang stabil karena memulai dari *cluster* terkecil hingga membentuk *cluster* besar. Penggabungan bertahap ini membuat hasil akhir lebih konsisten, terutama jika data memiliki distribusi dan jarak antar *cluster* yang jelas [10].
3. Robust Terhadap *Outlier*
Karena tidak tergantung pada centroid atau pusat *cluster* tertentu, agglomerative clustering cenderung lebih robust terhadap *outlier* dibandingkan dengan algoritma berbasis *centroid*. *Outlier* dapat terpisah sebagai *cluster* tunggal atau bergabung dengan *cluster* lain tanpa banyak mempengaruhi hasil utama [11].

3. HASIL DAN PEMBAHASAN

Penelitian ini dijalankan dengan melakukan eksperimen sebanyak 10 kali, menggunakan jumlah *cluster* yang berbeda-beda. Di bawah ini, pada Tabel 1. Dapat dilihat perbandingan hasil *silhouette score* kesepuluh eksperimen dari kedua algoritma yang digunakan, yaitu Algoritma K-Means Clustering dengan Algoritma *Hierarchical Clustering* dengan pendekatan *Agglomerative (Bottom-Up)*.

Tabel 1. Perbandingan Nilai *Silhouette Score*

<i>Experiment</i>	Total Cluster	<i>Silhouette Score K-Means</i>	<i>Silhouette Score Hierarchical</i>
1	4	0.53	0.349
2	6	0.637	0.326
3	8	0.668	0.357
4	10	0.716	0.376
5	12	0.756	0.396
6	14	0.772	0.32
7	16	0.768	0.32
8	18	0.766	0.32
9	20	0.809	0.27
10	22	0.815	0.27

Berdasarkan hasil evaluasi menggunakan *Silhouette Score* untuk metode *K-Means Clustering* dan *Hierarchical Clustering*, pada eksperimen pertama dengan jumlah *cluster* 4, *Silhouette Score* untuk *K-Means Clustering* adalah 0.53, sementara untuk *Hierarchical Clustering* adalah 0.349. Pada

eksperimen kedua dengan jumlah *cluster* 6, *K-Means Clustering* menghasilkan *Silhouette Score* sebesar 0.637, sedangkan *Hierarchical Clustering* memiliki skor 0.326. eksperimen ketiga, dengan 8 *cluster*, menunjukkan peningkatan skor *K-Means Clustering* menjadi 0.668, dan *Hierarchical Clustering* sebesar 0.357.

Eksperimen keempat, dengan 10 *cluster*, memberikan *Silhouette Score* 0.716, dan *Hierarchical* sebesar 0.376, eksperimen kelima dengan jumlah *cluster* 12, *K-Means Clustering* mencapai skor sebesar 0.756, sementara *Hierarchical Clustering* memperoleh 0.396. Selanjutnya eksperimen keenam, dengan 14 *cluster* memberikan *Silhouette Score* 0.772, dan *Hierarchical* sebesar 0.32, eksperimen ketujuh, dengan 16 *cluster* memberikan *Silhouette Score* 0.768, dan *Hierarchical* sebesar 0.32, eksperimen kedelapan, dengan 18 *cluster* memberikan *Silhouette Score* 0.766, dan *Hierarchical* sebesar 0.32, eksperimen kesembilan, dengan 20 *cluster* memberikan *Silhouette Score* 0.809, dan *Hierarchical* sebesar 0.27. Terakhir eksperimen kesepuluh, dengan 22 *cluster* memberikan hasil *Silhouette Score* tertinggi yaitu 0.815, dan *Hierarchical* terendah sebesar 0.27. Dari hasil ini, terlihat bahwa metode *K-Means Clustering* secara konsisten menghasilkan *Silhouette Score* yang lebih tinggi dibandingkan *Hierarchical Clustering* untuk semua jumlah *cluster*, dengan skor terbaik tercapai pada 22 *cluster* dengan *K-Means Clustering* (0.815). Hal ini menunjukkan bahwa *K-Means* cenderung lebih optimal dalam menghasilkan *cluster* yang lebih jelas dan terpisah dibandingkan dengan *Hierarchical Clustering* pada data ini.

Tabel 2. Perbandingan Nilai *Davies-Bouldin Index* (DBI)

<i>Experiment</i>	Total Cluster	Davies-Bouldin Index K-Means	Davies-Bouldin Index Hierarchical
1	4	0.73	0.83
2	6	0.55	0.9
3	8	0.49	0.82
4	10	0.44	0.82
5	12	0.36	0.76
6	14	0.384	0.79
7	16	0.384	0.79
8	18	0.387	0.75
9	20	0.394	0.73
10	22	0.396	0.74

Pada Tabel 2. Dapat dilihat juga perbandingan nilai *Davies-Bouldin Index* (DBI) kesepuluh eksperimen dari kedua algoritma yang digunakan, yaitu Algoritma *K-Means Clustering* dengan Algoritma *Hierarchical Clustering* dengan pendekatan *Agglomerative (Bottom-Up)*. Berdasarkan hasil evaluasi *clustering* menggunakan *Davies-Bouldin Index* (DBI) untuk metode *K-Means Clustering* dan *Hierarchical Clustering*, didapatkan hasil sebagai berikut. Pada eksperimen pertama dengan jumlah *cluster* 4, DBI untuk *K-Means Clustering* adalah 0.73, sedangkan untuk *Hierarchical Clustering* adalah 0.83. Pada eksperimen kedua dengan jumlah *cluster* 6, DBI *K-Means Clustering* menurun menjadi 0.55, sementara DBI *Hierarchical Clustering* berada pada 0.9. Pada eksperimen ketiga dengan 8 *cluster*, DBI *K-Means Clustering* adalah 0.49 dan *Hierarchical Clustering* mencapai 0.82.

Selanjutnya, pada eksperimen keempat dengan jumlah *cluster* 10, DBI *K-Means Clustering* tercatat sebesar 0.44 sedangkan DBI untuk *Hierarchical Clustering* adalah 0.82. Pada eksperimen kelima dengan 12 *cluster*, DBI untuk *K-Means Clustering* adalah 0.36 yang merupakan DBI terendah, sementara DBI *Hierarchical Clustering* adalah 0.76. Eksperimen Keenam dengan 14 *cluster*, DBI *K-Means Clustering* adalah 0.384, sementara DBI *Hierarchical Clustering* adalah 0.79

Pada eksperimen ketujuh dengan jumlah *cluster* 16, DBI *K-Means Clustering* dan DBI *Hierarchical Clustering* tetap sama dengan eksperimen sebelumnya yaitu 0.384 untuk *K-Means Clustering* dan 0.79 untuk *Hierarchical Clustering*. Eksperimen kedelapan dengan jumlah *cluster* 18, DBI untuk *K-Means Clustering* adalah 0.387, sedangkan untuk *Hierarchical Clustering* adalah

0.75. Selanjutnya eksperimen kesembilan dengan jumlah *cluster* 20, DBI untuk *K-Means Clustering* adalah 0.394, sedangkan untuk *Hierarchical Clustering* adalah 0.73. Terakhir eksperimen kesepuluh dengan jumlah *cluster* 20, DBI untuk *K-Means Clustering* adalah 0.396, sedangkan untuk *Hierarchical Clustering* adalah 0.74.

Dari hasil ini juga, terlihat bahwa metode *K-Means Clustering* secara konsisten menghasilkan nilai *Davies-Bouldin Index* yang lebih rendah dibandingkan metode *Hierarchical Clustering* pada semua jumlah *cluster*, yang menunjukkan bahwa *K-Means* mampu menghasilkan *clustering* yang lebih kompak dan terpisah dengan lebih baik. Nilai DBI terendah dicapai oleh *K-Means Clustering* pada 12 *cluster* dengan nilai 0.36, yang menandakan konfigurasi *cluster* terbaik dalam eksperimen ini.

4. KESIMPULAN

Berdasarkan analisis menggunakan *Silhouette Score* dan *Davies-Bouldin Index* (DBI), algoritma *K-Means Clustering* secara konsisten menunjukkan hasil yang lebih baik dibandingkan *Hierarchical Clustering* pada semua jumlah *cluster* yang diuji. *K-Means Clustering* menghasilkan *Silhouette Score* tertinggi sebesar 0.815 pada 22 *cluster* dan *Davies-Bouldin Index* terendah sebesar 0.36 pada 12 *cluster*, yang menandakan bahwa *cluster* yang dihasilkan lebih jelas, kompak, dan terpisah dengan baik. Di sisi lain, *Hierarchical Clustering* memiliki *Silhouette Score* yang lebih rendah dan DBI yang lebih tinggi, menunjukkan bahwa *cluster* yang terbentuk cenderung kurang terpisah dan lebih tersebar. Kelebihan utama dari *K-Means Clustering* adalah kemampuannya menghasilkan *cluster* yang lebih baik dengan komputasi yang efisien, terutama pada dataset besar, sementara *Hierarchical Clustering* lebih lambat dan kurang optimal dalam kasus ini. Berdasarkan hasil ini, *K-Means Clustering* dengan 20 atau 22 *cluster* adalah pilihan yang optimal untuk data ini karena memberikan *cluster* yang lebih baik dalam hal kepadatan dan keterpisahan.

Untuk pengembangan selanjutnya, metode alternatif seperti menggabungkan *K-Means Clustering* dan *Hierarchical Clustering* melalui teknik *ensemble clustering* juga bisa menjadi opsi untuk meningkatkan hasil *clustering*. Selain itu, penggunaan *dimensionality reduction* seperti PCA atau t-SNE sebelum *clustering* mungkin membantu meningkatkan kualitas *cluster*, terutama untuk data berdimensi tinggi. Secara keseluruhan, *K-Means Clustering* lebih unggul untuk dataset ini, namun opsi pengembangan tersebut dapat dicoba untuk hasil yang lebih optimal.

UCAPAN TERIMA KASIH

Puji dan syukur penulis panjatkan kepada Tuhan Yang Maha Esa, oleh karena berkat, hikmat, dan anugerah-Nya, penulis dapat menyelesaikan jurnal ini dengan baik. Dalam setiap langkah perjalanan penulisan, penulis merasa diberkati dengan berbagai kemudahan dan inspirasi yang datang dari-Nya. Penulis juga ingin mengucapkan terima kasih yang sebesar-besarnya kepada keluarga tercinta, yang selalu memberikan dukungan moral dan emosional. Tanpa kasih sayang dan dorongan mereka, penulis tidak akan dapat menghadapi tantangan yang ada. Tak terlupakan, penulis sangat berterima kasih juga kepada pasangan yang selalu setia menemani dan memberikan semangat, serta sahabat-sahabat yang turut berkontribusi dalam proses penyusunan jurnal ini. Keterlibatan mereka tidak hanya memberikan motivasi, tetapi juga ide-ide berharga yang memperkaya isi jurnal ini. Harapan penulis ialah agar jurnal ini mampu memberikan manfaat penuh dan kontribusi yang positif bagi pembaca serta menjadi sumber inspirasi bagi mereka yang ingin mendalami topik ini lebih lanjut.

DAFTAR PUSTAKA

1. D. P. ILHAM, "Optimalisasi Sistem Manajemen Keselamatan Untuk Meminimalisir Terjadinya Kecelakaan Kerja Di Tug Boat Transko Murai Milik Pt. Pertamina Trans Kontinental," 2019
2. R. A. Aditia Yudhistira, "Pengelompokan Data Nilai Siswa Menggunakan Metode K-Means Clustering," *Journal of Artificial Intelligence and Technology Information (JAITI)*, 2023.
3. P. R. N. S. Ahmad Chusyairi, "Pengelompokan Data Puskesmas Banyuwangi Dalam Pemberian Imunisasi Menggunakan Metode K-Means Clustering," *Telematika*, vol. 12, 2019.
4. A. K. Jain, "Data clustering: 50 years beyond K-means," vol. 31, no. 8.
5. T. M. D. M. N. N. S. P. C. D. S. R. & W. A. Y. Kanungo, "An efficient k-means clustering algorithm: Analysis and implementation," *IEEE Transactions on Pattern Analysis and Machine Intelligence*.
6. M. E. K. H. A. & V. P. A. Celebi, "A comparative study of efficient initialization methods for the k-means clustering algorithm. Expert Systems with Applications".
7. A. T. R. D. S. W. Fachrian Bimantoro Putra, "Penerapan Algoritma Hierarchical Clustering Dalam Pengelompokan Kabupaten/Kota Di Papua Berdasarkan Indikator Kemiskinan," *Mathematics & Applications*, 2023.
8. A. S. R. S. E. F. C. S. J. C. W. A. C. Moises Silva, "Agglomerative concentric hypersphere clustering applied to structural damage detection," *ResearchGate*.
9. L. & M. O. Rokach, "Clustering methods. In Data Mining and Knowledge Discovery Handbook," *Springer*.
10. J. H. Ward, "Hierarchical grouping to optimize an objective function. Journal of the American Statistical Association".
11. A. K. & D. R. C. Jain, "Algorithms for clustering data".