



KLASTERISASI LAHAN SAWIT PADA SRIWIJAYA PALM OIL PALEMBANG MENGGUNAKAN ALGORITMA K-MEANS

Eni Andriani¹, Linda Zuni Alvi Nasution², Rezanía Agramanisti Azdy^{3*}

¹Program Studi Informatika, STMIK Palcomtech, Palembang

e-mail: *³rezania@palcomtech.ac.id

*Corresponding Author

Abstrak

Sriwijaya Palm Oil merupakan salah satu perusahaan sawit di Palembang yang memiliki beberapa cabang, salah satunya adalah PT Selatan Agro Makmur Lestari (SAML) yang terletak di Air Sugihan. Untuk melakukan kontrol lahan, hasil produksi harian pada cabang tersebut dimasukkan ke dalam Excel lalu dihitung menggunakan standar deviasi sehingga hasil produksi dapat dikelompokkan ke dalam kategori tertentu. Hal ini sudah cukup baik, namun pengelompokan lahan dengan cara seperti ini masih memerlukan waktu yang cukup banyak. Dengan perkembangan ilmu pengetahuan, hal ini dapat dilakukan dengan cara baru, seperti menggunakan algoritma K-Means dalam proses klusterisasi. Setelah dilakukan pengujian, klusterisasi menggunakan algoritma K-Means memiliki tingkat akurasi 83%. Selain itu, dengan penggunaan aplikasi ini, pekerja dapat lebih cepat mengetahui lahan sawit yang bermasalah.

Kata kunci— Data Mining, Klusterisasi Lahan, Algoritma K-Means.

Abstract

Sriwijaya Palm Oil is a palm oil corporation at Palembang which already has several branches in Sumatera Selatan. In one of its branches, PT Selatan Agro Makmur Lestari (SAML), which is located at Air Sugihan, the result of their unit production is calculated by Excel using standard deviation so that they could figure out which land produces good or bad unit production. This method could be considered good, but land grouping by this method is not efficient enough. With the latest science development, this case can be solved by building an application which used K-Means algorithm for its land clustering. After doing the testing, the result of land clustering by this algorithm hits 83% accuracy. Moreover, with the use of an application, it will make the workers find the land with problem easier.

Keywords— Data Mining, Land Clustering, K-Means Algorithm.



I. PENDAHULUAN

Perkebunan kelapa sawit merupakan salah satu lahan bisnis yang ditekuni oleh banyak perusahaan di Indonesia. Keuntungan yang dihasilkan dari produksi kelapa sawit tidak hanya menguntungkan bagi perusahaan, namun juga dalam perkembangan ekonomi nasional. Keuntungan ini hanya dapat diperoleh jika perkebunan kelapa sawit memberikan hasil yang maksimal. Kelapa sawit dapat memberikan hasil yang maksimal dengan memperhatikan beberapa faktor, yaitu bibit sawit yang baik dan perawatan lahan. Namun, kontrol lahan seperti melihat kondisi hasil panen pada blok-blok sawit juga berperan penting dalam memastikan produksi sawit tetap berjalan dengan baik. Selain menjadi estimasi hasil panen di hari berikutnya, kontrol lahan juga dapat membantu untuk memastikan apakah ada kesalahan pada blok-blok sawit tersebut.

Sriwijaya Palm Oil Group Palembang merupakan perusahaan yang bergerak di bidang perkebunan sawit yang terletak di Air Sugihan. Perusahaan ini memiliki 3 divisi lahan sawit dengan luas kurang lebih 400 – 750 Ha. Setiap divisi memiliki beberapa blok, dimana luas rata-rata dari setiap blok adalah 30 Ha. Setiap harinya, pengontrolan lahan dilakukan perblok, dimana dalam satu blok terdapat rata-rata 15 ancak (lorong) atau 30 baris.

Adapun jumlah pokok dalam setiap ancak adalah 250 pokok hidup. Hasil dari kontrol produksi harian ini nantinya dimasukkan ke dalam tabel *Excel* satu per satu, lalu dihitung menggunakan standar deviasi untuk nantinya dikategorikan menjadi beberapa zona yang ditandai dengan empat warna, yaitu merah, kuning, hijau, dan biru. Penggunaan *Excel* memang cukup baik dalam pengolahan data, namun pengelompokan lahan dengan cara seperti ini masih memerlukan waktu yang cukup banyak. Dengan perkembangan ilmu pengetahuan, hal ini dapat dilakukan dengan cara baru, seperti menggunakan metode klasterisasi dalam pengelompokan lahan sawit berdasarkan hasil produksinya.

Klasterisasi (*clustering*) atau lebih dikenal dengan analisis cluster merupakan

proses pengelompokan satu set benda fisik ataupun abstrak ke dalam satu kelas objek yang sama, sedangkan cluster merupakan kumpulan objek data yang memiliki kemiripan antar objek dalam kelompok yang sama dan berbeda objek dengan data kelompok lain (Irfiani dan Rani, 2018). *Clustering* merupakan sebuah metode mempartisi data ke dalam kelompok sehingga data yang memiliki karakteristik yang sama dikelompokkan ke dalam satu cluster yang sama (Priyatman et al, 2019). Metode yang dapat digunakan dalam proses *clustering* terbagi menjadi dua, yaitu algoritma Fuzzy C-Means dan algoritma K-Means. Algoritma Fuzzy C-Means adalah suatu teknik pengelompokan data yang mendasarkan pada derajat keanggotaan tiap-tiap data (Nugraha et al, 2017). Algoritma K-Means merupakan algoritma klasterisasi yang mengelompokan data berdasarkan titik pusat klaster (*centroid*) terdekat dengan data (Tendean dan Purba, 2020). Tujuan dari K-Means adalah pengelompokan data dengan memaksimalkan kemiripan data dalam satu klaster dan meminimalkan kemiripan data antar klaster. Metode ini banyak digunakan di berbagai bidang karena sederhana, mudah dipahami, memiliki kemampuan untuk mengklaster data yang besar, mampu menangani data outlier, dan kompleksitas waktunya linear terhadap jumlah dokumen atau data, jumlah klaster, dan jumlah iterasi (Sirait, 2017).

Pada penelitian ini, data-data yang digunakan dalam klasterisasi ditentukan terlebih dahulu. Setelah menentukan data masukan dan data yang menjadi pusat klaster, proses klasterisasi dijalankan menggunakan algoritma K-Means sehingga menghasilkan output berupa klaster yang terbagi menjadi beberapa zona lahan.

Terdapat beberapa penelitian lainnya yang menggunakan algoritma K-Means untuk klasterisasi. Rohmatullah et al (2019) melakukan penelitian klasterisasi Data Pertanian di Kabupaten Lamongan menggunakan Algoritma K-Means dan Fuzzy C-Means. Pada penelitian ini, hasil klasterisasi pada kabupaten Lamongan dapat memberikan informasi mengenai potensi pertanian di daerah tersebut

berdasarkan luas lahan pertanian dan hasil produksi pertanian. Penelitian lainnya dilakukan oleh Tendean dan Purba (2020) yang menggunakan algoritma K-Means untuk klasterisasi provinsi di Indonesia berdasarkan produksi bahan pangan. Hasil dari proses *clustering* pada penelitian ini dapat menjadi masukan untuk pemerintah dalam menganalisa provinsi mana saja yang memiliki jumlah produksi bahan pangan yang dominan terhadap produksi bahan pangan seperti jagung, kacang tanah, ubi kayu, kedelai, dan padi.

Pada penelitian ini dilakukan klasterisasi lahan sawit pada Sriwijaya Palm Oil Group Palembang dengan algoritma K-Means dengan hasil klasterisasi lahan terbagi menjadi 4 macam, yaitu zona merah, kuning, hijau, dan biru.

II. METODOLOGI PENELITIAN

Metodologi penelitian yang dilakukan peneliti adalah dengan menerapkan metode pengembangan perangkat lunak *prototype*. Menurut Dwi (2017), metode *prototype* adalah suatu metode pengembangan perangkat lunak yang berupa model fisik kerja aplikasi dan berfungsi sebagai versi awal dari aplikasi. *Prototype* berperan sebagai perantara pengembang dan pengguna agar dapat berinteraksi dalam proses kegiatan pengembangan aplikasi.

1. Analisis kebutuhan. Pada tahap ini peneliti melakukan studi kelayakan dan studi terhadap kebutuhan pengguna.
2. Data Selection dan Pre-processing Data. Pada tahap ini ditentukan variabel-variabel yang dibutuhkan dalam klasterisasi lahan berdasarkan tahap sebelumnya.
3. Desain *prototype*. Pada tahap ini penulis melakukan perencanaan dan pembangun aplikasi berdasarkan informasi yang diperoleh dari tahap sebelumnya.
4. Evaluasi dan perbaikan. Pada tahap ini, *prototype* yang telah selesai dibuat dievaluasi oleh pengguna, sedangkan pengembang bertugas menyesuaikan atau memperbaiki hasil dari evaluasi pengguna sampai *prototype* tersebut

memenuhi seluruh kebutuhan pengguna.

5. Implementasi. *Prototype* yang telah dibentuk beserta perubahan-perubahan yang dilakukan pada tahap evaluasi diimplementasikan ke dalam bentuk aplikasi.

Sumber data yang digunakan dalam penelitian ini adalah data primer dan data sekunder.

1. Data Primer

Data primer yang diperoleh dalam penelitian ini berupa data jumlah panjang dan berat panjang (kg) yang akan diolah untuk klasterisasi lahan sawit.

2. Data Sekunder

Data sekunder yang diperoleh penulis berupa data cabang, divisi, blok, dan lorong pada cabang SAML (Sungai Raden, Bukit Batu, dan Rengas Abang) dalam bentuk Excel.

Tahapan klasterisasi pada penelitian ini menggunakan algoritma K-Means (Rohmawati et al, 2015):

1. Melakukan normalisasi data pada variabel yang digunakan dalam proses klasterisasi. Adapun normalisasi yang digunakan adalah Min- Max Normalization. Min-Max Normalization merupakan metode normalisasi dengan melakukan transformasi linear terhadap data asli sehingga menghasilkan keseimbangan nilai perbandingan antar data saat sebelum dan sesudah proses (Nasution et al, 2019).

$$n' = \frac{n - \min_A}{\max_A - \min_A} \dots \dots \dots (1)$$

dimana:

- n' adalah data normalisasi,
 - n adalah data yang dinormalisasi,
 - $\min$ adalah nilai minimum dari data yang dinormalisasi, dan
 - $\max$ adalah nilai maksimum dari data yang dinormalisasi.
2. Menentukan jumlah klaster (k) pada data set sebagai nilai *centroid*.
 3. Menghitung jarak setiap data terhadap masing-masing *centroid* menggunakan *Euclidian Distance* hingga ditemukan jarak yang paling dekat dari setiap data



dengan *centroid*. Berikut adalah persamaan *Euclidian Distance*.

$$D_e = \sqrt{(x_i - s_i)^2 + (y_i - t_i)^2} \dots\dots(2)$$

dimana:

- D_e adalah *Euclidian Distance*,
 - i adalah banyaknya obyek,
 - (x, y) adalah koordinat objek, dan
 - (s, t) adalah koordinat *centroid*.
4. Mengelompokkan setiap data berdasarkan kedekatannya dengan jarak terdekat.
 5. Memperbaharui nilai *centroid* yang diperoleh dari rata-rata kluster yang bersangkutan dengan menggunakan persamaan berikut.

$$\mu_{j_n} = \frac{\sum x_{ji}}{Ns_{ji}} \dots\dots\dots(3)$$

dimana:

- μ_{j_n} adalah *centroid* baru pada iterasi ke n ,
 - $\sum x_{ji}$ adalah jumlah data pada kluster i , dan
 - Ns_{ji} adalah banyak data pada kluster i .
6. Melakukan perulangan dari langkah 3 hingga 5 sampai nilai titik pusat (*centroid*) tidak ada yang berubah.

Pengujian pada aplikasi ini dilakukan dengan cara mengukur tingkat akurasi antara algoritma K-Means dan perhitungan manual oleh Sriwijaya Palm Oil Group terhadap hasil dari lahan sawit. Menurut Indrayana, dkk (2016:4), rumus akurasi dapat dijabarkan sebagai berikut:

$$P = \frac{C}{R} \times 100\% \dots\dots\dots(4)$$

dimana

- P adalah persentase akurasi,
- C adalah jumlah data yang memiliki hasil yang sama dengan perhitungan manual oleh Sriwijaya Palm Oil Group, dan
- R adalah jumlah keseluruhan data.

III. HASIL DAN PEMBAHASAN

1. Pemilihan Data (*Data Selection*)

Dalam melakukan klusterisasi, variabel-variabel yang dipakai ditentukan terlebih dahulu. Setelah melakukan proses observasi dan wawancara pada para pembimbing lapangan Sungai Raden Estate, penulis mengambil kesimpulan bahwa pada

Tabel 1 terdapat variabel yang diperlukan dalam menentukan klusterisasi lahan sawit.

Tabel 1. Format Data Hasil Sawit

No	Cbg	Div	Blok	Lrg	Tgl	Jjg	Kg
1	BBE	B/AF D-1	B/1/W1 6	B/1/W 16/ PN011	2020-10-06	146	1419
2	BBE	B/AF D-1	B/1/W1 6	B/1/W 16/ PN011	2020-10-06	159	1545
3	BBE	B/AF D-1	B/1/W1 6	B/1/W 16/ PN011	2020-10-06	168	1633
4	BBE	B/AF D-1	B/1/W1 6	B/1/W 16/ PN011	2020-10-06	104	1011
...	...	...	...	...	...	...	...
3004	BBE	B/AF D-3	B/3/V0 6C	B/3/V0 6C /PN033	2020-11-01	121	1208
3005	BBE	B/AF D-3	B/3/T04 B	B/3/T0 4B/ PN033	2020-11-01	151	1253

2. Preprocessing

Dalam tahap ini, disimpulkan jumlah panjang dan berat panjang (kg) sebagai variabel yang diproses dalam perhitungan klusterisasi.

Tabel 2. Variabel Sampel yang akan Diproses

No	Blok	Jjg	Kg
1	B/1/W16	7215	70131
2	B/1/X16	7314	65897
3	B/1/V17	5150	46350
4	B/1/V19	3893	35813
5	B/1/W19	10191	79801
6	B/1/X17	9653	76260
7	B/1/W20	6419	49620
8	B/1/W21	5813	40053
...	...	...	...
21	B/1/W15	4399	37787
22	B/1/V16	2987	32738
23	B/1/V18	4797	44275
24	B/1/X19	9777	72054

Berdasarkan data yang diperoleh pada tahap pre-processing, jumlah panjang (unit) dan berat panjang (kg) merupakan variabel-variabel yang digunakan pada perhitungan kluster. Berikut langkah-



langkah perhitungan algoritma K-Means menggunakan sampel data seluruh blok pada divisi 1 cabang Bukit Batu.

a. Normalisasi Data

Pada Tabel 3 terdapat data-data yang akan dinormalisasi.

Tabel 3. Data Sampel

No	Blok	Jjg	Kg
1	B/1/W16	7215	70131
2	B/1/X16	7314	65897
3	B/1/V17	5150	46350
4	B/1/V19	3893	35813
5	B/1/W19	10191	79801
6	B/1/X17	9653	76260
7	B/1/W20	6419	49620
8	B/1/W21	5813	40053
9	B/1/X18	6684	56281
10	B/1/X20	8062	56755
11	B/1/W22	3759	30337
12	B/1/W23	1437	11236
13	B/1/V20	3608	33192
14	B/1/X21	4210	29679
15	B/1/W17	8372	69739
16	B/1/V21	5045	35315
17	B/1/X22	1042	7482
18	B/1/W18	8994	73754
19	B/1/X15	7195	63533
20	B/1/V15	3008	37393
21	B/1/W15	4399	37787
22	B/1/V16	2987	32738
23	B/1/V18	4797	44275
24	B/1/X19	9777	72054

Berdasarkan rumus (1), ditentukan masing-masing nilai minimum dan maksimum pada variabel janjang dan kg. Misal:

$$\text{janjang } W16' = \frac{\text{janjang} - \text{nilai janjang min}}{\text{nilai janjang maks} - \text{nilai janjang min}}$$

$$\text{janjang } W16' = \frac{7215 - 1042}{10191 - 1042} = 0.6747$$

Hasil normalisasi terhadap Data Sampel diperlihatkan pada Tabel 4.

Tabel 4. Normalisasi Data

No.	Blok	Jjg	Kg
1	B/1/W16	0.67471855	0.86628687
2	B/1/X16	0.6855394	0.8077407
3	B/1/V17	0.44901082	0.53745212

4	B/1/V19	0.31161876	0.39175044
5	B/1/W19	1	1
6	B/1/X17	0.94119576	0.95103638
7	B/1/W20	0.5877145	0.58266846
8	B/1/W21	0.52147776	0.45037957
9	B/1/X18	0.61667942	0.67477426
10	B/1/X20	0.76729697	0.68132856
11	B/1/W22	0.29697235	0.31603037
12	B/1/W23	0.04317412	0.0519089
13	B/1/V20	0.28046781	0.35550823
14	B/1/X21	0.34626735	0.30693179
15	B/1/W17	0.80118046	0.86086644
16	B/1/V21	0.43753416	0.38486428
17	B/1/X22	0	0
18	B/1/W18	0.86916603	0.91638435
19	B/1/X15	0.67253252	0.7750522
20	B/1/V15	0.21488687	0.41359809
21	B/1/W15	0.36692535	0.41904617
22	B/1/V16	0.21259154	0.34923049
23	B/1/V18	0.41042737	0.5087598
24	B/1/X19	0.95474915	0.89287739

b. Penentuan Kluster Awal

Empat macam kluster yang dibentuk ditentukan berdasarkan rentang nilai 0 – 1, yaitu 0.2, 0.4, 0.6, dan 0.8 (berurutan dari kluster terendah sampai kluster tertinggi) untuk masing-masing variabel janjang dan kg.

c. Perhitungan Jarak Data ke Kluster

Jarak data ke kluster 1, 2, 3, maupun 4 ditentukan menggunakan rumus (2). Misalkan saja jarak data W16 ke kluster 1 dapat dihitung seperti berikut:

$$\text{jarak } W16_e = \sqrt{(\text{janjang } W16' - 0.2)^2 + (\text{kg } W16' - 0.2)^2}$$

$$\text{jarak } W16_e = \sqrt{(0.674718 - 0.2)^2 + (0.866286 - 0.2)^2}$$

$$\text{jarak } W16_e = 0.818105$$

Tabel 5 merupakan jarak setiap variabel terhadap masing-masing kluster.

Tabel 5. Jarak Variabel Terhadap Kluster

No.	K1	K2	K3	K4	Jarak Terpendek
1	0.81810506	0.54119657	0.27657107	0.14173705	0.141737049
2	0.77787999	0.49778031	0.22466239	0.11472204	0.114722041



3	0.41938088	0.14592857	0.16343185	0.43832042	0.145928566
4	0.22187153	0.08876542	0.35571284	0.63654061	0.088765418
5	1.13137085	0.84852814	0.56568542	0.28284271	0.282842712
6	1.05519041	0.7723561	0.4895315	0.20675645	0.206756452
7	0.5447547	0.26192461	0.0212442	0.30380608	0.021244196
8	0.40747746	0.13151025	0.16897342	0.4470001	0.131510252
9	0.63169007	0.3499298	0.07661197	0.22200883	0.076611967
10	0.74397785	0.46265843	0.18601777	0.12309508	0.123095082
11	0.15121733	0.13291199	0.41528847	0.69804257	0.132911989
12	0.21569731	0.49848984	0.7813187	1.06415493	0.215697312
13	0.17509392	0.12754396	0.40233946	0.68372994	0.127543959
14	0.18118649	0.10746576	0.38764576	0.67006685	0.10746576
15	0.89339935	0.61101852	0.32943114	0.06087788	0.060877885
16	0.30099382	0.04047101	0.26958955	0.55110721	0.040471013
17	0.28284271	0.56568542	0.84852814	1.13137085	0.282842712
18	0.98030083	0.69768873	0.41539067	0.13538559	0.135385587
19	0.74429296	0.4636142	0.18948414	0.12988592	0.129885919
20	0.21411623	0.1856119	0.42785254	0.70118743	0.185611901
21	0.27540025	0.0381666	0.29507301	0.57678373	0.038166599
22	0.14976077	0.19416352	0.46148744	0.74043355	0.149760765
23	0.37364728	0.10925852	0.21038668	0.4864028	0.10925852
24	1.02456106	0.74207462	0.46002622	0.18048133	0.180481329

d. Pengelompokan Data dari Jarak Terdekat

Pengelompokan data ini dapat dilakukan dengan menghitung nilai paling kecil data pada masing-masing klaster. Berdasarkan hasil yang diperoleh pada Tabel 1, kita dapat menyimpulkan blok-blok yang dikelompokkan pada masing-masing klaster menjadi Tabel 6.

Tabel 6. Pengelompokan Blok dari Jarak Terdekat

	Nama Blok
Klaster 1	W23, X22, V16

Klaster 2	V17, V19, W21, W22, V20, X21, V21, V15, W15, V18
Klaster 3	W20, X18
Klaster 4	W16, X16, W19, X17, X20, W17, W18, X15, X19

e. Inisialisasi Klaster Baru

Apabila pengelompokan hasil klaster masih berubah, maka akan dilakukan lagi perhitungan klaster pada iterasi selanjutnya. Pada tahap ini, nilai klaster yang digunakan tidak dalam rentang 0 - 1 seperti sebelumnya, akan tetapi ditentukan dengan menggunakan rumus pada persamaan (3).

$$\text{klaster 1 baru} = \frac{\text{jumlah data pada Klaster 1}}{\text{banyak data pada Klaster 1}}$$

$$\text{klaster 1 baru} = \frac{W23 + X22 + V16}{3}$$

$$\text{klaster 1 baru} = \frac{0.21569731 + 0.28284271 + 0.14976077}{3}$$

$$\text{klaster 1 baru} = 0.085255219$$

Nilai inisialisasi klaster baru dapat dilihat pada Tabel 7.

Tabel 7. Inisialisasi Klaster Baru

K1	K2	K3	K4
0.085255219	0.363558859	0.602196961	0.818486538
0.133713132	0.408432086	0.62872136	0.861285877

Setelah dilakukan inisialisasi klaster baru, dilakukan perhitungan jarak data ke klaster seperti tahap 3. Inisialisasi klaster baru ini akan berhenti dilakukan ketika pengelompokan klaster tidak berubah lagi. Pada pengujian sampel ini, proses iterasi dilakukan sebanyak 5 kali dengan hasil klaster sebagai berikut.

Tabel 8. Pengelompokan Klaster pada Iterasi Akhir

	Nama Blok
Klaster 1	W23, X22
Klaster 2	V17, V19, W21, W22, V20, X21, V21, V15, W15, V16, V18
Klaster 3	W16, X16, W20, X18, X20, X15
Klaster 4	W19, X17, W17, W18, X19

f. Pengujian Akurasi



Setelah melakukan pengujian fungsional, seluruh komponen pada aplikasi dapat digunakan dengan baik. Selain pengujian fungsional, dilakukan pengujian persentase akurasi berdasarkan perhitungan manual menggunakan algoritma K-Means dan perhitungan manual yang dilakukan

oleh pihak Sriwijaya Palm Oil Group (SPOG) Palembang.

Tabel 9 merupakan perhitungan manual yang dilakukan oleh pihak Sriwijaya Palm Oil Group serta perbandingan hasilnya dengan klusterisasi menggunakan algoritma K-Means.

Tabel 9. Perhitungan Manual Janjang oleh Pihak SPOG

No.	Blok	Jjg	Avg	Std	Min	Med	Max	Hasil Klaster oleh SPOG	Hasil Klaster dengan K-Means
1	W16	7215	5792.67	2602.34	3190.33	5792.67	8395.00	7215	K3
2	X16	7314	5792.67	2602.34	3190.33	5792.67	8395.00	7314	K3
3	V17	5150	5792.67	2602.34	3190.33	5792.67	8395.00	5150	K2
4	V19	3893	5792.67	2602.34	3190.33	5792.67	8395.00	3893	K2
5	W19	10191	5792.67	2602.34	3190.33	5792.67	8395.00	10191	K4
6	X17	9653	5792.67	2602.34	3190.33	5792.67	8395.00	9653	K4
7	W20	6419	5792.67	2602.34	3190.33	5792.67	8395.00	6419	K3
8	W21	5813	5792.67	2602.34	3190.33	5792.67	8395.00	5813	K2
9	X18	6684	5792.67	2602.34	3190.33	5792.67	8395.00	6684	K3
10	X20	8062	5792.67	2602.34	3190.33	5792.67	8395.00	8062	K3
11	W22	3759	5792.67	2602.34	3190.33	5792.67	8395.00	3759	K2
12	W23	1437	5792.67	2602.34	3190.33	5792.67	8395.00	1437	K1
13	V20	3608	5792.67	2602.34	3190.33	5792.67	8395.00	3608	K2
14	X21	4210	5792.67	2602.34	3190.33	5792.67	8395.00	4210	K2
15	W17	8372	5792.67	2602.34	3190.33	5792.67	8395.00	8372	K4
16	V21	5045	5792.67	2602.34	3190.33	5792.67	8395.00	5045	K2
17	X22	1042	5792.67	2602.34	3190.33	5792.67	8395.00	1042	K1
18	W18	8994	5792.67	2602.34	3190.33	5792.67	8395.00	8994	K4
19	X15	7195	5792.67	2602.34	3190.33	5792.67	8395.00	7195	K3
20	V15	3008	5792.67	2602.34	3190.33	5792.67	8395.00	3008	K2
21	W15	4399	5792.67	2602.34	3190.33	5792.67	8395.00	4399	K2
22	V16	2987	5792.67	2602.34	3190.33	5792.67	8395.00	2987	K2
23	V18	4797	5792.67	2602.34	3190.33	5792.67	8395.00	4797	K2
24	X19	9777	5792.67	2602.34	3190.33	5792.67	8395.00	9777	K4

Berdasarkan hasil yang didapatkan pada Tabel 9, tingkat akurasi pada aplikasi klusterisasi dapat dihitung dengan menggunakan rumus (4).

$$P = \frac{\text{jumlah data yang benar}}{\text{jumlah keseluruhan data}} \times 100\%$$

$$P = \frac{20}{24} \times 100\% = 83\%$$

Jadi, tingkat akurasi aplikasi klusterisasi lahan sawit dengan menggunakan algoritma K-Means adalah 83%.

IV. KESIMPULAN

Berdasarkan pembahasan pada bab-bab sebelumnya terhadap Aplikasi Klusterisasi Lahan Sawit pada Sriwijaya Palm Oil Group Palembang menggunakan Algoritma K-means, aplikasi dapat



digunakan pekerja untuk mengolah data hasil produksi serta mengklasterisasikannya menjadi 4 klaster, yaitu klaster 1 (zona merah), klaster 2 (zona kuning), klaster 3 (zona hijau), dan klaster 4 (zona biru). Berdasarkan hasil pengujian, aplikasi dapat digunakan dengan baik dan memiliki tingkat akurasi sebesar 83%.

V. SARAN

Saran terhadap hasil penelitian dan pengembangan penelitian selanjutnya adalah:

1. Perusahaan dapat mengetahui pekerja mana yang bertanggung jawab atas lahan sawit yang bermasalah berdasarkan klaster hasil sawit oleh aplikasi.
2. Hasil klasterisasi dapat dikembangkan menggunakan *mapping*.



VI. DAFTAR PUSTAKA

- Irfiani, Eni dan Siti Sulistia Rani. 2018. Algoritma K-Means Clustering untuk Menentukan Nilai Gizi Balita. *Jurnal Sistem dan Teknologi Informasi*, 6(4): 165 – 172.
- Nugraha, Faizal Widya, Silmi Fauziati, dan Adhistya Erna Permanasari. 2017. Sistem Pendukung Keputusan Pemilihan Varietas Kelapa Sawit dengan Metode Fuzzy C-Means. *Seminar Nasional Inovasi dan Aplikasi Teknologi Industri*.
- Priyatman, Hendro, Fahmi Sajid, dan Dannis Haldivany. 2019. Klasterisasi Menggunakan Algoritma K-Means Clustering untuk Memprediksi Waktu Kelulusan Mahasiswa. *Jurnal Edukasi dan Penelitian Informatika*, 5(1): 62 – 66.
- Purnomo, Dwi. 2017. Model Prototyping pada Pengembangan Sistem Informasi. *Jurnal Informatika Merdeka Pasuruan*, 2(2): 54 – 61.
- Rohmatullah, Arif, Dinita Rahmalia, dan Mohammad Syaiful Pradana. 2019. Klasterisasi Data Pertanian di Kabupaten Lamongan menggunakan Algoritma K-Means dan Fuzzy C-Means. *Jurnal Ilmiah Teknosains*, V(2): 86 – 93.
- Rohmawati, Nurul, Sofi Defiyanti, dan Mohamad Jajuli. 2015. Implementasi Algoritma K-Means dalam Pengklasteran Mahasiswa Pelamar Beasiswa. *Jurnal Ilmiah Teknologi Informasi Terapan*, 1(2): 62 – 68.
- Sirait, Nenci. 2017. Implementasi K-Means Clustering pada Pengelompokan Mutu Biji Sawit. *Jurnal Pelita Informatika*, 6(2): 170 – 174.
- Tendean, Tonny, dan Windania Purba. 2020. Analisis Cluster Provinsi Indonesia berdasarkan Produksi Bahan Pangan menggunakan Algoritma K-Means. *Jurnal Sains dan Teknologi*, 1(2): 5 – 11.