

PERBANDINGAN ALGORITMA SUPPORT VECTOR MACHINE DAN K-NEAREST NEIGHBORS PADA KEMATANGAN BUAH SAWIT

Syawaluddin Kadafi Parinduri ¹⁾, Rika Rosnelly ²⁾, Anton Purnama ³⁾Ameliana Sihotang ⁴⁾, Mimi Chintya Adelina⁵⁾

^{1), 2), 3), 4), 5)}Universitas Potensi Utama

Email : parindurikadafi30@gmail.com ¹⁾, rika@potensi-utama.ac.id ²⁾, antonpurnama515@gmail.com ³⁾, amelianasihotang123@gmail.com⁴⁾, mimichintya8@gmail.com⁵⁾,

ABSTRAK

Berdasarkan pengamatan dan hasil dari *observasi*, buah kelapa sawit memiliki suatu warna buah yang hampir sama yaitu berwarna hitam pekat atau hitam agak kekuning-kuningan saat mentah, dan berwarna merah tua saat matang. Sangat sulit untuk membedakan buah kelapa sawit yang matang dan mentah. Tandan buah kelapasawit memiliki jumlah buah yang banyak, dalam satu tandan diperkirakan beratnya mencapai kurang lebih 20 sampai 30 Kilogram. Untuk dapat mengetahui kematangan buah sawit tersebut, dibutuhkan suatu sistem untuk melakukan klasifikasi kematangan buah secara otomatis. Metode *Support Vector Machine (SVM)* dan *K-Nearest neighbors (K-NN)* dapat digunakan untuk klasifikasi buah kelapa sawit yang matang dan mentah. Penelitian ini bertujuan untuk membandingkan model prediksi *Support Vector Machine* dengan *K-Nearest Neighbor* dalam memprediksi kematangan buah sawit. Dalam penelitian ini, model *Support Vector Machine* dan *K-Nearest Neighbor* disajikan dalam dataset dari 40 data *Image*, yang terdiri dari 35 *Image* data Uji, dan 34 *Image* data Latih, Data latih yang digunakan berjumlah 34 data *image* dengan 2 categories. Data uji yang digunakan berjumlah 6 data *image*. Data latih dan data uji akan dilakukan proses mengekstraksi fitur gambar menggunakan *Inception-V3*. Setelah ekstraksi fitur gambar dilakukan, maka dilaksanakan hitungan metode *Support Vector Machine (SVM)* dan *K-Nearest neighbors (K-NN)*. Hasil evaluasi menunjukkan bahwa kedua model prediksi kematangan buah sawit memiliki nilai akurasi yang sama. *Support Vector Machine (SVM)* dan *K-Nearest neighbors (K-NN)* yaitu 0.500 dan 0.500. Hal tersebut menunjukkan bahwa kedua metode tersebut sama baiknya.

. Penggunaan model *Support Vector Machine* dan *K-Nearest Neighbor* digunakan untuk memperoleh hasil yang relevan serta akurat dalam prediksi struktur sekunder. Kedua Metode ini akan digunakan untuk melihat kelebihan akurasi tertinggi. Sehingga Kedua metode ini, akan dibandingkan. dan bekerja baik dengan ruang dimensi yang tinggi dengan menggunakan bantuan aplikasi *Orange data mining*. Hasil yang diperoleh pada metode *Support Vector Machine (SVM)* skenario satu, mendapatkan nilai akurasi yang sangat baik, yaitu 100%. Pada skenario dua, dengan menggunakan metode *K-Nearest neighbors (K-NN)* mendapatkan nilai akurasi yang sangat baik juga sebesar 100%.. Hal ini membuktikan bahwa kedua metode tersebut dapat digunakan untuk mengklasifikasikan kematangan buah kelapa sawit dengan hasil yang sangat baik.

Kata Kunci : *Support Vector Machine (SVM), K-Nearest Neighbor (KNN), buah sawit, Orange data mining*

ABSTRACT

Bunaithe ar thuairimí agus torthaí ó bharúlacha, tá beagnach an dath torthaí céanna ag torthaí pailme ola, is é sin dubh dorchá nó beagán buí dubh nuair a bhíonn siad amh, agus dorchá dearg nuair a bhíonn siad aibí. Tá sé an-deacair idirdhealú a dhéanamh idir torthaí pailme aibí agus neamhaibí. Tá líon mór torthaí ag bunches torthaí pailme ola, meastar go bhfuil bun amháin ag meáchan thart ar 20 go 30 cileagram. Chun a bheith in ann aibíocht na dtorthaí pailme a chinneadh, tá gá le córas chun aibíocht na dtorthaí a rangú go huathoibríoch. Is féidir modhanna Meaisín Veicteoir Tacaíochta (SVM) agus K-Nearest Neighbors (K-NN) a úsáid chun torthaí pailme ola aibí agus neamhaibí a rangú. Tá sé mar aidhm ag an taighde seo an

tsamhail tuar Meaisín Veicteoir Tacaíochta a chur i gcomparáid le K-Nearest Neighbour maidir le haibíocht torthaí pailme a thuar. Sa taighde seo, cuirtear an Meaisín Vector Tacaíochta agus na samhlacha K-Nearest Neighbour i láthair i tacar sonraí de 40 sonraí Íomhá, comhdhéanta de 35 sonraí Íomhá Tástála, agus 34 sonraí Íomhá Oiliúna. Is iad na sonraí oiliúna a úsáidtear ná 34 sonraí íomhá le 2 chatagóir. Is iad na sonraí tástála a úsáidtear ná 6 shonraí íomhá. Próiseálfar na sonraí oiliúna agus na sonraí tástála chun gnéithe íomhá a bhaint as Inception-V3. Tar éis eastóscadh gné íomhá a dhéanamh, ríomhtar na modhanna Meaisín Veicteoir Tacaíochta (SVM) agus K-Nearest Neighbours (K-NN). Léiríonn na torthaí meastóireachta go bhfuil na luachanna cruinne céanna ag an dá mhúnla tuar aibíochta torthaí pailme. Is é 0.500 agus 0.500 an Meaisín Veicteoir Tacaíochta (SVM) agus K-na comharsana is gaire (K-NN) K-NN. Léiríonn sé seo go bhfuil an dá mhodh chomh maith céanna.

. Úsáidtear múnlaí Meaisín Veicteoir Tacaíochta agus K-Nearest Neighbour chun torthaí ábhartha agus cruinne a fháil i dtuar struchtúir tánaisteacha. Bainfear úsáid as an dá mhodh seo chun na buntáistí a bhaineann le cruinneas is airde a fheiceáil. Mar sin déanfar comparáid idir an dá mhodh seo. agus oibríonn sé go maith le spásanna ardtoiseacha ag baint úsáide as cabhair ó fheidhmchlár mianadóireachta sonraí Orange. Fuair na torthaí a fuarthas sa mhodh Meaisín Veicteoir Tacaíochta (SVM) i gcás a haon, luach cruinneas an-mhaith, is é sin 100%. I gcás a dó, ag baint úsáide as an modh K-Nearest Neighbours (K-NN), fuair muid luach cruinneas an-mhaith de 100%. Cruthaíonn sé seo gur féidir an dá mhodh a úsáid chun aibíocht torthaí pailme ola a rangú le torthaí an-mhaith.

Keywords: Support Vector Machine (SVM), K-Nearest Neighbor (KNN), palm fruit, data mining people

1. PENDAHULUAN

Buah sawit merupakan salah satu tanaman produksi dibidang perkebunan, yang termasuk taman komoditi di Indonesia. Produksi tanaman sawit di Indonesia termasuk komoditas ekspor yang cukup tinggi. Indonesia menempati urutan pertama sebagai penghasil kelapa sawit terbesar di dunia. Pada tahun 2022 Indonesia tercatat menghasilkan 45,5 juta ton CPO per tahunnya, dengan luas perkebunan kelapa sawit seluas 16,38 juta ha. (Sumber: <https://indonesiabaik.id/infografis/indonesia-produsen-minyak-sawit-terbesar-dunia#:~:text=Produksi%20Minyak%20Sawit,minyak%20sawit%20terbesar%20di%20dunia>).

Gambar 1. Urutan Indonesia sebagai Penghasil Minyak Sawit



Penelitian ini berkaitan dengan prediksi kematangan buah sawit dengan menggunakan dua Algoritma pembandingan yaitu dengan menggunakan Algoritma Support Vector Machine (SVM) dan Algoritma K-Nearest Neighbor (KNN).

Metode Support Vector Machine (SVM), yaitu merupakan sistem pembelajaran dengan menggunakan ruang hipotesis yang berupa fungsi-fungsi linear didalam sebuah fitur yang memiliki dimensi tinggi dan dilatih menggunakan algoritma pembelajaran yang berdasarkan teori optimasi (Monika Parapat & Tanzil Furqon, 2018).

K-Nearest Neighbor Algoritma K-Nearest Neighbor (KNN) merupakan sebuah metode

untuk melakukan klasifikasi terhadap objek berdasarkan data pembelajaran yang jaraknya paling dekat dengan objek tersebut (Raysyah et al., 2021)

2. METODE

Orange Data Mining adalah perangkat lunak *open-source* yang dirancang untuk menganalisis dan melakukan pengembangan data dengan pendekatan yang *intuitif* dan mudah digunakan. Aplikasi ini memungkinkan pengguna dari berbagai latar belakang, termasuk ilmuwan data, peneliti, mahasiswa, dan pengguna pemula, untuk memproses dan memahami data mereka dengan lebih *efisien*.

2.1. Data Mining Process

Hasil simulasi 2 model prediksi, data yang telah melalui tahap preprocessing selanjutnya dilakukan pengujian untuk mendapatkan model prediksi terbaik. Model prediksi yang telah dilakukan pengujian dan evaluasi dengan menggunakan kumpulan data uji pada aplikasi *orange* dimana 1 atribut sebagai targetnya.

Model prediksi buah matang dapat dianalisa menggunakan *orange tool* untuk memilih metode terbaik. Model prediksi *dataimage*.

Setelah dilakukan preprocessing data, maka selanjutnya *dataimage* diproses kedalam model prediksi pada software *orange* dengan menggunakan metode K-NN dan SVM.

2.2. Support Vector Machine (SVM)

Support Vector Machine (SVM) adalah pengklasifikasi *diskriminatif* yang menghasilkan *hyperplane* pemisah. Toleransi kesalahan disertakan untuk membuat *hyperplane* pemisah menjadi kuat jika terjadi data kelas yang tidak dapat dipisahkan.

Dalam *Algoritma SVM* ada trik *kernel* dimana ada *SVM linear* dan *SVM nonlinear*. *SVM hyperplane linear* bekerja hanya pada data yang dapat dipisahkan dengan cara *linear*. *SVM Non Linear* yaitu data yang berdistribusi pada kelas yang tidak linear sering digunakan pendekatan *kernel* pada fitur awal set (Krisandi et al., 2013).

Kernel dapat diartikan sebagai suatu fungsi yang memetakan fitur data yang memiliki dimensi awal rendah fitur lainnya

yang berdimensi lebih tinggi bahkan jauh lebih tinggi.

2.3. K-NN

K-Nearest Neighbor (K-NN) Algoritma K-NN memiliki langkah langkah sebagai berikut :

- Tentukan parameter K (jumlah banyaknya tetangga terdekat)
- Hitung jarak antara sample data uji dan seluruh sample data penelitian
- Urutkan jaraknya Ambil tetangga terdekat
- Kumpulkan/tentukan kategori tetangga terdekat
- Gunakan mayoritas kategori sederhana/tentukan kategori yang paling sering muncul (mayoritas) sebagai nilai prediksi dari data baru.

Ada banyak cara untuk mengukur jarak kedekatan antara data baru dengan data lama, diantaranya euclidean distance dan manhattan distance (*city block distance*), yang paling sering digunakan adalah *euclidean distance* (Ichsan et al., 2022).

2.4. AUC (Area Under the Curve)

AUC mengukur sejauh mana model mampu membedakan antara dua kelas (misalnya, kelas positif dan kelas negatif) berdasarkan skor prediksinya. AUC adalah ukuran dari area di bawah kurva Receiver Operating Characteristic (ROC), yang menampilkan perbandingan antara true positive rate (sensitivitas) dan false positive rate (1 - spesifisitas) pada berbagai ambang batas prediksi. Semakin tinggi nilai AUC, semakin baik model dalam membedakan kelas-kelas tersebut. (Ichsan et al., 2022)

2.5. CA (Accuracy):

Akurasi adalah ukuran persentase prediksi yang benar dari keseluruhan jumlah sampel. Dinyatakan sebagai: $(TP + TN) / (TP + TN + FP + FN)$, di mana TP adalah True Positive (jumlah sampel positif yang diprediksi benar), TN adalah *True Negative* (jumlah sampel negatif yang diprediksi benar), FP adalah False Positive (jumlah sampel *negatif* yang salah diprediksi sebagai positif), dan FN adalah *False Negative* (jumlah sampel positif yang salah diprediksi sebagai negatif). Semakin tinggi nilai akurasi, semakin baik kinerja model.

2.6. F1 Score

F1 score adalah ukuran keseimbangan antara precision dan recall. Precision (presisi) adalah perbandingan antara true positive dengan total prediksi positif ($TP / (TP + FP)$), sementara *recall* (*sensitivitas*) adalah perbandingan antara *true positive* dengan total jumlah sampel positif sebenarnya ($TP / (TP + FN)$). F1 score dinyatakan sebagai $2 * (precision * recall) / (precision + recall)$ dan memberikan keseluruhan ukuran kinerja model dalam mengidentifikasi kedua kelas secara seimbang.

2.7. Precision (Presisi)

Precision mengukur sejauh mana prediksi positif yang sebenarnya benar. Semakin tinggi nilai precision, semakin sedikit false positive yang terjadi.

2.8. Confusion matrix

Confusion matrix adalah matriks yang cukup *intuitif* dan mudah untuk mengetahui tingkat ketepatan dan akurasi dari model yang dihasilkan. Gambar 2 menunjukkan ilustrasi dari confusion matrix. Gambar 2 Confusion Matrix Berdasarkan confusion matrix, banyaknya total prediksi yang benar ditunjukkan oleh variabel TP (*True Positive*) dan TN (*True Negative*). Sedangkan banyaknya prediksi yang salah ditunjukkan oleh variabel FP (*False Positive*) dan FN (*False Negative*). Adapun

perhitungan indikator performa diperoleh dengan menghitung nilai accuracy, precision, recall, dan F1-Score. yang masing-masing ditunjukkan oleh persamaan (2), (3), (4), dan (5) berikut:

$$Accuracy = \frac{TP + TN}{TP + FP + FN + TN} \quad (2)$$

$$Precision = \frac{TP}{TP + FP} \quad (3)$$

$$Recall = \frac{TP}{TP + FN} \quad (4)$$

$$F1 - Score = \frac{2 * Precision * Recall}{Precision + Recall} \quad (5)$$

TP = banyaknya data dari kelas positive (1) yang benar diprediksi sebagai kelas

positive (1).

TN= banyaknya data dari kelas negative (0) yang benar diprediksi sebagai kelasnegative (0).

FP = banyaknya data dari kelas negative (0) yang salah diprediksi sebagai kelas positive (1).

FN= banyaknya data dari kelas positive (1) yang salahdiprediksi sebagai kela negative (0).

Dengan menggunakan antarmuka grafis yang mudah dipahami, *Orange* memungkinkan pengguna untuk membangun aliran kerja analisis data yang kompleks tanpa perlu menulis kode pemrograman. Pengguna dapat memanfaatkan berbagai algoritma analisis data yang telah disediakan dalam aplikasi ini, termasuk pemodelan *prediktif*, *klustering*, dan *visualisasi* data.

Tujuan dari penelitian ini untuk melakukan analisa perbandingan metode KNN dan SVM dalam memprediksi kematangan buah sawit, dalam penelitian ini dilakukan proses pengumpulan data, proses pengumpulan data pada penelitian ini menggunakan data set public. *Dataset public* yang dilakukan adalah dengan mengunduh (*download*) gambar terkait dari *google* yaitu berupa gambar citra buah sawit yang ada pada *variabel* penelitian ini. Dengan mengunduh Satu-Satu atau *Scraping*, dengan cara pengambilan data dari Web secara manual. Selanjutnya penelitian melakukan tahap pembuatan *data set* yaitu dengan melakukan *preprocessing* dimana tahapan yang di lakukan pada *preprocessing* kali ini adalah mengubah ukuran citra menjadi 500x500 *pixel*, mengubah warna *background*, dan mengubah jenis file citra menjadi JPG. Dari 20 data citra, hasil dari *preprocessing* tersebut terdapat 2 data yaitu data latih atau data *training* dan data uji atau data *testing*. Dimana terdapat sebanyak 17 data citra untuk data latih dan 3 data citra untuk data uji dan dikategorikan menjadi 2 kelas yaitu mentah dan matang. Setelah mendapatkan data citra buah sawit yang akan di uji selanjutnya peneliti membuat sistem yang akan digunakan untuk mengklasifikasikan kematangan buah sawit berdasarkan deteksi warna.

Tahapan terakhir yaitu kesimpulan yang berisi rangkuman dari hasil pengujian sistem, dan rangkuman dari hasil akhir dari penelitian yang di lakukan, bisa berupa keakurasian dari sistem klasifikasi yang sudah kita buat serta hasil akurasi dari metodologi yang di gunakan.

2.1. Perancangan Penelitian

Perancangan penelitian yang dilakukan dengan menggunakan *Orange Tools* untuk simulasinya.ditunjukkan oleh Gambar 2 sebagai berikut :

Gambar 2. Tahapan Penelitian



Tahapan pada penelitian ini terdiri dari 4 tahap yaitu pengumpulan data image, preprocessing, pembuatan data latih dan data uji, proses prediksi kematangan buah sawit, evaluasi kinerja dan hasil perbandingan metode.

2.2. Preprocessing Data

Data *image* yang digunakan merupakan data *image* kematangan buah sawit dengan fitur yang digunakan berupa data latih dan data Uji terdiri dari 34 *Image* ,pada penelitian ini terlihat pada Tabel 1,2, dan 3 antara lain :

No.	Matang	Mentah
1		
2		
3		
4		
---	---	---
36		

Tabel .1 Data Latih Dan Data Uji

No.	Matang	Mentah
1		
2		
3		
4		
---	---	---
34		

Tabel 2. Data Latih

No.	Matang	Mentah
35		
---	---	---
40		

Tabel 3. Data Uji

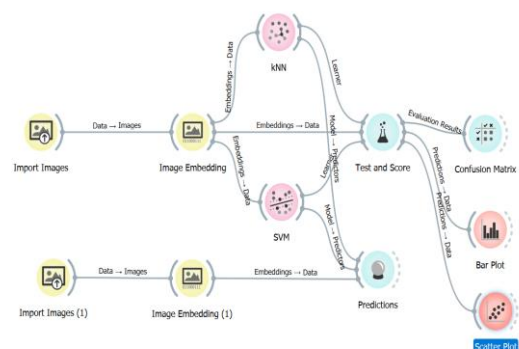
3. HASIL DAN PEMBAHASAN

3.1. Hasil

Hasil simulasi 2 model prediksi, data yang telah melalui tahap preprocessing selanjutnya dilakukan pengujian untuk mendapatkan model prediksi terbaik. Model prediksi yang telah dilakukan pengujian dan evaluasi dengan menggunakan kumpulan data uji pada aplikasi orange dimana 1 atribut sebagai targetnya didapatkan hasil simulasi seperti pada Gambar 3.

Model prediksi buah matang dapat dianalisa menggunakan *orange tool* untuk memilih metode terbaik. Model prediksi *data image* dapat dilihat pada Gambar 3. dibawah ini.

Gambar 3. Design model prediksi Kematangan Buah



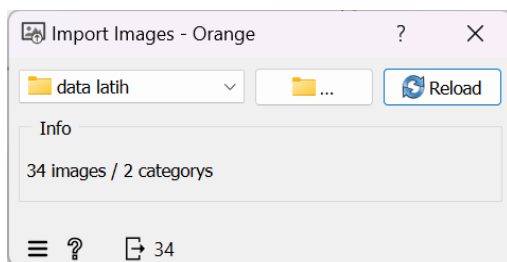
Setelah dilakukan preprocessing data, maka selanjutnya *data image* diproses kedalam model prediksi pada software orange dengan menggunakan metode K-NN dan SVM.

Perhitungan akurasi dilakukan dengan membandingkan antara total data yang diklasifikasikan benar dengan keseluruhan data. Precision menunjukkan tingkat kepastian dengan membandingkan antara sampel kelas positif yang diklasifikasikan dengan benar terhadap keseluruhan sampel kelas positif. Recall mengukur rasio antara sampel kelas positif yang diklasifikasikan dengan benar terhadap sampel kelas positif yang salah diklasifikasikan ke dalam kelas negatif. Adapun F1-Score merupakan harmonic mean dari precision dan recall (Hidayatullah et al., 2019).

3.2 PEMBAHASAN

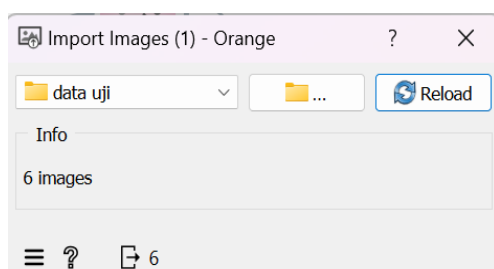
Data latih yang digunakan berjumlah 34 data image dengan 2 categorys, dapat dilihat pada gambar 4 dibawah ini :

Gambar 4. Data Latih



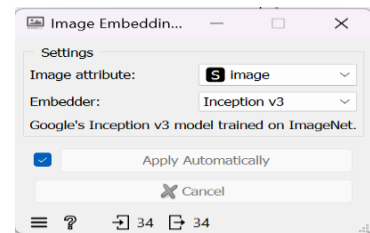
Data uji yang digunakan berjumlah 6 data image, yang dapat dilihat pada gambar 5 :

Gambar 5. Data Uji



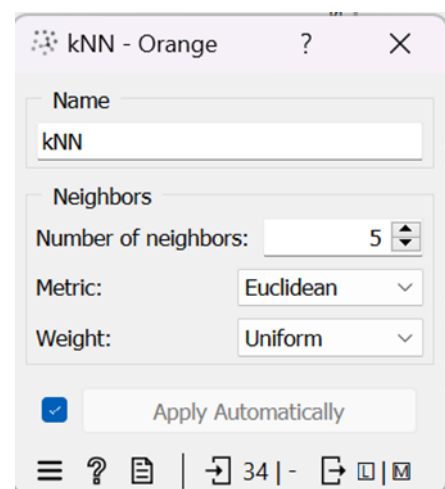
Data latih dan data uji akan dilakukan proses mengekstrasi fitur gambar menggunakan Inception-V3 yang terlihat pada gambar 6 berikut :

Gambar 6. Inception-V3

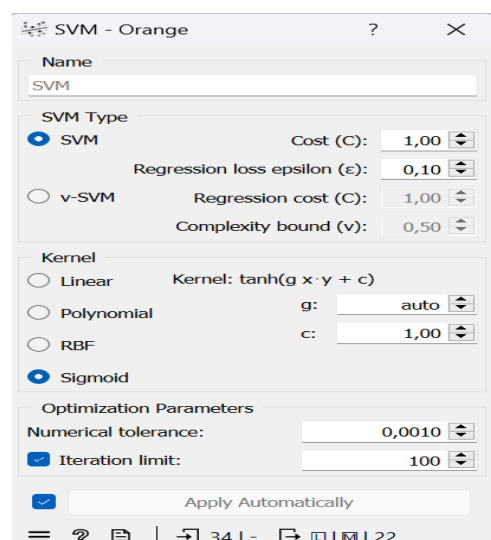


Setelah ekstrasi fitur gambar dilakukan, maka dilaksanakan hitungan metode *Support Vector Machine (SVM)* dan *K-Nearest neighbors (K-NN)*, yang dapat dilihat pada gambar 7 dan 8 pada gambar berikut ini :

Gambar 7. K-Nearest neighbors (K-NN)

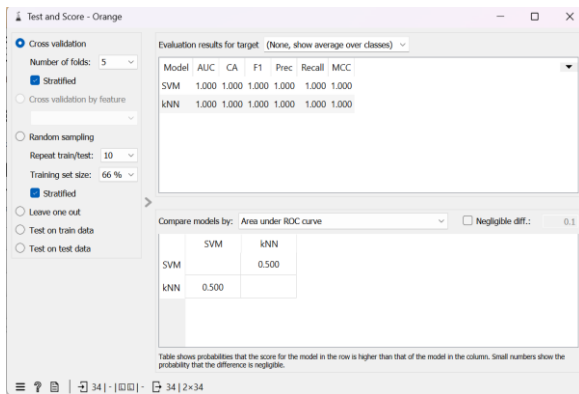


Gambar 8. Support Vector Machine (SVM)



Hasil dari perhitungan ke dua model tersebut dapat dilihat pada gambar 9. Sebagai berikut :

Gambar 9. Test And Score Support Vector Machine (SVM) dan K-Nearest neighbors (K-NN)



Pada Gambar 9. Merupakan proses evaluasi hasil perbandingan dari model prediksi yang diuji. Proses perhitungan keberhasilan model prediksi pada orange tool dapat menggunakan widget Test and Score dan hasil prediksinya dapat dilihat menggunakan widget Predictions.

Hasil simulasi 2 model prediksi data yang telah melalui tahap preprocessing selanjutnya dilakukan pengujian untuk mendapatkan model prediksi terbaik. Model prediksi yang telah dilakukan pengujian dan evaluasi dengan menggunakan kumpulan data uji pada aplikasi orange dimana 1 atribut sebagai targetnya didapatkan hasil simulasi seperti pada Gambar 10.

Gambar 10. prediksi widget Predictions

	SVM	KNN	image name	image	size	width	height
1	0.98 : 0.02 → Matang	1.00 : 0.00 → Matang	matang 18	matang 18.JPG	97644	922	615
2	0.51 : 0.49 → Mentah	0.60 : 0.40 → Matang	matang 19	matang 19.JPG	103717	922	615
3	0.90 : 0.10 → Matang	1.00 : 0.00 → Matang	matang 20	matang 20.JPG	101485	922	615
4	0.20 : 0.80 → Mentah	0.20 : 0.80 → Mentah	mentah 18	mentah 18.JPG	91451	922	615
5	0.03 : 0.97 → Mentah	0.00 : 1.00 → Mentah	mentah 19	mentah 19.JPG	94542	922	615
6	0.04 : 0.96 → Mentah	0.00 : 1.00 → Mentah	mentah 20	mentah 20.JPG	94537	922	615

Gambar 10 merupakan penampakan dari widget Test and Score aplikasi Orange. Pada penelitian ini proses pengujian menerapkan K-Fold Cross Validation (K=5) yang dapat diatur pada widget Test and Score seperti yang

terlihat pada gambar. Pada widget tersebut juga diperlihatkan hasil evaluasi kedua metode, dimana hasil perbandingan dari kedua model prediksi buah matang ditunjukkan dengan nilai Accuracy, Precision, Recall, dan F1-Score. Penelitian kali ini hanya memperhatikan nilai. Hasil evaluasi menunjukkan bahwa kedua model prediksi kematangan buah sawit memiliki nilai akurasi yang sama. Support Vector Machine (SVM) dan K-Nearest neighbors (K-NN) yaitu 0.500 dan 0.500. Hal tersebut menunjukkan bahwa kedua metode tersebut sama baiknya.

5. Kesimpulan dan Saran

Berdasarkan hasil penelitian dari 34 data image 2 kategori yang diekstraksi fitur menggunakan Inception-V3 dan dilakukan perhitungan dengan Support Vector Machine (SVM) dan K-Nearest neighbors (K-NN), Hasil evaluasi menunjukkan bahwa kedua model prediksi kematangan buah sawit memiliki nilai akurasi yang sama. Support Vector Machine (SVM) dan K-Nearest neighbors (K-NN) yaitu 0.500 dan 0.500. didapatkan model dengan Accuracy 100%, F1 100%, Recall 100%, Precision 100%.

Hasil klasifikasi dari 6 data image menghasilkan klasifikasi yang bernilai benar semua.

Pada penelitian selanjutnya diharapkan agar menambahkan jumlah dataset agar meningkatkan keakuratan dari model yang dihasilkan.

6. DAFTAR PUSTAKA

Achmad, N. D., Soleh, A. M., & Rizki, A. (2022). Perbandingan Pengklasifikasian Metode Support Vector Machine dan Random Forest (Kasus Perusahaan Kebun Kelapa Sawit). *Xplore: Journal of Statistics*, 11(2), 147–156. <https://doi.org/10.29244/xplore.v11i2.919>

Hidayatullah, A. F., Aulia, A., Yusuf, F., Juwairi, K. P., Abida, R., & Nayoan, N. (2019). Identifikasi Konten Kasar pada Tweet Bahasa Indonesia. In *JLK* (Vol. 2, Issue 1). <https://t.co/YQCC0CM4gG>

- Ichsan, N., Fatah, H., Wahyuni, T., & Ermawati, E. (2022). IMPLEMENTASI ORANGE DATA MINING UNTUK PREDIKSI HARGA BITCOIN. *JURNAL RESPONSIF*, 4(2), 118–125. <https://investing.com/crypto/bitcoin/historical->
- Krisandi, N., Helmi, B., & Prihandono, I. (2013). *ALGORITMA k -NEAREST NEIGHBOR DALAM KLASIFIKASI DATA HASIL PRODUKSI KELAPA SAWIT PADA PT. MINAMAS KECAMATAN PARINDU* (Vol. 02, Issue 1).
- Monika Parapat, I., & Tanzil Furqon, M. (2018). *Penerapan Metode Support Vector Machine (SVM) Pada Klasifikasi Penyimpangan Tumbuh Kembang Anak* (Vol. 2, Issue 10). <http://j-ptiik.ub.ac.id>
- Raysyah, S., Arinal, V., & Mulyana, D. I. (2021). KLASIFIKASI TINGKAT KEMATANGAN BUAH KOPI BERDASARKAN DETEKSI WARNA MENGGUNAKAN METODE KNN DAN PCA. *Sistem Informasi* |, 8(2), 88–95.