

Klasifikasi Permintaan Keripik Nenas Di PT Agritek Desa Indonesia Sipahutar Menggunakan Algoritma *Random Forest* Berdasarkan Bahan Baku

Fatner S Simanjuntak¹, Doli Hasibuan², Margaretha Yohanna³
^{1,2,3}Fakultas Ilmu Komputer, Universitas Methodist Indonesia

Info Artikel

Histori Artikel:

Received, Agst 25, 2025

Revised, Sept 10, 2025

Accepted, Sept 22, 2025

Keywords:

Keripik nenas,
Klasifikasi permintaan,
Random Forest,
Bahan baku.

ABSTRAK

Keripik nenas merupakan produk olahan buah nenas yang memiliki daya simpan lebih lama dibandingkan nenas segar, namun tetap memiliki batas kedaluwarsa sehingga perencanaan produksi yang tepat menjadi hal yang penting. PT Agritek Desa Indonesia (AGDIN) Sipahutar menghadapi permasalahan berupa fluktuasi permintaan pasar dan keterbatasan ketersediaan bahan baku, sehingga diperlukan metode yang mampu membantu dalam mengklasifikasikan permintaan keripik nenas secara akurat. Penelitian ini bertujuan untuk menerapkan algoritma *Random Forest* dalam mengklasifikasikan permintaan keripik nenas berdasarkan variabel produksi dan bahan baku. Data penelitian menggunakan periode Oktober 2023 hingga Mei 2025 dengan total 87 data permintaan dan bahan baku (dalam kilogram). Pengujian model dilakukan menggunakan tiga skenario pembagian data, yaitu 80:20, 70:30, dan 60:40. Hasil penelitian menunjukkan bahwa algoritma *Random Forest* mampu menangkap hubungan *non-linear* antara produksi dan ketersediaan bahan baku dalam mengklasifikasikan permintaan. Nilai akurasi yang diperoleh masing-masing sebesar 83% pada pembagian 80:20, 85% pada pembagian 70:30, dan 89% pada pembagian 60:40, dengan nilai *F1-Score* berada pada rentang 0,82 hingga 0,90. Selain itu, nilai *cross-validation 5-fold* sebesar 0,6771 menunjukkan bahwa model memiliki tingkat kestabilan yang cukup baik. Berdasarkan pola yang terbentuk, permintaan cenderung meningkat ketika produksi dan bahan baku berada pada kategori tinggi, serta menurun pada kondisi produksi dan bahan baku kategori sedang atau rendah. Dengan demikian, algoritma *Random Forest* terbukti efektif sebagai alat bantu pengambilan keputusan dalam perencanaan produksi dan pengadaan bahan baku keripik nenas di PT AGDIN Sipahutar.

This is an open access article under the [CC BY-SA](#) license.



Penulis Koresponden:

Fatner S Simanjuntak
Fakultas Ilmu Komputer,
Universitas Methodist Indonesia, Medan,
Jl. Hang Tuah No.8, Medan - Sumatera Utara.
Email: fatnersimanjuntak123@gmail.com

1. PENDAHULUAN

Keripik nenas merupakan salah satu produk olahan hasil pertanian yang berasal dari buah nenas segar melalui proses penggorengan dengan teknik khusus sehingga menghasilkan tekstur

renyah dan rasa khas buah nenas yang tetap terjaga. Produk ini memiliki daya simpan lebih lama dibandingkan nenas segar, yaitu berkisar antara 6–12 bulan tergantung kemasan dan kondisi penyimpanan. Namun demikian, keripik nenas tetap memiliki masa kedaluwarsa (*expired*), karena kualitasnya dapat menurun seiring waktu akibat faktor lingkungan seperti kelembapan, paparan udara, dan penyimpanan yang tidak sesuai standar. Oleh karena itu, perencanaan produksi yang tepat sangat penting agar tidak terjadi penumpukan stok berlebihan yang berisiko mendekati masa kedaluwarsa dan menimbulkan kerugian.

PT Agritek Desa Indonesia (AGDIN) yang berlokasi di Sipahutar merupakan perusahaan agribisnis yang berfokus pada pengolahan buah nenas menjadi berbagai produk turunan, termasuk keripik nenas dengan dua jenis kemasan, yaitu ukuran 100 gram dan 65 gram. Dalam kegiatan produksinya, PT AGDIN menghadapi berbagai tantangan, khususnya dalam ketersediaan bahan baku yang sangat dipengaruhi oleh kondisi cuaca, hama, serta penyakit tanaman yang dapat menurunkan kualitas maupun kuantitas hasil panen. Selain itu, keterbatasan infrastruktur seperti pasokan listrik juga menjadi kendala, karena gangguan listrik ketika proses penggorengan berlangsung dapat menyebabkan gagalnya produksi. Dari sisi pasar, permintaan lokal bersifat fluktuatif. Hal ini dipengaruhi oleh kebiasaan masyarakat setempat yang lebih terbiasa mengonsumsi buah nenas segar dibandingkan produk olahannya. Di sisi lain, kapasitas mesin produksi yang tersedia di PT AGDIN terbatas, yaitu sekitar 80–150 kg/minggu, sehingga perencanaan produksi harus disesuaikan dengan kapasitas riil perusahaan. Apabila permintaan pasar tidak dapat diklasifikasi dengan baik, maka perusahaan berisiko mengalami kelebihan produksi yang dapat mendekati masa kedaluwarsa (*expired*), atau sebaliknya kekurangan produksi yang menurunkan kepuasan konsumen.

Untuk menghadapi permasalahan tersebut, dibutuhkan suatu sistem klasifikasi permintaan keripik nenas yang lebih akurat agar proses produksi dapat direncanakan dengan baik berdasarkan ketersediaan bahan baku. Salah satu metode yang dapat digunakan adalah data mining dengan algoritma *Random Forest*. Algoritma ini termasuk dalam metode *supervised learning* yang bekerja dengan menggabungkan banyak pohon keputusan (*decision trees*) sehingga mampu menghasilkan model klasifikasi yang lebih stabil, akurat, dan tahan terhadap data yang kompleks maupun *noise* [1]. *Random Forest* dapat memanfaatkan data historis, seperti ketersediaan bahan baku nenas, kondisi lingkungan, serta permintaan pasar, untuk membangun model klasifikasi yang mendukung pengambilan keputusan produksi.

Sejumlah penelitian sebelumnya telah menunjukkan efektivitas *Random Forest* dalam bidang pertanian. Penelitian yang dilakukan oleh [2] mengenai klasifikasi hasil panen kakao di Desa Minanga berhasil mencapai akurasi sebesar 98,95% dengan nilai R^2 sebesar 0,99, dengan variabel hama dan penyakit sebagai faktor paling dominan. Penelitian lain oleh [3] juga menemukan bahwa *Random Forest* memiliki kinerja yang sangat baik dalam klasifikasi hasil panen berbasis data deret waktu dengan tingkat akurasi tinggi dan kesalahan rendah. Metode *Random Forest* memiliki keunggulan dalam akurasi karena menggabungkan banyak pohon keputusan (*ensembling*) sehingga mampu mengurangi overfitting dan menangani data kompleks. Dibandingkan metode lain seperti *Regresi Linear* atau *Decesion Tree* [4], dengan menggunakan *Random Forest* lebih stabil dalam klasifikasi terhadap data *noise*.

2. METODE PENELITIAN

2.1 Decision Tree

Decision tree adalah salah satu metode dalam *machine learning* dan data mining yang digunakan untuk membuat klasifikasi berdasarkan aturan keputusan yang diperoleh dari data [5]. Metode ini berbentuk diagram pohon di mana setiap cabang merepresentasikan keputusan atau kondisi, dan setiap daun merepresentasikan hasil atau kategori. *Decision tree* sangat berguna karena mudah dipahami dan diinterpretasikan, bahkan oleh orang yang tidak memiliki latar belakang teknis yang kuat. Pohon keputusan ini dapat untuk berbagai masalah klasifikasi dan regresi [6][7].

2.2 Random Forest

Machine Learning adalah salah satu dari metodologi ilmiah modern yang dapat melakukan prosedur otomatis untuk menghasilkan klasifikasi pada suatu fenomena dengan melakukan

pengamatan dari kejadian yang terjadi sebelumnya yaitu mencari pola pada suatu kumpulan data yang diberikan[8]. Metode *Random Forest* menerapkan metode *bootstrap aggregating (bagging)* dan *random feature selection* [9] [10] [11]. Saat ini *machine learning* telah menjadi metode yang umum digunakan untuk menyelesaikan suatu tugas atau masalah dalam kehidupan sehari-hari yang membutuhkan proses ekstraksi sekumpulan data yang besar [12]. Secara garis besar ada dua tipe *machine learning*, yaitu *Supervised Learning* dan *Unsupervised Learning*. *Supervised learning* mengacu pada *machine learning* dimana data yang digunakan untuk belajar sudah diberi label *output* yang harus dikeluarkan mesin, sedangkan *Unsupervised learning* sebaliknya mengacu pada *machine learning* yang belajar dari data yang tidak diberi label *output*.

2.2.1 Langkah-Langkah Metode Random Forest

Adapun tahapan dari penerapan Random Forest adalah [13]:

1. Persiapan Data

Persiapan data dimulai dengan pengumpulan dan pemilihan dataset yang relevan. Data kemudian diproses untuk memastikan kualitasnya, seperti menghapus duplikasi dan data yang hilang. Data yang digunakan harus representatif terhadap masalah yang ingin diselesaikan.

2. Preprocessing Data

Preprocessing data meliputi langkah-langkah untuk membersihkan dan menormalkan data, seperti menangani nilai yang hilang, encoding variabel kategorikal, dan scaling fitur. Ini penting untuk memastikan algoritma dapat bekerja dengan baik pada data yang diberikan.

3. Pembagian Data *Training* dan *Testing*

Data dibagi menjadi dua bagian: data pelatihan (*training set*) dan data pengujian (*testing set*). Pembagian ini umumnya dilakukan secara acak, dengan proporsi 70-80% untuk pelatihan dan 20-30% untuk pengujian, untuk menghindari *overfitting*. Untuk memilih atribut akar, didasarkan pada nilai *gain* tertinggi dari atribut-atribut yang ada. Untuk menghitung *gain* digunakan rumus menghitung *Gain* seperti yang tertera dalam persamaan 1 berikut [14]:

$$Gain(S, A) = Entropy(s) - \sum_{i=1}^n \frac{|S_i|}{|S|} * Entropy(S_i) \quad (1)$$

Di mana :

S : himpunan kasus

A : atribut

N : jumlah partisi atribut A

|S_i| : jumlah kasus pada partisi ke-i

|S| : jumlah kasus dalam S

Sementara itu, perhitungan nilai entropi dapat dilihat pada persamaan 2 berikut.

$$Entropy(S) = \sum_{i=1}^n -p_i * \log_2 p_i \quad (2)$$

Di mana :

S : himpunan kasus

A : fitur

N : jumlah partisi S

p_i : proporsi dari S_i terhadap S

4. Pembentukan *Decision Tree*

Random Forest membangun banyak pohon keputusan dengan memilih fitur secara acak pada setiap cabang. Setiap pohon dibangun menggunakan subset data yang diambil dengan teknik *bootstrap*, sehingga setiap pohon dapat memberikan keputusan yang berbeda.

5. Hasil Klasifikasi

Setelah semua pohon keputusan dibangun, *Random Forest* menghasilkan klasifikasi dengan cara menggabungkan hasil dari setiap pohon, menggunakan mekanisme *voting* untuk klasifikasi atau rata-rata untuk regresi. Klasifikasi ini mencerminkan keputusan kolektif dari seluruh hutan pohon.

6. Evaluasi Akurasi

Evaluasi dilakukan dengan mengukur akurasi model menggunakan data uji. Metode seperti *confusion matrix*, *precision*, *recall*, *F1-score*, atau AUC-ROC digunakan untuk mengukur kinerja model dalam memklasifikasi label yang benar atau hasil kontinu yang akurat.

2.2.2 Confusion Matrix

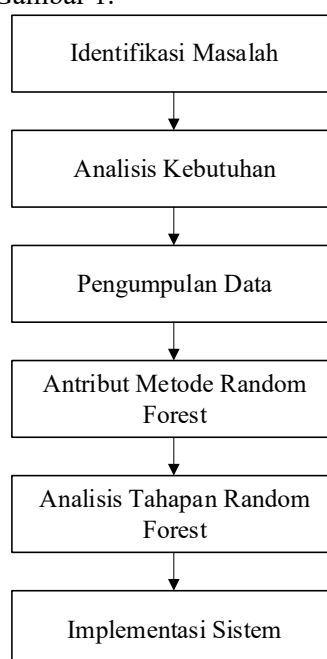
Confusion matrix adalah sebuah tabel yang digunakan untuk mengevaluasi performa model klasifikasi dengan membandingkan hasil klasifikasi model terhadap data sebenarnya. Matriks ini menampilkan empat komponen utama, yaitu *True Positive* (TP), *True Negative* (TN), *False Positive* (FP), dan *False Negative* (FN). Dengan melihat distribusi data pada keempat komponen tersebut, kita dapat menilai seberapa baik model dalam mengklasifikasikan data ke dalam kelas yang benar [15]. Salah satu metrik evaluasi yang paling umum digunakan adalah akurasi, yang mengukur proporsi klasifikasi yang benar dari seluruh klasifikasi yang dilakukan. Rumus akurasi dapat dinyatakan pada rumus 3:

$$\text{Akurasi} = \frac{TP + TN}{TP + TN + FP + FN} \quad (3)$$

Dimana TP dan TN adalah jumlah klasifikasi yang tepat untuk kelas positif dan negatif, sedangkan FP dan FN adalah jumlah kesalahan klasifikasi. Akurasi memberikan gambaran umum tentang performa model, namun perlu diperhatikan juga metrik lain terutama jika data tidak seimbang antar kelas.

2.3 Framework Penelitian

Dalam penelitian ini kerangka kerja digunakan dalam menjelaskan seluruh langkah-langkah aktifitas yang dikerjakan saat penelitian sehingga sesuai dengan tujuan yang telah ditentukan. Kerangka kerja dapat dilihat pada Gambar 1.



Gambar 1. Framework Penelitian

2.4 Pengumpulan Data

Pengumpulan data dilakukan secara sistematis melalui wawancara dan observasi langsung di PT Agritek Desa Indonesia (AGDIN) Sipahutar guna memastikan kesesuaian data dengan tujuan penelitian. Wawancara dilakukan dengan manajer produksi, staf pengadaan bahan baku, dan tenaga kerja lapangan untuk memperoleh informasi mendalam mengenai ketersediaan bahan baku nanas, pola permintaan keripik nanas, serta kendala dalam pemenuhan kebutuhan pasar, yang menghasilkan data kualitatif dan kontekstual. Selain itu, observasi dilakukan dengan mengamati aktivitas operasional pabrik, khususnya pengelolaan bahan baku dan dinamika permintaan, untuk mengidentifikasi hubungan antara ketersediaan bahan baku, fluktuasi permintaan, serta dampaknya terhadap keputusan produksi dan distribusi. Seluruh data yang diperoleh kemudian dipadukan dengan data kuantitatif historis sebagai dasar dalam pembangunan model klasifikasi permintaan

keripik nanas berbasis algoritma *Random Forest*. Adapun data yang digunakan dapat diamati pada tabel 1

Tabel 1. Data Bahan Baku

No	Minggu Dimulai	Produksi Keripik Nenas (kg)	Bahan Baku (kg) Nenas	Permintaan (kg)
1	02-Oct-23	100	822	105
2	09-Oct-23	100	767	111
3	16-Oct-23	95	760	104
4	23-Oct-23	95	726	104
5	30-Oct-23	95	744	108
6	06-Nov-23	90	701	83
7	13-Nov-23	90	687	84
8	20-Nov-23	95	772	84
9	27-Nov-23	100	769	87
10	04-Dec-23	105	823	94
...	...	...	...	...
87	26-May-25	90	727	76

Data pada Tabel 1, merupakan data kuantitatif yang menunjukkan jumlah produksi, ketersediaan bahan baku, dan permintaan keripik nanas di PT Agritek Desa Indonesia (AGDIN) Sipahutar. Nilai pada setiap variabel bersifat fluktuatif dari minggu ke minggu, menyesuaikan dengan kondisi nyata di lapangan seperti keterbatasan bahan baku, pengaruh cuaca, serta perubahan tren permintaan pasar. Dengan demikian, data ini tidak bersifat konstan, melainkan dinamis sehingga dapat dijadikan dasar yang relevan dalam membangun model klasifikasi permintaan menggunakan algoritma *Random Forest*.

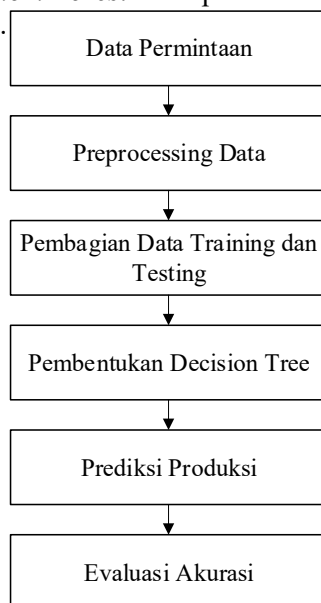
2.5 Tahapan Random Forest

Penerapan metode Random Forest pada penelitian ini diawali dengan pengumpulan data historis permintaan keripik nanas di PT Agritek Desa Indonesia (AGDIN) Sipahutar yang mencakup faktor bahan baku, jumlah produksi, dan kondisi permintaan pasar (naik/turun) sebagai label klasifikasi. Data kemudian melalui tahap preprocessing untuk mengatasi nilai kosong, duplikasi, dan inkonsistensi, sebelum dibagi menjadi data training dan testing. Selanjutnya, model Random Forest dilatih dengan membangun sejumlah pohon keputusan berdasarkan kombinasi acak fitur dan data, kemudian menentukan hasil klasifikasi melalui mekanisme voting mayoritas. Evaluasi kinerja model dilakukan menggunakan data testing dengan metrik akurasi serta didukung Precision, Recall, dan F1-Score. Pembagian data dilakukan secara acak dengan beberapa skenario rasio, yaitu 80:20, 70:30, dan 60:40, untuk memastikan model mampu mengenali pola permintaan secara representatif dan menghasilkan klasifikasi yang akurat terhadap data yang belum pernah digunakan sebelumnya.

Tabel 2. Bentuk *Dataset* Untuk Kebutuhan *Training* Dan *Testing*

Skenario <i>Dataset Training</i> dan <i>Dataset Testing</i>	Jumlah data pada <i>Datasets training</i>	Jumlah data pada <i>Datasets testing</i>	Jumlah Data
<i>Dataset Training</i> 80% dan <i>Dataset Testing</i> 20%	69 (80%)	18 (20%)	87
<i>Dataset Training</i> 70% dan <i>Dataset Testing</i> 30%	60 (70%)	27 (30%)	87
<i>Dataset Training</i> 60% dan <i>Dataset Testing</i> 40%	52 (60%)	35 (40%)	87

Adapun tahapan metode *Random Forest* ini dapat dilihat pada Gambar 2.



Gambar 2. Tahapan Metode *Random Forest*

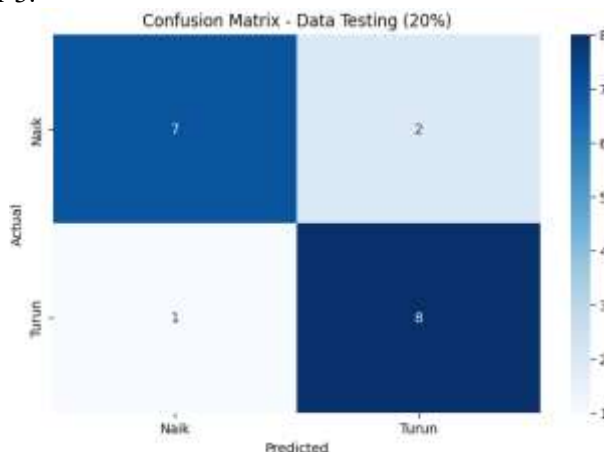
3. HASIL DAN PEMBAHASAN

3.1 Pengujian Sistem

Pengujian dilakukan dengan menggunakan metode seperti *confusion matrix*, akurasi, *precision*, *recall*, *F1-score*, dan *cross validation*[16]. Berdasarkan hasil pengujian tersebut, diperoleh performa model dalam memklasifikasi permintaan Keripik Nenas sesuai dengan kombinasi nilai Produksi dan Bahan Baku, sehingga dapat digunakan sebagai dasar pengambilan keputusan dalam perencanaan produksi:

1. Hasil Pengujian *Confusion Matrix*

Adapun hasil pengujian *confusion matrix* untuk analisis klasifikasi permintaan keripik nenas seperti pada Gambar 3.



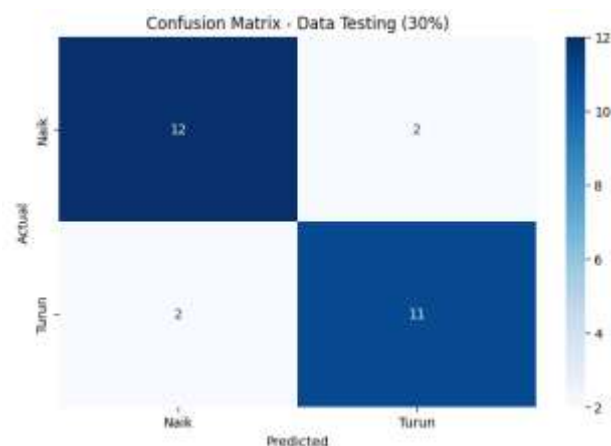
Gambar 3. Hasil Pengujian *Confusion Matrix* Data 80:20

Secara lengkap seperti pada Tabel 2.

Tabel 2. Hasil *Confusion Matrix* Data 80:20

	Klasifikasi:	
	Permintaan Naik	Permintaan Turun
Aktual: Permintaan Naik	7 (<i>True Positive</i>)	2 (<i>False Negative</i>)
Aktual: Permintaan Turun	1 (<i>False Positive</i>)	8 (<i>True Negative</i>)

Pengujian *confusion matrix* menggunakan data 70:30 seperti pada Gambar 4.

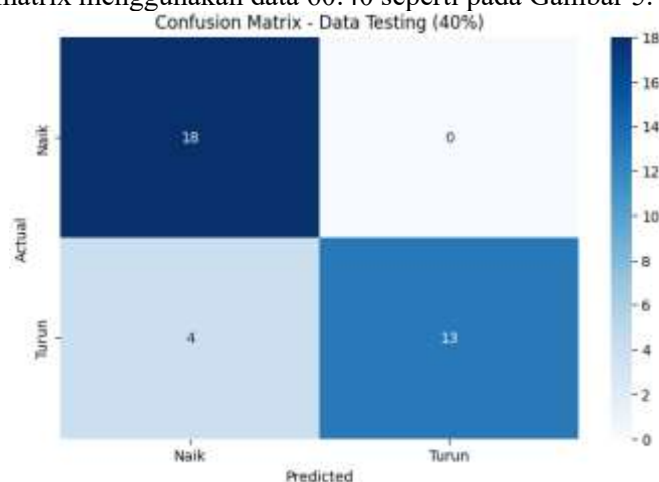


Gambar 4. Hasil Pengujian *Confusion Matrix* Data 70:30
Secara lengkap seperti pada Tabel 4.

Tabel 3. Hasil *Confusion Matrix* Data 70:30

	Klasifikasi: Permintaan Naik	Klasifikasi: Permintaan Turun
Aktual: Permintaan Naik	12 (<i>True Positive</i>)	2 (<i>False Negative</i>)
Aktual: Permintaan Turun	2 (<i>False Positive</i>)	11 (<i>True Negative</i>)

Pengujian *confusion matrix* menggunakan data 60:40 seperti pada Gambar 5.



Gambar 5. Hasil Pengujian *Confusion Matrix* Data 60:40
Secara lengkap seperti pada Tabel 4.

Tabel 4. Hasil *Confusion Matrix* Data 60:40

	Klasifikasi: Permintaan Naik	Klasifikasi: Permintaan Turun
Aktual: Permintaan Naik	18 (<i>True Positive</i>)	0 (<i>False Negative</i>)
Aktual: Permintaan Turun	4 (<i>False Positive</i>)	13 (<i>True Negative</i>)

3.2 Hasil Pengujian

Berikut tabel perbandingan hasil evaluasi model *Random Forest* berdasarkan pembagian data 80:20, 70:30, dan 60:40 seperti pada Tabel 5.

Tabel 5. Perbandingan Klasifikasi Model *Random Forest*

No	Pembagian Data	Akurasi	F1-Score Naik	F1-Score Turun	Cross Validation
1	80 : 20	0.83 (83%)	0.82	0.84	0.6771
2	70 : 30	0.85 (85%)	0.86	0.85	0.6771
3	60 : 40	0.89 (89%)	0.90	0.87	0.6771

Berdasarkan hasil penelitian klasifikasi permintaan keripik nenas di PT Agritek Desa Indonesia Sipahutar dengan *Random Forest* berdasarkan bahan baku, didapati bahwa algoritma *Random Forest* mampu mengklasifikasikan kondisi permintaan keripik nenas dengan performa yang baik pada berbagai skenario pembagian data. Dari ketiga pengujian, pembagian data 60:40 menghasilkan nilai akurasi tertinggi sebesar 89%, sedangkan pembagian 70:30 menunjukkan performa yang paling seimbang antara kelas permintaan naik dan turun. Nilai *cross-validation* sebesar 0,6771 pada seluruh skenario mengindikasikan bahwa model memiliki kemampuan generalisasi yang cukup stabil. Dengan demikian, algoritma *Random Forest* dapat dijadikan sebagai metode yang efektif dalam membantu perusahaan memprediksi dan mengelola permintaan keripik nenas berdasarkan ketersediaan bahan baku.

4. KESIMPULAN

Berdasarkan hasil penelitian, algoritma *Random Forest* terbukti efektif dalam mengklasifikasi permintaan keripik nenas di PT Agritek Desa Indonesia Sipahutar dengan memanfaatkan variabel produksi dan bahan baku, serta mampu menangkap pola hubungan non-linear yang memengaruhi fluktuasi permintaan. Hasil pengujian dengan skenario pembagian data 80:20, 70:30, dan 60:40 menunjukkan kinerja model yang baik dan konsisten, dengan akurasi berturut-turut sebesar 83%, 85%, dan 89%, serta nilai *F1-Score* pada kisaran 0,82–0,90. Pembagian data 70:30 memberikan keseimbangan klasifikasi terbaik, sementara 60:40 menghasilkan akurasi tertinggi, dan nilai *cross-validation 5-fold* sebesar 0,6771 menandakan kestabilan serta kemampuan generalisasi model yang memadai. Dengan demikian, *Random Forest* dapat digunakan sebagai alat bantu yang andal dalam mendukung pengambilan keputusan terkait perencanaan produksi dan pengadaan bahan baku keripik nenas di perusahaan.

REFERENSI

- [1] S. Saadah and H. Salsabila, "Prediksi Harga Bitcoin Menggunakan Metode Random Forest," *J. Komput. Terap.*, vol. 7, no. 1, pp. 24–32, 2021, doi: 10.35143/jkt.v7i1.4618.
- [2] Andrian, "PREDIKSI HASIL PANEN KAKAO DI DESA MINANGA MENGGUNAKAN ALGORITMA RANDOM FOREST REGRESSION," UNIVERSITAS SULAWESI BARAT MAJENE, 2025.
- [3] A. Chahal *et al.*, "Predictive analytics technique based on hybrid sampling to manage unbalanced data in smart cities," *Heliyon*, vol. 10, no. 24, p. e39275, 2024, doi: 10.1016/j.heliyon.2024.e39275.
- [4] D. H. Depari, Y. Widiastiwi, and M. M. Santoni, "Perbandingan Model Decision Tree, Naive Bayes dan Random Forest untuk Prediksi Klasifikasi Penyakit Jantung," *Inform. J. Ilmu Komput.*, vol. 18, no. 3, p. 239, 2022, doi: 10.52958/iftk.v18i3.4694.
- [5] A. P. Silalahi and H. G. Simanullang, *Decission Tree dan Penerapannya dengan Algoritma Iterative Dichotomizes versi 3 (ID3)*. 2025.
- [6] I. Khoeri and D. Iskandar Mulyana, "Implementasi Machine Learning dengan Decision Tree Algoritma C4.5 dalam Penerimaan Karyawan Baru pada PT. Gitareksa Dinamika Jakarta," *J. Sos. Teknol.*, vol. 1, no. 7, pp. 615–623, 2021, doi: 10.59188/jurnalsostech.v1i7.126.
- [7] O.- Pahlevi, A.- Amrin, and Y.- Handrianto, "Implementasi Algoritma Klasifikasi Random Forest Untuk Penilaian Kelayakan Kredit," *J. Infortech*, vol. 5, no. 1, pp. 71–76, 2023, doi: 10.31294/infortech.v5i1.15829.
- [8] Sriyanto and A. R. Supriyatna, "Prediksi Penyakit Diabetes Menggunakan Algoritma Random Forest," *J. Tek.*, vol. 17, no. 1, pp. 163–172, 2023.
- [9] T. Tambunan, M. Yohanna, and A. P. Silalahi, "Penerapan Metode Random Forest Dalam Mendeteksi Berita Hoax," *METHOMIKA J. Manaj. Inform. dan Komputerisasi Akunt.*, vol. 7, no. 2, pp. 301–306, 2023, doi: 10.46880/jmika.vol7no2.pp301-306.
- [10] W. Apriliah, I. Kurniawan, M. Baydhowi, and T. Haryati, "Prediksi Kemungkinan Diabetes pada Tahap Awal Menggunakan Algoritma Klasifikasi Random Forest," *Sistemasi*, vol. 10, no. 1, p. 163, 2021, doi: 10.32520/stmsi.v10i1.1129.
- [11] B. Prasojoa and E. Haryatmi, "Analisa Prediksi Kelayakan Pemberian Kredit Pinjaman

dengan Metode Random,” *J. Nas. Teknol. dan Sist. Inf.*, vol. 02, pp. 79–89, 2021.

[12] D. N. A. Fadlilah, Y. Wahyuningsih, and Y. A. Khawarga, “Prediksi Spesies Burung Menggunakan Random Forest,” *Nas. Teknol. Inf. dan Komputer*, vol. 6, no. 1, pp. 15–19, 2022, doi: 10.30865/komik.v6i1.5783.

[13] T. Posangi, L. Yahya, and D. Wungguli, “Implementasi Algoritma Random Forest dengan Forward Selection untuk Klasifikasi Indeks Pembangunan Manusia,” *Jambura J. Probab. Stat.*, vol. 4, no. 2, pp. 85–91, 2023, doi: 10.37905/jjps.v4i2.18460.

[14] E. M. Hutabarat, “Analisis Dan Klasifikasi Lalu Lintas Data Pada Jaringan Internet di Universitas Methodist Indonesia Dengan Metode Random Forest,” Universitas Methodist Indonesia, 2024.

[15] Suci Amaliah, M. Nusrang, and A. Aswi, “Penerapan Metode Random Forest Untuk Klasifikasi Varian Minuman Kopi di Kedai Kopi Konijiwa Bantaeng,” *VARIANSI J. Stat. Its Appl. Teach. Res.*, vol. 4, no. 3, pp. 121–127, 2022, doi: 10.35580/variantsiunm31.

[16] L. M. Barus, H. G. Simanullang, and A. P. Silalahi, “Classification of English-Language Racial Hate Speech Using Random Forest ”,” *J. Intell. Comput. Inf. Syst.*, vol. 1, no. 1, pp. 1–17, 2025.