

Evaluasi Performa *Naïve Bayes* dan *CART* pada Klasifikasi Kualitas Tahu

Luthfy Akmal Nugraha¹, Jupriyanto², Haris Nizhomul Haq², Anderias Eko Wijaya^{1*},
Hermansyah Nur Ahmad¹

¹Teknik Informatika, Fakultas Teknik, Universitas Mandiri, Indonesia

²Sistem Informasi, Fakultas Teknik, Universitas Mandiri, Indonesia

E-mail: nugrahala@gmail.com¹, jupriyanto.kahar@gmail.com², harisnizhom@gmail.com²,
ekowjy09@gmail.com^{1*}, hermansyahna@gmail.com¹

Received: 2025-09-03 | Accepted: 2025-09-11 | Published: 2025-10-01

Abstrak

Untuk tetap bersaing di pasar global, produsen tahu harus memastikan kualitas produk yang konsisten. Pabrik Tahu Sumber Barokah, sebagai pemasok tahu bernutrisi tinggi yang telah lama beroperasi, menghadapi tantangan dalam menjaga kualitas sepanjang proses produksi. Penelitian ini membandingkan kinerja algoritma *Naïve Bayes* dan *Classification and Regression Trees (CART)* dalam mengklasifikasikan kualitas tahu menggunakan dataset yang dikumpulkan dari pabrik, yang berisi sampel tahu berkualitas tinggi dan rendah. Metodologi penelitian mencakup identifikasi masalah, pengumpulan data, preprocessing, klasifikasi, validasi, evaluasi, dan penarikan kesimpulan. *Cross-validation* digunakan untuk validasi model, dan *confusion matrix* digunakan untuk menilai *precision*, *recall*, dan *F1-score*. Hasil eksperimen menunjukkan bahwa *Naïve Bayes* mencapai akurasi 91%, *precision* 100%, *recall* 85%, dan *F1-score* 92%, sedangkan *CART* mencapai akurasi 86%, *precision* 70%, *recall* 100%, dan *F1-score* 82%. Hasil ini menunjukkan bahwa *Naïve Bayes* lebih cocok untuk mengklasifikasikan kualitas tahu dalam konteks ini.

Kata Kunci: *Naïve Bayes, CART, Klasifikasi Kualitas Tahu, Cross-Validation.*

Abstract

To remain competitive in the global market, tofu producers must ensure consistent product quality. Sumber Barokah Tofu Factory, a longstanding supplier of high-nutrient tofu, faces challenges in maintaining quality throughout the production process. This study compares the performance of *Naïve Bayes* and *Classification and Regression Trees (CART)* algorithms in classifying tofu quality using a dataset collected from a factory, which contains both high-quality and low-quality tofu samples. The research methodology encompasses problem identification, data collection, preprocessing, classification, validation, evaluation, and conclusion. *Cross-validation* was employed for model validation, and *confusion matrices* were utilised to assess *precision*, *recall*, and *F1-score*. Experimental results indicate that *Naïve Bayes* achieved an Accuracy of 91%, *precision* of 100%, *recall* of 85%, and *F1-score* of 92%, while *CART* achieved an Accuracy of 86%, *precision* of 70%, *recall* of 100%, and *F1-score* of 82%. These results suggest that *Naïve Bayes* is more suitable for classifying tofu quality in this context.

Keywords: *Naïve Bayes, CART, Tofu Quality Classification, Cross-Validation.*

1. Pendahuluan

Industri tahu di Indonesia memegang peran penting dalam memenuhi kebutuhan pangan bergizi bagi masyarakat [1], [2]. Tahu merupakan salah satu sumber protein nabati yang murah dan mudah diakses, sehingga kontribusinya terhadap ketahanan pangan nasional sangat signifikan. Namun, tantangan utama yang dihadapi oleh produsen tahu adalah menjaga konsistensi kualitas produk di tengah fluktuasi bahan baku dan variasi dalam proses produksi yang kompleks, mulai dari pemilihan kedelai, proses penggilingan, koagulasi, hingga pencetakan dan pengemasan [2], [3]. Ketidakkonsistenan kualitas dapat berdampak pada nilai gizi, tekstur, dan keamanan produk, sehingga diperlukan pengawasan yang ketat dan sistem klasifikasi yang efektif untuk memastikan kualitas tahu tetap terjaga [3], [4].

Dalam konteks ini, algoritma pembelajaran mesin seperti *Naïve Bayes* dan *Classification and Regression Trees (CART)* telah banyak diterapkan dalam berbagai studi untuk klasifikasi kualitas produk pangan [1], [5], [6]. *Naïve Bayes* dikenal karena kesederhanaannya, kecepatan proses, dan kemampuannya

dalam menangani data besar dengan performa yang baik, terutama ketika variabel input bersifat independen [5], [7]. Sementara itu, CART menawarkan interpretabilitas yang tinggi melalui struktur pohon keputusan yang mudah dipahami, memudahkan produsen dan peneliti untuk menafsirkan faktor-faktor yang memengaruhi kualitas produk, dan telah diterapkan dalam pengklasifikasian kualitas tahu maupun produk pangan lain [6], [8], [9].

Beberapa penelitian sebelumnya telah membandingkan kedua algoritma ini dalam konteks klasifikasi produk pangan. Misalnya, studi oleh Azis dan Jabir menggunakan model pohon keputusan untuk menganalisis kualitas aroma tahu yang diawetkan dengan formalin, dengan akurasi sebesar 36,79% [1]. Studi lain mengembangkan varian pohon keputusan seperti *HHCART* untuk meningkatkan akurasi klasifikasi dengan memanfaatkan refleksi *matriks Householder* pada setiap simpul pohon, yang dapat mengatasi keterbatasan pemisahan sumbu pada pohon keputusan tradisional [8]. Pendekatan pohon keputusan teracak optimal juga telah diusulkan untuk mengatasi keterbatasan metode *greedy* tradisional dalam membangun pohon dan meningkatkan akurasi klasifikasi [9].

Penelitian ini bertujuan untuk membandingkan kinerja algoritma Naïve Bayes dan CART dalam mengklasifikasikan kualitas tahu menggunakan dataset dari Pabrik Tahu Sumber Barokah [10]. Dengan membandingkan akurasi, *precision*, *recall*, dan F1-score dari kedua algoritma, diharapkan penelitian ini dapat memberikan kontribusi dalam pengembangan sistem klasifikasi otomatis yang efektif, sehingga produsen tahu dapat memastikan kualitas produk yang konsisten dan meningkatkan daya saing industri tahu di Indonesia [2], [3].

2. Metode Penelitian

2.1. Alur Penelitian

Penelitian ini menggunakan pendekatan kuantitatif eksperimen untuk membandingkan kinerja dua algoritma klasifikasi, yaitu *Naïve Bayes* dan *CART*, dalam mengklasifikasikan kualitas tahu [1], [5], [6]. Fokus penelitian adalah menentukan algoritma yang paling akurat dan efisien dalam mengidentifikasi kualitas tahu berdasarkan data sensor dan parameter produksi [2], [3]. Dataset penelitian diperoleh dari Pabrik Tahu Sumber Barokah, yang mencakup sampel tahu berkualitas tinggi dan rendah. Data dikumpulkan melalui pengukuran fisik dan kimia, seperti tekstur, warna, kadar air, pH, serta parameter produksi, termasuk lama koagulasi, suhu, dan jenis kedelai [10]. Total dataset yang digunakan terdiri dari sejumlah sampel yang proporsinya seimbang antara kelas kualitas tinggi dan rendah, dan disimpan dalam format CSV untuk diolah menggunakan perangkat lunak *Python* [5], [6].

Sebelum diterapkan pada algoritma klasifikasi, dataset melalui tahap *preprocessing* untuk memastikan kualitas data. Proses ini mencakup pembersihan data dari nilai yang hilang atau outlier ekstrem, normalisasi fitur numerik ke rentang 0–1 agar algoritma dapat bekerja lebih efektif, serta konversi variabel kategori seperti jenis kedelai menjadi format numerik melalui teknik *one-hot encoding* [4], [7]. Setelah *preprocessing*, dataset diimplementasikan pada dua algoritma klasifikasi. Algoritma Naïve Bayes dipilih karena kesederhanaannya, kecepatan proses, dan kemampuannya dalam menangani data besar dengan performa yang baik, terutama ketika fitur bersifat independen [1], [5]. Sedangkan CART dipilih karena interpretabilitasnya tinggi; algoritma ini membagi dataset berdasarkan nilai fitur paling informatif menggunakan Gini Index, menghasilkan struktur pohon keputusan yang mudah dipahami dan digunakan untuk menilai faktor-faktor yang memengaruhi kualitas tahu [6], [8], [9].

Validasi model dilakukan menggunakan *k-fold cross-validation* dengan $k = 5$, yang membagi dataset menjadi lima subset agar setiap data diuji pada model. Teknik ini digunakan untuk mengurangi bias dan memperoleh estimasi performa yang lebih andal [1], [5]. Kinerja masing-masing algoritma kemudian dievaluasi menggunakan metrik standar dalam klasifikasi, yaitu akurasi, *precision*, *recall*, dan F1-score [5], [1], [6]. Akurasi mengukur persentase prediksi yang benar terhadap total prediksi, *precision* menilai rasio prediksi positif yang benar terhadap seluruh prediksi positif, *recall* mengevaluasi rasio prediksi positif yang benar terhadap semua data positif sesungguhnya, dan F1-score memberikan evaluasi seimbang antara *precision* dan *recall*.

Alur penelitian dimulai dengan identifikasi masalah dan tujuan penelitian, dilanjutkan dengan pengumpulan dataset dari Pabrik Tahu Sumber Barokah, *preprocessing* data, implementasi algoritma *Naïve Bayes* dan *CART*, validasi model menggunakan *k-fold cross-validation*, evaluasi kinerja berdasarkan metrik yang telah disebutkan, dan diakhiri dengan analisis hasil serta penarikan kesimpulan untuk menentukan algoritma terbaik dalam klasifikasi kualitas tahu [1], [2], [3], [5], [6]. Dengan metodologi ini, penelitian bertujuan memberikan kontribusi signifikan dalam pengembangan sistem klasifikasi otomatis yang dapat membantu produsen tahu menjaga kualitas produk secara konsisten dan meningkatkan daya saing industri tahu di Indonesia [2], [3].

2.2. Dataset Penelitian

Sebagian dataset yang digunakan dalam penelitian ini ditampilkan pada Tabel 1. Tabel ini menunjukkan sampel produk tahu beserta atributnya, seperti aroma, tekstur, cita rasa, masa kadaluarsa, dan klasifikasi kualitas. Informasi ini memberikan gambaran mengenai variasi karakteristik produk yang akan dianalisis oleh algoritma klasifikasi.

Tabel 1. Dataset Produk

No	Produk Tahu	Aroma	Tekstur	Cita Rasa	Masa Kadaluarsa	Kualitas
1	Tahu Putih	Bagus	Bagus	Rendah	1 Hari	Ya
2	Tahu Kuning	Bagus	Bagus	Bagus	2 Hari	Ya
3	Tahu Matang	Sedang	Bagus	Rendah	3 Hari	Tidak
4	Tahu Sumedang	Sedang	Bagus	Sedang	2 Hari	Tidak
5	Tahu Pletok	Bagus	Rendah	Rendah	3 hari	Tidak
6	Tahu Mencos	Bagus	Bagus	Bagus	1 Hari	Ya
7	Tahu Pong	Rendah	Bagus	Bagus	3 Hari	Ya
8	Tahu Putih	Sedang	Rendah	Rendah	3 Hari	Tidak
9	Tahu Kuning	Sedang	Bagus	Bagus	1 Hari	Ya
10	Tahu Matang	Sedang	Bagus	Bagus	3 Hari	Tidak
11	Tahu Sumedang	Sedang	Bagus	Bagus	1 Hari	Ya
...	...	...	...	...	...	...
210	Tahu Pletok	Sedang	Bagus	Bagus	1 Hari	Ya

2.3. Preprocessing Data

Tahap *preprocessing* data diawali dengan proses *data cleaning* yang bertujuan untuk mengurangi jumlah dan kompleksitas data akibat adanya data non-relevan maupun *noise* yang berupa record hilang, invalid, atau kesalahan penulisan. Proses ini penting dilakukan karena kualitas data akan memengaruhi kinerja algoritma klasifikasi yang digunakan [5], [6].

Dataset awal yang diperoleh dari Pabrik Tahu Sumber Barokah berjumlah 210 record. Setelah melalui tahapan *cleaning*, jumlah data berkurang menjadi 110 record yang dianggap valid dan relevan untuk dianalisis lebih lanjut. Reduksi data ini terjadi karena sejumlah entri ditemukan tidak lengkap atau tidak sesuai dengan atribut penelitian, sehingga harus dihapus dari dataset.

Hasil dari proses *data cleaning* ditunjukkan pada Tabel 3.2, yang memuat sampel data pasca pembersihan. Dataset yang tersisa inilah yang kemudian digunakan pada tahap implementasi algoritma *Naïve Bayes* dan *CART*.

Tabel 2. Data Cleaning

No	Produk Tahu	Aroma	Tekstur	Cita Rasa	Masa Kadaluarsa	Kualitas
1	Tahu Putih	Bagus	Bagus	Rendah	1 Hari	Ya
2	Tahu Kuning	Bagus	Bagus	Bagus	2 Hari	Ya
3	Tahu Matang	Sedang	Bagus	Rendah	3 Hari	Tidak
4	Tahu Sumedang	Sedang	Bagus	Sedang	2 Hari	Tidak
5	Tahu Pletok	Bagus	Rendah	Rendah	3 hari	Tidak
...	...	...	...	...	...	...
110	Tahu Pletok	Sedang	Bagus	Bagus	1 Hari	Ya

3. Hasil dan Pembahasan

3.1. Implementasi Naïve Bayes

Algoritma *Naïve Bayes* digunakan dalam penelitian ini karena kemampuannya yang sederhana namun efektif dalam mengklasifikasikan data berbasis probabilitas. Prinsip dasar *Naïve Bayes* mengacu pada *Teorema Bayes*, yang dirumuskan sebagai berikut:

$$P(C|X) = \frac{P(X|C).P(C)}{P(X)} \quad (1)$$

dengan:

- ($P(C|X)$) adalah probabilitas posterior dari kelas (C) terhadap data (X),
- ($P(X|C)$) adalah probabilitas likelihood dari data (X) pada kelas (C),
- ($P(C)$) adalah probabilitas prior dari kelas (C),

- ($P(X)$) adalah probabilitas evidence dari data (X).

Dalam konteks penelitian ini, kelas (C) terdiri dari dua kategori yaitu Kualitas = Ya (tahu berkualitas baik) dan Kualitas = Tidak (tahu berkualitas rendah). Sedangkan fitur (X) meliputi atribut aroma, tekstur, cita rasa, dan masa kadaluarsa.

Langkah pertama adalah menghitung probabilitas prior dari masing-masing kelas menggunakan formula (1). Dari total 110 record data hasil *preprocessing*, diperoleh:

$$P(A) = \frac{\text{Jumlah Sampel di Kelas A}}{\text{Jumlah Total Sampel}}$$

$$P(B) = \frac{\text{Jumlah Sampel di Kelas B}}{\text{Jumlah Total Sampel}}$$

$$P(A) = \frac{64}{110} = 0,58182$$

$$P(B) = \frac{46}{110} = 0,41818$$

Hasil probabilitas *prior* adalah:

$$P(A) = 0,58182$$

$$P(B) = 0,41818$$

Menghitung Probabilitas *Likelihood*

$$P(X|C) = P(x_1|C) * P(x_2|C) \dots P(x_n|C)$$

Ya|Aroma

$$P(\text{Aroma} = \text{Bagus}|Ya) = \frac{53}{64} = 0,09375$$

$$P(\text{Aroma} = \text{Sedang}|Ya) = \frac{6}{64} = 0,09375$$

$$P(\text{Aroma} = \text{Rendah}|Ya) = \frac{5}{64} = 0,078125$$

Ya|Tekstur

$$P(\text{Tekstur} = \text{Bagus}|Ya) = \frac{55}{64} = 0,859375$$

$$P(\text{Tekstur} = \text{Sedang}|Ya) = \frac{6}{64} = 0,09375$$

$$P(\text{Tekstur} = \text{Rendah}|Ya) = \frac{3}{64} = 0,046875$$

Ya|Cita Rasa

$$P(\text{Cita Rasa} = \text{Bagus}|Ya) = \frac{55}{64} = 0,859375$$

$$P(\text{Cita Rasa} = \text{Sedang}|Ya) = \frac{4}{64} = 0,0625$$

$$P(\text{Cita Rasa} = \text{Rendah}|Ya) = \frac{5}{64} = 0,07813$$

Ya|Masa Kadaluarsa

$$P(\text{Masa Kadaluarsa} = \text{Bagus}|Ya) = \frac{56}{64} = 0,875$$

$$P(\text{Masa Kadaluarsa} = \text{Sedang}|Ya) = \frac{6}{64} = 0,09375$$

$$P(\text{Masa Kadaluarsa} = \text{Rendah}|Ya) = \frac{2}{64} = 0,03125$$

Tidak|Aroma

$$P(\text{Aroma} = \text{Bagus}|Tidak) = \frac{9}{46} = 0,195652174$$

$$P(\text{Aroma} = \text{Sedang}|Tidak) = \frac{33}{46} = 0,717391304$$

$$P(\text{Aroma} = \text{Rendah}|Tidak) = \frac{4}{46} = 0,086956522$$

Tidak|Tekstur

$$\begin{aligned}
 P(\text{Aroma} = \text{Bagus}|\text{Tidak}) &= \frac{14}{46} = 0,304347826 \\
 P(\text{Aroma} = \text{Sedang}|\text{Tidak}) &= \frac{8}{46} = 0,173913043 \\
 P(\text{Aroma} = \text{Rendah}|\text{Tidak}) &= \frac{24}{46} = 0,52173913 \\
 \text{Tidak|Cita Rasa} \\
 P(\text{Aroma} = \text{Bagus}|\text{Tidak}) &= \frac{12}{46} = 0,260869565 \\
 P(\text{Aroma} = \text{Sedang}|\text{Tidak}) &= \frac{15}{46} = 0,326086957 \\
 P(\text{Aroma} = \text{Rendah}|\text{Tidak}) &= \frac{19}{46} = 0,413043478 \\
 \text{Tidak|Masa Kadaluarsa} \\
 P(\text{Aroma} = \text{Bagus}|\text{Tidak}) &= \frac{3}{46} = 0,065217391 \\
 P(\text{Aroma} = \text{Sedang}|\text{Tidak}) &= \frac{26}{46} = 0,565217391 \\
 P(\text{Aroma} = \text{Rendah}|\text{Tidak}) &= \frac{17}{46} = 0,369565217
 \end{aligned}$$

Langkah serupa diterapkan pada semua kombinasi fitur dan kelas. Hasil perhitungan ini kemudian digunakan untuk memperoleh nilai probabilitas posterior setiap kelas terhadap sampel baru. Prediksi kelas dilakukan dengan memilih kelas yang memiliki nilai probabilitas posterior terbesar.

Hal serupa juga terlihat pada atribut Masa Kadaluarsa, di mana kategori 1 Hari memiliki probabilitas ($P(1 \text{ Hari} | \text{Ya}) = 0,875$), sedangkan pada kelas "Tidak" hanya sebesar 0,0652. Dengan demikian, tahu dengan masa kadaluarsa satu hari cenderung lebih banyak ditemukan pada kategori tahu berkualitas baik.

Tabel 3 menyajikan nilai probabilitas likelihood untuk seluruh atribut penelitian, yaitu aroma, tekstur, cita rasa, dan masa kadaluarsa. Informasi ini kemudian digunakan dalam perhitungan probabilitas posterior berdasarkan rumus *Naïve Bayes* untuk menentukan prediksi kelas dari sampel tahu baru yang diuji.

Hasil probabilitas Likelihood:

Tabel 3. Probabilitas Likelihood

Kategori/Atribut	Subset	Ya	Tidak
Aroma	Bagus	0,828125	0,19565217
	Sedang	0,09375	0,7173913
	Rendah	0,078125	0,08695652
Tekstur	Bagus	0,859375	0,30434783
	Sedang	0,09375	0,17391304
	Rendah	0,046875	0,52173913
Cita Rasa	Bagus	0,859375	0,26086957
	Sedang	0,0625	0,32608696
	Rendah	0,07813	0,41304348
Masa Kadaluarsa	1 Hari	0,875	0,06521739
	2 Hari	0,09375	0,56521739
	3 Hari	0,03125	0,36956522

Dengan pendekatan ini, *Naïve Bayes* dapat mengklasifikasikan kualitas tahu berdasarkan atribut yang tersedia secara sistematis dan terukur [1], [5], [6].

Setelah diperoleh nilai probabilitas prior dan likelihood, tahap selanjutnya adalah menghitung hasil klasifikasi menggunakan rumus *Naïve Bayes*:

$$P(c | F) \propto P(c) \prod_{i=1}^n P(f_i | c) \quad (2)$$

di mana c adalah kelas (dalam penelitian ini: Ya atau Tidak), dan $F = \{f_1, f_2, \dots, f_n\}$ adalah himpunan fitur sampel (aroma, tekstur, cita rasa, masa kadaluarsa). Karena pembagi $P(F)$ sama untuk semua kelas, kita cukup membandingkan nilai proporsional di kanan untuk menentukan kelas dengan probabilitas *posterior* terbesar.

Rumus (2) tersebut diaplikasikan pada setiap data uji untuk menentukan kelas “Ya” (tahu berkualitas baik) atau “Tidak” (tahu tidak layak). Sebagai contoh, pada salah satu data uji dengan atribut Aroma = Bagus, Tekstur = Bagus, Cita Rasa = Bagus, Masa Kadaluarsa = 1 Hari, probabilitas dihitung dengan mengalikan nilai *likelihood* setiap atribut dengan probabilitas prior dari kelas. Hasil perhitungan menunjukkan bahwa nilai probabilitas untuk kelas *Ya* lebih tinggi dibandingkan dengan kelas *Tidak*, sehingga data tersebut diklasifikasikan ke dalam kategori *Ya*.

Selanjutnya, seluruh hasil perhitungan untuk 110 data uji disajikan dalam Tabel 4. Tabel ini memperlihatkan nilai probabilitas masing-masing kelas, label fakta, hasil klasifikasi, serta kesesuaian prediksi. Dari tabel tersebut terlihat bahwa sebagian besar data dapat diklasifikasikan dengan tepat oleh algoritma *Naïve Bayes*. Namun, masih terdapat beberapa kasus salah klasifikasi (misalnya pada baris ke-9 dan 22), di mana prediksi berbeda dengan fakta sebenarnya. Temuan ini menunjukkan bahwa meskipun *Naïve Bayes* memiliki akurasi yang cukup tinggi, tetap ada keterbatasan dalam menangani data dengan distribusi tertentu.

Tabel 4. Hasil Klasifikasi Naive Bayes

No	Ya	Tidak	Fakta	Klasifikasi	Prediksi
1	0,186813354	0,00798878	Ya	Ya	Sesuai
2	0,249084473	0,00582515	Ya	Ya	Sesuai
3	0,00453186	0,08462236	Tidak	Tidak	Sesuai
4	3,424911499	0,00121569	Ya	Ya	Sesuai
5	0,000606537	0,07167728	Tidak	Tidak	Sesuai
6	0,186813354	0,00798878	Ya	Ya	Sesuai
7	0,00125885	0,15330994	Tidak	Tidak	Sesuai
8	0,323104858	0,00207118	Ya	Ya	Sesuai
9	0,041542053	0,14808926	Tidak	Ya	Tidak
10	3,424911499	0,00121569	Ya	Ya	Sesuai
15	0,055330013	0,38772583	Tidak	Tidak	Sesuai
16	0,38772583	0,01708713	Ya	Ya	Sesuai
17	0,001602173	0,00562175	Tidak	Tidak	Sesuai
18	0,033966064	0,00421629	Ya	Ya	Sesuai
19	0,122318268	0,02640752	Ya	Ya	Sesuai
20	0,001455688	0,08654543	Tidak	Tidak	Sesuai
21	0,373626709	0,00266293	Ya	Ya	Sesuai
22	0,028198242	0,02135888	Ya	Tidak	Tidak
23	3,424911499	0,00121569	Ya	Ya	Sesuai
...	...	...	...	...	...
110	3,424911499	0,00121569	Ya	Ya	Sesuai

Untuk mengukur performa model secara kuantitatif, dilakukan perhitungan *confusion matrix*. Berdasarkan hasil pengujian, diperoleh nilai *True Positive* (TP) = 11, *True Negative* (TN) = 9, *False Positive* (FP) = 0, dan *False Negative* (FN) = 2. Dari nilai ini, dihitung akurasi sebesar 91%, *precision* 100%, *recall* 85%, *specificity* 100%, dan *F1-score* 92%.

Hasil klasifikasi ini akan dibandingkan dengan algoritma *CART* pada tahap analisis berikutnya, untuk menentukan algoritma mana yang memberikan kinerja lebih baik dalam mengklasifikasikan kualitas tahu di Pabrik Tahu Sumber Barokah.

3.2. Implementasi *CART*

Berdasarkan Tabel 5, ditentukan sebanyak 12 calon cabang yang diperoleh dari atribut Aroma, Tekstur, Cita Rasa, dan Masa Kadaluarsa.

Tabel 5 Calon Cabang (*CART*)

No	Calon Cabang Kiri	Calon Cabang Kanan
1	Aroma = Bagus	Aroma = (Sedang,Rendah)
2	Aroma = Sedang	Aroma = (Bagus,Rendah)
3	Aroma = Rendah	Aroma = (Bagus,Sedang)
4	Tekstur = Bagus	Tekstur = (Sedang,Rendah)
5	Tekstur = Sedang	Tekstur = (Bagus,Rendah)
6	Tekstur = Rendah	Tekstur = (Bagus,Sedang)
7	Cita Rasa = Bagus	Cita Rasa = (Sedang,Rendah)
8	Cita Rasa = Sedang	Cita Rasa = (Bagus,Rendah)
9	Cita Rasa = Rendah	Cita Rasa = (Bagus,Sedang)
10	MK = 1 hari	MK = (2,3) hari
11	MK = 2 hari	MK = (1,3) hari
12	MK = 3 hari	MK = (1,2) hari

Selanjutnya, untuk masing-masing calon cabang dihitung jumlah data pada sisi kiri (PL) dan sisi kanan (PR), sebagaimana ditampilkan pada Tabel 6.

Tabel 6 Calon Cabang $\Phi(s/t)$

NO	PL	PR
1	62	48
2	39	71
3	9	101
4	62	41
5	14	96
6	27	83
7	67	43
8	19	91
9	24	86
10	59	51
11	32	78
12	19	91

Langkah berikutnya adalah menghitung nilai probabilitas $P(L)$ dan $P(R)$ dengan rumus.

$$P(L) = \frac{P_L}{P_L + P_R}, P(R) = \frac{P_R}{P_L + P_R} \quad (3)$$

Hasil perhitungan probabilitas ini menjadi bobot pada proses perhitungan impurity node kiri dan kanan. Dengan demikian, nilai probabilitas $P(L)$ dan $P(R)$ akan memengaruhi besarnya nilai kesesuaian $\Phi(s/t)$ yang digunakan untuk menentukan pemisahan terbaik dalam algoritma *CART*.

$$P(j|t) = \frac{N_{jt}}{N_t} \quad (4)$$

dengan:

- N_{jt} = jumlah sampel kelas j pada node t
- N_t = total sampel pada node t
- j = kelas (dalam kasus penelitianmu: “Ya” = tahu berkualitas, “Tidak” = tahu tidak layak)

Perhitungan ini bertujuan untuk mengetahui distribusi kelas pada masing-masing sisi cabang. Setelah itu dilakukan perhitungan nilai kesesuaian $\Phi(s/t)$ menggunakan rumus. Nilai kesesuaian terbesar menunjukkan atribut terbaik yang dipilih sebagai pemisah pada node.

$$i(t) = 1 - \sum_j P(j|t)^2 \quad (5)$$

Sehingga alurnya jadi:

1. Tentukan calon cabang (Tabel 5)
2. Hitung jumlah data PL dan PR (Tabel 6)
3. Dari PL dan PR → dapatkan P(L) dan P(R)
4. Hitung distribusi kelas di masing-masing node dengan rumus (4)
5. Hitung impurity $i(t_L), i(t_R)$, lalu dapatkan nilai kesesuaian $\Phi(s/t)$
- 6.

Tabel 7 Data Perhitungan $P(j|t_L)$ dan $P(j|t_R)$

No	Keputusan	$P(j t_L)$	$P(j t_R)$
1	YA	0,85483871	0,229166667
	TIDAK	0,14516129	0,770833333
2	YA	0,153846154	0,816901408
	TIDAK	0,846153846	0,183098592
3	YA	0,555555556	0,584158416
	TIDAK	0,444444444	0,415841584
4	YA	0,797101449	0,219512195
	TIDAK	0,202898551	0,780487805
5	YA	0,428571429	0,604166667
	TIDAK	0,571428571	0,395833333
6	YA	0,111111111	0,734939759
	TIDAK	0,888888889	0,265060241
7	YA	0,820895522	0,209302326
	TIDAK	0,179104478	0,790697674
8	YA	0,210526316	0,659340659
	TIDAK	0,421052632	0,340659341
9	YA	0,208333333	0,686046512
	TIDAK	0,791666667	0,313953488
10	YA	0,949152542	0,156862745
	TIDAK	0,050847458	0,843137255
11	YA	0,1875	0,743589744
	TIDAK	0,8125	0,256410256
12	YA	0,105263158	0,681318681
	TIDAK	0,894736842	0,318681319

Hasil perhitungan menunjukkan bahwa atribut (contoh: Cita Rasa = Bagus) memberikan nilai $\Phi(s/t)$ tertinggi, sehingga dipilih sebagai pemisah pertama pada pohon keputusan *CART*. Dari pemisah ini, pohon dibangun secara rekursif hingga seluruh data dapat dipisahkan dengan optimal.

$$\Phi(s | t) = Q\left(\frac{s}{t}\right) \times 2 \cdot PL \cdot PR \quad (6)$$

Berdasarkan hasil perhitungan probabilitas kelas pada sisi kiri dan kanan (Tabel 7), selanjutnya dihitung nilai *goodness of split* $Q\left(\frac{s}{t}\right)$, bobot pemisahan ($2 \cdot PL \cdot PR$), dan nilai akhir kesesuaian. Hasil perhitungan tersebut ditampilkan pada Tabel 8.

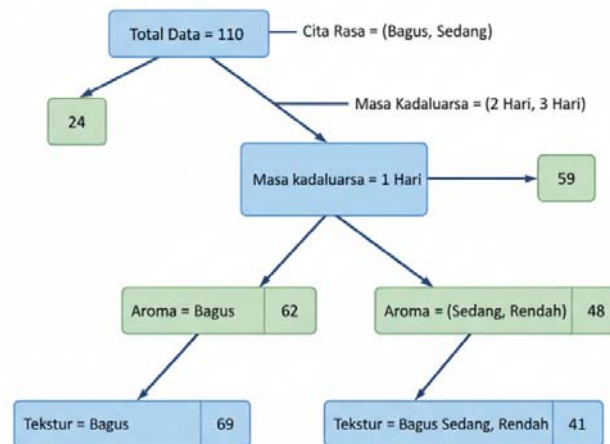
Tabel 8 Data Calon Cabang Kesesuaian $\Phi(s/t)$

NO	$Q(s/t)$	$2PL \cdot PR$	$\Phi(s/t)$
1	1,251344086	0,491900826	0,61553719
2	1,326110509	0,45768595	0,606942149
3	0,057205721	0,150247934	0,008595041
4	1,155178508	0,467603306	0,540165289
5	0,351190476	0,22214876	0,078016529
6	1,247657296	0,370413223	0,46214876
7	1,223186394	0,476198347	0,582479339
8	0,529207634	0,285785124	0,151239669
9	0,955426357	0,952396694	0,909944904
10	1,584579595	0,497355372	0,788099174
11	0	0,412561983	0
12	1,152111047	0,285785124	0,329256198

Setelah pohon terbentuk, dilakukan pengujian menggunakan data uji sebanyak 110 record. Dengan cara ini, performa *CART* dapat dibandingkan langsung dengan algoritma *Naïve Bayes* pada kasus klasifikasi kualitas tahu di Pabrik Tahu Sumber Barokah.

Tabel 9 Pohon Keputusan (*CART*)

Cita Rasa = Rendah	24
Cita Rasa = (Bagus,Sedang)	86
Masa kadaluarsa = 1 Hari	59
Masa Kadaluarsa = (2 Hari,3 Hari)	51
Aroma = Bagus	62
Aroma = (Sedang,Rendah)	48
Tekstur = Bagus	69
Tekstur = (Sedang,Rendah)	41



Gambar 1. Pohon Keputusan (CART)

Tabel 9 menyajikan hasil pemodelan pohon keputusan (CART) berdasarkan atribut yang digunakan dalam klasifikasi kualitas produk. Pada atribut cita rasa, data menunjukkan bahwa 24 sampel dikategorikan dalam tingkat rendah, sedangkan mayoritas, yaitu 86 sampel, masuk dalam kategori bagus dan sedang. Dari sisi masa kadaluarsa, terdapat 59 sampel dengan daya tahan hanya 1 hari, sementara 51 sampel lainnya memiliki masa kadaluarsa 2 hingga 3 hari.

Selanjutnya, atribut aroma memperlihatkan distribusi 62 sampel yang berkualitas bagus dan 48 sampel dengan aroma sedang maupun rendah. Sedangkan pada atribut tekstur, sebanyak 69 sampel dinilai memiliki tekstur bagus, sementara 41 sampel lainnya masuk dalam kategori tekstur sedang hingga rendah.

Hasil ini menunjukkan bahwa secara umum distribusi data lebih banyak terkonsentrasi pada kategori kualitas yang baik, baik dari aspek cita rasa, aroma, maupun tekstur. Sementara itu, variasi masa kadaluarsa relatif seimbang antara produk dengan daya tahan singkat (1 hari) dan lebih lama (2–3 hari). Temuan ini mengindikasikan bahwa atribut-atribut tersebut berperan penting dalam membentuk pola keputusan klasifikasi.

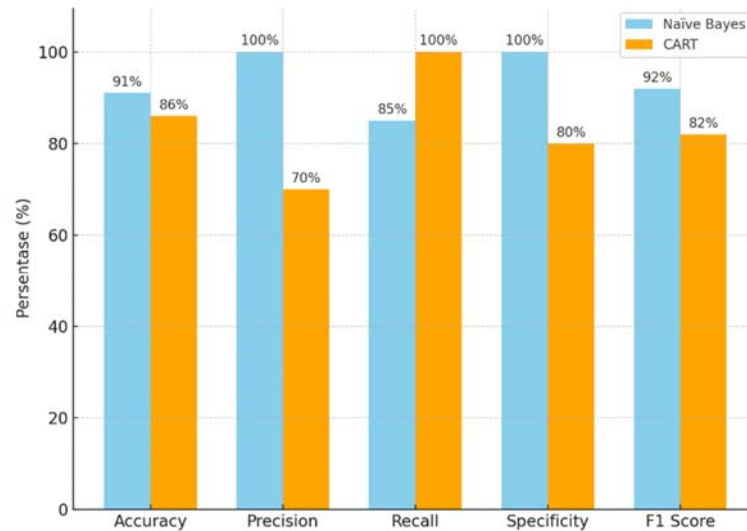
Hasil pengujian kemudian disajikan dalam bentuk *confusion matrix*, yang berisi jumlah data (TP), (TN), (FP), dan (FN). Dari *confusion matrix* tersebut dihitung nilai *akurasi*, *precision*, *recall*, dan *F1-score* menggunakan rumus (16)–(20). Berdasarkan hasil evaluasi menggunakan *confusion matrix*, metode CART memperoleh nilai akurasi sebesar 86%. Hal ini menunjukkan bahwa secara keseluruhan model mampu mengklasifikasikan data dengan tingkat ketepatan yang cukup tinggi. Nilai *precision* sebesar 70% mengindikasikan bahwa dari seluruh prediksi positif yang dihasilkan, 70% di antaranya benar sesuai dengan kelas aktual. Nilai *recall* mencapai 100%, artinya seluruh data positif berhasil diprediksi dengan benar oleh model tanpa ada yang terlewat.

3.3. Evaluasi Hasil Pengujian

Berdasarkan hasil evaluasi menggunakan *confusion matrix*, kedua algoritma—*Naive Bayes* dan CART—menunjukkan kinerja yang baik dalam mengklasifikasikan kualitas tahu di Pabrik Tahu Sumber Barokah, meskipun dengan karakteristik performa yang berbeda.

Naive Bayes menghasilkan akurasi sebesar 91% dengan *precision* 100%, *recall* 85%, *specificity* 100%, dan *F1-score* 92%. Hal ini menunjukkan bahwa *Naive Bayes* memiliki kemampuan yang sangat baik dalam mengenali data negatif secara tepat tanpa menghasilkan kesalahan *false positive*, sekaligus menjaga keseimbangan antara *precision* dan *recall*. Dengan kata lain, setiap prediksi positif yang diberikan oleh model dapat dipastikan benar, meskipun masih terdapat sebagian kecil data positif yang tidak terdeteksi.

Sementara itu, CART memberikan akurasi 86% dengan *precision* 70%, *recall* 100%, *specificity* 80%, dan *F1-score* 82%. Keunggulan utama CART terletak pada nilai *recall* yang mencapai 100%, yang berarti seluruh data positif dapat teridentifikasi dengan sempurna tanpa ada yang terlewat. Namun, nilai *precision* yang lebih rendah dibandingkan *Naive Bayes* menunjukkan bahwa CART lebih sering menghasilkan *false positive*, sehingga tingkat ketepatannya dalam memberikan prediksi positif relatif lebih rendah. Untuk lebih jelasnya terlihat pada gambar 1.

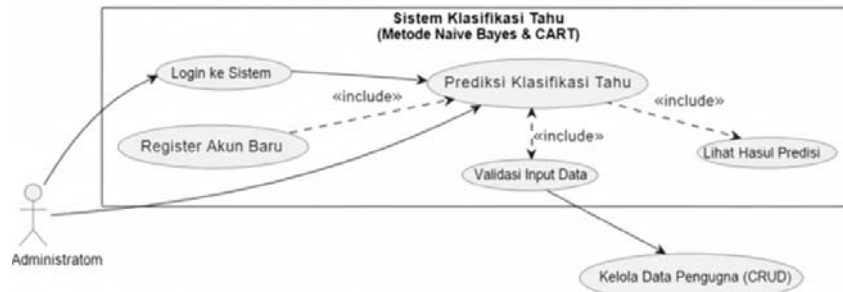


Gambar 1 Perbandingan Kinerja Naïve Bayes dan CART

Dari perbandingan ini dapat disimpulkan bahwa *Naive Bayes* lebih unggul dalam hal akurasi, precision, dan specificity, sedangkan *CART* lebih menonjol pada aspek recall. Dengan demikian, pilihan algoritma terbaik sangat bergantung pada kebutuhan praktis: apabila tujuan utamanya adalah memastikan tidak ada data positif yang terlewat, maka *CART* lebih sesuai, sebaliknya, apabila ketepatan prediksi positif lebih diutamakan, *Naive Bayes* menjadi pilihan yang lebih tepat.

3.4. Implementasi Sistem

Diagram use case di atas menggambarkan interaksi antara aktor utama, yaitu Administrator, dengan sistem Klasifikasi Tahu yang menggunakan metode *Naive Bayes* dan *CART*.



Gambar 2. Use Case

Administrator memiliki beberapa fungsi utama di dalam sistem, yaitu:

1. Login ke Sistem
Sebelum menggunakan fitur utama, administrator perlu melakukan autentikasi melalui proses login. Setelah berhasil login, administrator dapat mengakses seluruh modul dalam sistem.
2. Register Akun Baru
Administrator dapat melakukan pendaftaran akun baru bagi pengguna sistem. Fitur ini memastikan hanya pengguna terotorisasi yang dapat memanfaatkan sistem klasifikasi.
3. Prediksi Klasifikasi Tahu
Fitur utama dari sistem adalah melakukan prediksi kualitas tahu berdasarkan data input yang diberikan. Pada tahap ini, terdapat integrasi dengan dua komponen pendukung:
4. Validasi Input Data: sebelum proses prediksi dijalankan, data yang dimasukkan terlebih dahulu divalidasi untuk memastikan kesesuaian format dan kelengkapan.
5. Lihat Hasil Prediksi: setelah proses klasifikasi, administrator dapat menampilkan hasil prediksi kualitas tahu.
6. Kelola Data Pengguna (CRUD)

Administrator memiliki kewenangan untuk mengelola data pengguna melalui operasi *Create, Read, Update, dan Delete* (CRUD). Fitur ini menjaga integritas dan keamanan data pengguna sistem.

Relasi antar use case ditunjukkan dengan hubungan «include», yang menandakan bahwa proses prediksi selalu mencakup validasi data input dan penampilan hasil prediksi. Dengan demikian, alur sistem menjadi lebih terstruktur dan terjamin kualitas hasil klasifikasinya.

Secara keseluruhan, diagram ini memperlihatkan bagaimana administrator dapat mengoperasikan sistem klasifikasi tahu, mulai dari autentikasi, pengelolaan akun, pengelolaan data pengguna, hingga menjalankan prediksi kualitas tahu dengan metode *Naive Bayes* dan *CART*.

4. Kesimpulan

Berdasarkan hasil penelitian, dapat disimpulkan bahwa kedua algoritma, *Naive Bayes* dan *CART*, mampu digunakan untuk mengklasifikasikan kualitas tahu di Pabrik Tahu Sumber Barokah dengan tingkat performa yang cukup baik. Namun, terdapat perbedaan karakteristik pada hasil evaluasi masing-masing algoritma.

Naive Bayes menunjukkan kinerja yang lebih unggul dengan akurasi 91%, *precision* 100%, *recall* 85%, *specificity* 100%, dan *F1-score* 92%. Hal ini menegaskan bahwa algoritma *Naive Bayes* lebih andal dalam menghasilkan prediksi yang tepat, terutama dalam meminimalkan kesalahan *false positive* sekaligus mempertahankan keseimbangan antara *precision* dan *recall*.

Di sisi lain, *CART* menghasilkan akurasi 86%, *precision* 70%, *recall* 100%, *specificity* 80%, dan *F1-score* 82%. Keunggulan utama *CART* terletak pada kemampuannya mendeteksi seluruh data positif (*recall* 100%), meskipun tingkat *precision*nya lebih rendah dibandingkan *Naive Bayes*.

Dengan demikian, *Naive Bayes* lebih sesuai digunakan untuk klasifikasi kualitas tahu dalam konteks penelitian ini, karena mampu memberikan performa yang lebih stabil dan akurat. Sementara itu, *CART* dapat dijadikan alternatif jika tujuan analisis lebih menekankan pada deteksi semua data positif tanpa terlewat.

Daftar Pustaka

- [1] H. Azis and S. R. Jabir, "Chemical Composition and Aroma Profiling: Decision Tree Modeling of Formalin Tofu," *J. Embed. Syst. Secur. Intell. Syst.*, pp. 206–211, Nov. 2023, doi: 10.59562/jessi.v4i2.1162.
- [2] Z. Huang *et al.*, "Evaluating the effect of different processing methods on fermented soybean whey-based tofu quality, nutrition, and flavour," *LWT*, vol. 158, p. 113139, Mar. 2022, doi: 10.1016/j.lwt.2022.113139.
- [3] M. Herrmann, E. Mehner, L. Egger, R. Portmann, L. Hammer, and T. Nemecek, "A comparative nutritional life cycle assessment of processed and unprocessed soy-based meat and milk alternatives including protein quality adjustment," *Front. Sustain. Food Syst.*, vol. 8, Jun. 2024, doi: 10.3389/fsufs.2024.1413802.
- [4] F. Ali, K. Tian, and Z.-X. Wang, "Modern techniques efficacy on tofu processing: A review," *Trends Food Sci. Technol.*, vol. 116, pp. 766–785, Oct. 2021, doi: 10.1016/j.tifs.2021.07.023.
- [5] B. Phatcharathada and P. Srisuradetthai, "Randomized Feature and Bootstrapped Naive Bayes Classification," *Appl. Syst. Innov.*, vol. 8, no. 4, p. 94, Jul. 2025, doi: 10.3390/asi8040094.
- [6] A. Malik, B. Ram, D. Arumugam, Z. Jin, X. Sun, and M. Xu, "Predicting gypsum tofu quality from soybean seeds using hyperspectral imaging and machine learning," *Food Control*, vol. 160, p. 110357, Jun. 2024, doi: 10.1016/j.foodcont.2024.110357.
- [7] N. Zhang, E. Zhang, and F. Li, "A Soybean Classification Method Based on Data Balance and Deep Learning," *Appl. Sci.*, vol. 13, no. 11, p. 6425, May 2023, doi: 10.3390/app13116425.
- [8] D. C. Wickramarachchi, B. L. Robertson, M. Reale, C. J. Price, and J. Brown, "HH CART: An Oblique Decision Tree," Apr. 2015, [Online]. Available: <http://arxiv.org/abs/1504.03415>.
- [9] R. Blanquero, E. Carrizosa, C. Molero-Río, and D. R. Morales, "Optimal randomized classification trees," Oct. 2021, doi: 10.1016/j.cor.2021.105281.
- [10] C. A. Döttinger *et al.*, "Unravelling the genetic architecture of soybean tofu quality traits," *Mol. Breed.*, vol. 45, no. 1, p. 8, Jan. 2025, doi: 10.1007/s11032-024-01529-x.