

IMPLEMENTASI KLASIFIKASI KUALITAS SUSU MENGGUNAKAN ALGORITMA DECISION TREE, K-NEAREST NEIGHBORS DAN NAIVE BAYES

Jidan Haviar Saviola, Nofrian Deny Hendrawan

Sistem Informasi, Universitas Merdeka Malang

Jalan Terusan Dieng, Malang, Indonesia

wildanavi2@gmail.com

ABSTRAK

Susu mengandung banyak nutrisi yang penting bagi manusia. Nutrisi didalam susu sangat mempengaruhi tingkat kualitas sebuah susu Oleh sebab itu, penting untuk melakukan pengujian pada kualitas susu agar tidak ada susu berkualitas rendah yang dikonsumsi. Namun dalam melakukan pengujian kualitas susu membutuhkan hasil yang akurat dan tepat. Dengan menggunakan machine learning kita dapat membuat model untuk membantu melakukan otomatisasi dalam mengklasifikasikan kualitas susu. Dengan menggunakan algoritma decision tree, k-nearest neighbors dan naïve bayes. Ketiga model tersebut akan dibandingkan untuk mencari model yang terbaik untuk dataset Milk Quality pada penelitian kali ini. Metode *smote totem links* diterapkan saat proses pengolahan data dalam melakukan klasifikasi agar menjadi lebih akurat dan adil. Hasil akurasi pada model k-nearest neighbors mencapai 1.00 dan dapat disimpulkan bahwa model mencapai akurasi yang sempurna dan hampir tidak melakukan kesalahan dalam melakukan prediksi.

Kata kunci : susu, machine learning , smote totem links, decision tree, k-nearest neighbors, naïve bayes

1. PENDAHULUAN

Susu merupakan makanan yang kaya akan nutrisi yang penting bagi kehidupan manusia, seperti protein, lemak, karbohidrat, mineral, vitamin, dan faktor pertumbuhan. Nutrisi dalam susu sapi bisa dipengaruhi oleh tingkat kualitasnya. Jika susu yang diminum memiliki mutu yang tidak baik, hal ini bisa menimbulkan masalah kesehatan bagi manusia, seperti gangguan pencernaan [1].

Pengujian untuk penilaian kualitas susu merupakan salah satu proses pengolahan susu yang sangat penting. Penilaian kualitas susu perlu dilakukan untuk mengetahui kelayakan sebuah susu sebelum dikonsumsi [2]. Dalam hal penilaian kualitas susu, salah satu tantangan utama yang dihadapi adalah kemampuan untuk mengklasifikasikan kualitas susu dengan akurat dan tepat, yang mana dibutuhkan analisis mendalam terhadap beragam parameter atau fitur. Dalam konteks ini, machine learning muncul sebagai solusi yang sangat relevan, khususnya melalui pendekatan supervised learning.

Machine learning, sebagai bagian dari artificial intelligence (AI) digunakan dalam penelitian ini untuk membangun sistem klasifikasi kualitas susu yang mampu belajar dari data sebelumnya tanpa perlu pemrograman eksplisit. Ada beberapa jenis machine learning yang ada hingga saat ini, diantaranya *supervised learning*, *unsupervised learning*, dan *reinforcement learning*, namun jenis machine learning yang digunakan pada penelitian kali ini adalah *supervised learning* [3].

Supervised learning adalah cabang dari machine learning yang bergantung pada kehadiran ahli domain untuk memberikan bimbingan atau supervisi kepada sistem pembelajaran. Tujuan utama supervised learning adalah untuk memetakan input menjadi

output tertentu berdasarkan pola yang ditemukan dalam data latih atau training set. Dalam konteks klasifikasi kualitas susu, teknik ini dapat diterapkan untuk mengenali pola dari berbagai parameter susu dan memprediksi kelas kualitasnya [4].

Supervised learning dibagi menjadi dua kategori utama: klasifikasi dan regresi. Klasifikasi menghasilkan output dalam bentuk label, sedangkan regresi menghasilkan output dalam bentuk nilai kontinu. Pada penelitian ini, fokus utama adalah pada pendekatan klasifikasi, karena kualitas susu dikategorikan ke dalam kelas-kelas tertentu. Algoritma yang digunakan antara lain *decision tree*, *naive bayes*, dan *k-nearest neighbors*, yang masing-masing memiliki keunggulan dalam hal kecepatan, keakuratan, dan interpretasi. Ketiga algoritma ini dipilih untuk dibandingkan kinerjanya dalam mengklasifikasikan kualitas susu secara otomatis berdasarkan parameter yang tersedia [5].

Tujuan penelitian kali ini adalah untuk mengetahui model manakah yang mendapatkan hasil akurasi terbaik. Penelitian ini menggunakan dataset Milk Quality yang berasal dari *Kaggle* terdiri dari 1.059 sampel dan 7 fitur, seperti pH, suhu, rasa, bau dan 3 fitur lainnya dengan 1 label sebagai output. Dataset ini dipilih karena atribut-atributnya relevan untuk klasifikasi kualitas susu, mudah dinormalisasi, serta mempercepat proses pengumpulan data.

Penulis berencana menyelesaikan masalah ini dengan melakukan implementasi dataset tersebut pada ketiga model dengan menggunakan metode yang sama dan melakukan perbandingan hasil akurasi, presisi, recall dan F1-score yang merupakan matrik penting dalam tahapan evaluasi untuk mengukur performa model.

2. TINJAUAN PUSTAKA

Industri susu merupakan salah satu sektor pangan yang memiliki potensi besar di Indonesia seiring meningkatnya konsumsi masyarakat. Namun, kualitas susu yang rentan menurun akibat kontaminasi mikroba dan penanganan yang kurang tepat menjadi tantangan serius. Berbagai penelitian telah dilakukan untuk memprediksi kualitas susu menggunakan pendekatan Machine Learning, seperti algoritma K-Nearest Neighbor maupun Naive Bayes. Namun demikian, implementasi teknologi ini masih memerlukan pengembangan untuk meningkatkan akurasi dan efisiensi proses kontrol mutu di lapangan.

Pada penelitian Aditiya Rahman, menyatakan bahwa hasil pengumpulan dan pengolahan data penelitian mengenai dampak kualitas produk dan citra merek terhadap keputusan pembelian konsumen Susu Indomilk di Kecamatan Tanah Sareal Bogor menunjukkan bahwa kualitas produk dan citra merek, secara bersamaan dan signifikan, mempengaruhi keputusan pembelian (Y). Pengaruh tersebut dapat dijelaskan melalui variabel kualitas produk (X1) dan citra merek (X2) [6].

Berdasarkan penelitian berjudul Milk Quality Prediction Using Machine Learning oleh Drashti Bhavsar dkk, mendapati hasil model dari algoritma SVM mencapai akurasi rata-rata 57 %, sedangkan RF mencapai 92% dalam mengklasifikasikan sampel ke dalam grade yang telah ditetapkan [7] [8]

Penelitian oleh Fauzi Wardah Ali, berjudul Seleksi Fitur Information Gain Untuk Klasifikasi Kualitas Susu Sapi Menggunakan Metode K-Nearest Neighbor Dan Naive Bayes, disimpulkan bahwa masing-masing algoritma memiliki hasil yang berbeda maka dari itu penentuan algoritma dalam melakukan klasifikasi kualitas susu harus tepat [9].

Penelitian [10], mendapatkan hasil pengujian menggunakan metode K-NN dengan pembagian data 80–20, 70–30, dan 60–40 menunjukkan bahwa performa terbaik dicapai pada skema 80–20. Dengan menggunakan jumlah kluster optimal $K = 3$, yang ditentukan melalui metode elbow, model ini berhasil mencapai akurasi sebesar 0,90, recall 0,80, dan precision 1,00.

2.1. Supervised Learning

Supervised learning dikenal sebagai pendekatan pembelajaran mesin yang memanfaatkan dataset dengan label yang jelas, memungkinkan algoritma untuk mengenali pola spesifik yang sesuai dengan keluaran yang diharapkan. Dalam pendekatan ini, sistem dilatih menggunakan dataset yang terdiri dari pasangan input dan output yang diharapkan, sehingga model dapat mempelajari hubungan antara keduanya dengan akurat. Berkat kemampuannya, supervised learning menjadi alat yang sangat bermanfaat bagi organisasi dalam mengembangkan model-model kompleks yang mampu menghasilkan prediksi yang andal [11].

Supervised learning telah banyak diterapkan di berbagai sektor, seperti layanan kesehatan, pemasaran, jasa keuangan, dan bidang lainnya, berkat efektivitasnya dalam mengolah data. Berbeda dengan unseprvised learning yang beroperasi tanpa label, supervised learning secara eksplisit menggunakan data pelatihan berlabel untuk mengekstraksi pola dan membuat keputusan berdasarkan data baru yang serupa.

Tujuan supervised learning adalah untuk membuat mesin memahami sistem klasifikasi yang telah kita buat. Secara umum, pembelajaran terarah memungkinkan adanya probabilitas untuk input yang tidak terdefinisi, seperti input di mana hasil yang diharapkan sudah diketahui. Proses ini menyajikan seperangkat data yang terdiri dari fitur-fitur dan label. Tugas utamanya adalah mengembangkan estimator yang mampu meramalkan label objek berdasarkan fitur yang diberikan. Selanjutnya, algoritme pembelajaran menerima fitur sebagai masukan bersama dengan hasil yang benar dan belajar dengan membandingkan hasil aktualnya dengan hasil yang sudah diperbaiki untuk mendeteksi kesalahan. Kemudian, model akan diperbaharui berdasarkan hal tersebut. Model yang dihasilkan tidak diperlukan selama masukan tersedia, namun jika ada beberapa nilai masukan yang hilang, tidak mungkin untuk menarik kesimpulan mengenai hasilnya [12].

2.2. Decision Tree

Metode Decision Tree adalah salah satu teknik dalam data mining yang digunakan untuk membangun model prediksi dalam bentuk struktur pohon keputusan. Model ini terdiri dari node yang mewakili variabel input dan edge yang menunjukkan hubungan antara input dan output. Proses prediksi dilakukan dengan mengikuti jalur dari akar hingga daun pada pohon keputusan, menghasilkan keputusan akhir berdasarkan data yang dimasukkan. Metode ini sangat berguna dalam menganalisis dan mengelompokkan data ke dalam kategori atau kelas tertentu, karena mampu memberikan visualisasi yang jelas mengenai proses pengambilan keputusan. decision Tree banyak digunakan di berbagai bidang, seperti bisnis, kesehatan, dan teknologi informasi, berkat kemampuannya menyederhanakan klasifikasi dan prediksi.

Untuk memprediksi kinerja algoritma pada data kualitas susu, algoritma decision tree CART digunakan. Ini didasarkan pada volume data, ruang memori yang tersedia pada komputer, dan skalabilitas algoritma. CART, singkatan dari Trees of Classification and Regression, dibuat oleh Breiman pada tahun 1984 dan berfungsi untuk membangun pohon klasifikasi dan regresi. Desain pohon klasifikasi CART didasarkan pada pemisahan biner dari atribut. Ini dapat digunakan secara berurutan dan juga didasarkan pada algoritma Hunt. CART menghasilkan pengklasifikasi pohon keputusan modern dengan akurasi klasifikasi dan prediksi yang tinggi berkat fitur

dan kemampuan yang disempurnakan untuk mengatasi kekurangan CART [13].

2.3. K-Nearest Neighbors

Metode K-Nearest Neighbor (KNN) adalah salah satu teknik klasifikasi dalam supervised learning yang mengkategorikan data baru berdasarkan kedekatannya dengan data pelatihan. Algoritma ini bekerja dengan mencari sejumlah k tetangga terdekat (yang diukur berdasarkan jarak minimum, seperti Euclidean distance) dari data yang ingin diklasifikasikan. Selanjutnya, KNN menentukan kelas dari data baru tersebut dengan melihat mayoritas kelas di antara tetangga-tetangga terdekatnya. Dengan pendekatan ini, KNN memanfaatkan pola dan distribusi data pelatihan untuk memprediksi atau mengklasifikasikan data uji. Semakin dekat jarak suatu data uji dengan data pelatihan tertentu, kemungkinan besar data uji tersebut akan tergolong ke dalam kelas yang sama semakin besar.

Keunggulan utama dari metode KNN terletak pada kesederhanaan dan efektivitasnya, terutama ketika diterapkan pada dataset yang besar dan bersih. Metode ini juga menunjukkan ketahanan terhadap data yang mengandung noise, sebab pengambilan keputusan didasarkan pada mayoritas tetangga terdekat daripada satu data tunggal. Namun, KNN tidak lepas dari kelemahan. Salah satunya adalah peningkatan waktu komputasi yang signifikan seiring bertambahnya jumlah data pelatihan, karena proses klasifikasi memerlukan penghitungan jarak dari setiap data uji ke semua data latih. Selain itu, KNN juga rentan terhadap data yang redundan atau tidak relevan, karena setiap atribut tetap dipertimbangkan dalam penghitungan jarak, meskipun kontribusi mereka terhadap klasifikasi tidak signifikan [14].

2.4. Naive Bayes

Naive Bayes adalah salah satu metode klasifikasi berbasis probabilitas yang sederhana namun efektif. Algoritma ini memanfaatkan Teorema Bayes untuk memperkirakan kemungkinan sebuah data masuk ke dalam suatu kelas tertentu dengan asumsi bahwa setiap atribut saling bebas atau independen satu sama lain. Proses kerjanya diawali dengan menetapkan probabilitas awal untuk setiap kelas yang ada pada data latih. Setelah itu, dihitung peluang bersyarat (likelihood) dari setiap atribut pada masing-masing kelas. Secara matematis, probabilitas posterior ($P(H|E)$) diperoleh dari rumus Teorema Bayes, di mana ($P(E|H)$) menunjukkan peluang terjadinya data E jika hipotesis H benar, $P(H)$ adalah probabilitas awal hipotesis, dan $P(E)$ adalah probabilitas total parameter tersebut. Berbekal pendekatan ini, Naive Bayes dapat memprediksi kelas dengan memanfaatkan frekuensi kemunculan serta kombinasi nilai atribut dalam dataset secara cepat dan praktis.

Salah satu keunggulan utama Naive Bayes terletak pada kemampuannya menangani data numerik maupun kategorikal. Meski asumsi independensi antar

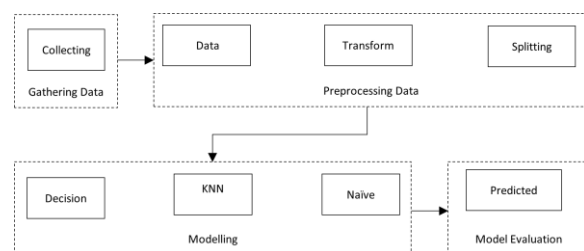
atribut sering kali tidak sepenuhnya terpenuhi di data nyata, metode ini tetap banyak digunakan karena perhitungannya efisien dan hasilnya cukup akurat. Selain itu, Naive Bayes tidak memerlukan jumlah data latih yang besar untuk mengestimasi parameter, seperti rata-rata dan variansi. Keunggulan lainnya, metode ini tetap dapat berjalan meskipun ada atribut yang nilainya kosong, karena perhitungannya dapat mengabaikan atribut yang hilang. Dari segi kecepatan dan penggunaan memori, algoritma ini juga tergolong ringan dan tangguh meskipun harus memproses atribut yang mungkin tidak relevan [15].

Namun demikian, ada beberapa kekurangan yang perlu diperhatikan. Jika suatu nilai probabilitas bersyarat bernilai nol, maka hasil prediksi juga akan otomatis bernilai nol, sehingga dapat memengaruhi keakuratan model. Selain itu, anggapan independensi atribut yang menjadi dasar Naive Bayes sering kali tidak sepenuhnya berlaku pada data dunia nyata, sehingga bisa menurunkan performa klasifikasi [16].

Secara umum, algoritma Naive Bayes memiliki tiga jenis yang sering digunakan, yaitu Multinomial Naive Bayes, Gaussian Naive Bayes, dan Bernoulli Naive Bayes. Dalam penelitian ini, metode yang digunakan adalah Gaussian Naive Bayes, karena cocok diterapkan pada data numerik kontinu. Berdasarkan dataset yang dianalisis, salah satu variabel penting yaitu pH memiliki tipe data float32, sehingga pendekatan Gaussian Naive Bayes dinilai tepat untuk mengolah dan memprediksi kualitas data tersebut.

3. METODE PENELITIAN

Perancangan penelitian digunakan untuk mengetahui alur proses penelitian yang akan dilakukan. Gambaran rancangan penelitian dapat dilihat pada Gambar 1, alur proses yang dilakukan diawali dengan pengumpulan dataset yang akan digunakan atau dikumpulkan berasal dari platform kaggle. Data yang sudah terkumpul akan dilakukan tahap preprocessing.



Gambar 1. Tahapan penelitian

3.1. Gathering Data

Dalam tahap pengolahan data, penelitian yang dilakukan saat ini menggunakan dataset yang digunakan untuk memprediksi kualitas susu dengan jumlah data keseluruhan sebanyak 1.059 data. Dimana data diambil dan diolah untuk mengambil atribut-atribut yang berkorelasi dengan pengendalian kualitas susu. Output dari metode prediksi ini adalah prediksi

terhadap kualitas susu, dengan atribut-atribut input, yaitu : pH, Temperature, Taste, Odor, Fat, Turbidity, Colour dan Grade sebagai label. Penjelasan dari atribut yang ada pada dataset milk quality beserta nilai nya terdapat pada Tabel 2.

3.2. Preprocessing Data

Langkah-langkah sebelum menguji model meliputi gathering data dan preprocessing data. Data diunduh dari Kaggle, lalu dipisahkan menjadi dua segmen, yaitu untuk training dan testing[17]. Preprocessing data memiliki tujuan untuk menghasilkan himpunan data yang bersih dan berkualitas dengan cara menangani nilai-nilai yang salah, hilang, atau tidak diinginkan. Beberapa teknik - teknik yang umumnya sering digunakan dalam melakukan preprocessing data dapat dikelompokkan ke dalam tiga kategori: cleaning data, transformasi data, dan splitting data. Setiap kategori tersebut dikaji melalui berbagai algoritma, dan penelitian menunjukkan bahwa jika tanpa penerapan preprocessing yang tepat pada data mentah, atau jika metode yang tidak sesuai digunakan, teknik analisis data tidak akan mampu mencapai tingkat kinerja yang memadai.

Dalam penelitian ini, preprocessing data meliputi pembersihan data, transformasi data, dan pemisahan data. Pembersihan data merujuk pada langkah-langkah untuk menghapus nilai yang tidak relevan atau yang tidak tersedia (null). Dalam kumpulan dataset kami, proses ini dilakukan dengan memeriksa nilai-nilai yang termasuk dalam kategori nilai yang hilang untuk setiap atribut. Hasil yang diperoleh menunjukkan status False pada setiap atribut, yang berarti tidak ada nilai yang hilang. Hal ini juga mengindikasikan bahwa kualitas dataset tersebut adalah baik.

Setelah melakukan pengecekan terhadap nilai yang hilang, kemudian dataset dibagi menjadi dua bagian melalui proses pemisahan data, yaitu data pelatihan dan data pengujian, dengan rasio 8:2. Selanjutnya, kami menentukan objek data yang terbagi, di mana masing-masing bagian terdiri dari data fitur (x) dan data target (y). Setelah pemisahan, didapati sebanyak 847 data untuk pelatihan dan 212 data untuk pengujian. Data yang diperoleh tersebut merupakan data murni yang berasal dari dataset tanpa melakukan teknik apapun dengan total keseluruhan 1059 dataset.

Tabel 1. Tabel sebelum smote tometk links

Grade	Count
Low	429
Medium	374
High	256
Total	1059

Pada tabel diatas merupakan rincian total dataset pada kelas "Grade" sebelum dilakukan teknik smote tometk links yang mana kelas tersebut merupakan label pada klasifikasi penelitian ini.

Pada tabel dibawah merupakan rincian total dataset pada kelas "Grade" setelah dilakukan teknik smote tometk links yang mana kelas tersebut merupakan label pada klasifikasi penelitian ini.

Tabel 2. Tabel setelah smote tometk links

Grade	Count
Low	429
Medium	429
High	429
Jumlah	1287

Setelah diterapkan teknik SMOTE Tomek links dataset asli nya akan di adjust secara adil dan menunjukkan sebanyak 1287 total dataset, didapati sebanyak 1029 data untuk pelatihan dan 258 data untuk pengujian. *Smote Tomek links* bekerja dengan cara menambahkan contoh dari kategori minoritas melalui smote, setelah itu menerapkan Tomek Links untuk menghapus contoh dari kategori mayoritas yang terlalu dekat dengan kategori minoritas, sehingga mengurangi adanya tumpang tindih antar kelas. Dengan cara ini, mutu dataset menjadi lebih baik, dan pada akhirnya, algoritma pembelajaran menjadi lebih efektif[18].

Tahap akhir dalam preprocessing data adalah transformasi data, yang bertujuan untuk mengubah skala pengukuran dari nilai asli menjadi bentuk yang berbeda. Proses ini diterapkan pada kumpulan data yang digunakan dalam penelitian ini. Kumpulan data tersebut sudah dalam format int64, tetapi untuk kolom "Grade" terdapat nilai *NaN*. Oleh karena itu, dilakukan transformasi data dengan memanfaatkan *LabelEncoder*, di mana nilai-nilai kategorik awalnya diubah menjadi nilai numerik. Proses ini berfungsi untuk mencegah overfitting, menjamin evaluasi kinerja model yang tepat, dan mendukung pemilihan parameter model yang terbaik.

3.3. Modeling

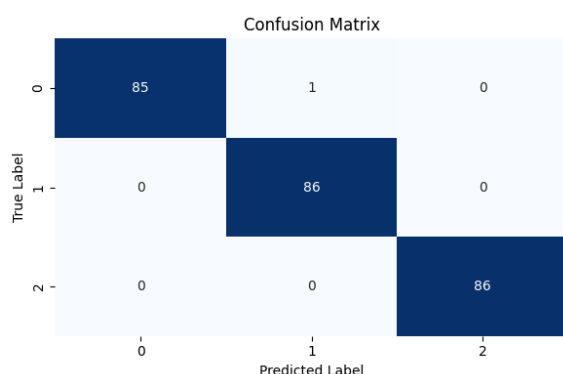
Dalam penelitian ini, tahap awal pengembangan model akan difokuskan pada penerapan tiga algoritma supervised learning, yaitu Decision Tree, K-Nearest Neighbors (KNN), dan Naive Bayes, untuk mengembangkan dan mengevaluasi model klasifikasi kualitas susu. Setiap algoritma akan dilatih menggunakan data pelatihan dan diuji pada data pengujian guna mengukur kinerjanya secara objektif. Perbandingan antara ketiga metode ini akan dilakukan dengan menganalisis empat metrik utama: akurasi, yang mengukur ketepatan prediksi secara keseluruhan; presisi, untuk menentukan sejauh mana prediksi positif benar-benar relevan; recall, yang menilai kemampuan model dalam mengenali semua kasus positif; dan F1-score, yang mencerminkan keseimbangan antara presisi dan recall. Dengan pendekatan ini, penelitian ini bertujuan untuk mengidentifikasi algoritma yang paling unggul dalam konteks klasifikasi kualitas susu, baik dari segi akurasi prediksi maupun kestabilan kinerja, sehingga dapat merekomendasikan metode

yang optimal untuk penerapan sistem kendali mutu otomatis.

Tahapan dalam melakukan modeling pada masing-masing algoritma yaitu inisialisasi model terlebih dahulu kemudian model di fit untuk melakukan proses prediksi atau inferensi. Langkah selanjutnya melakukan confusion matrix untuk melakukan evaluasi kinerja model, hal pertama yang di evaluasi merupakan data train. Tujuannya adalah untuk mengetahui akurasi dan melakukan prediksi pada tiap kelas, kemudian akan terlihat berapa banyak data yang diklasifikasikan dengan benar, berapa banyak yang salah dan salahnya ke arah mana (kelas mana). Langkah terakhir modeling yakni melakukan evaluasi menggunakan classification report, didalam nya akan menggunakan metrik-metrik penting dalam klasifikasi yakni accuracy, precision, recall, dan F1 score, ini adalah metrik dari tahapan evaluasi untuk mengukur performa model secara komprehensif.

4. HASIL DAN PEMBAHASAN

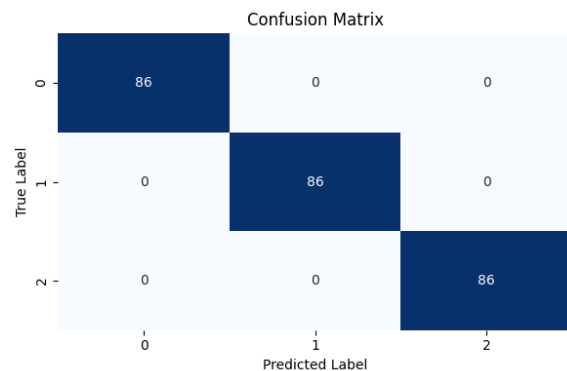
Setelah menyelesaikan pemodelan, penilaian dilakukan terhadap tiga algoritma klasifikasi yang digunakan, yaitu decision tree, k-nearest neighbors (KNN), dan naive bayes. Evaluasi ini bertujuan untuk menilai seberapa efektif masing-masing model dalam mengklasifikasikan mutu susu berdasarkan data yang tersedia. Dalam evaluasi tersebut, diterapkan metode matriks kebingungan untuk ketiga model, yang memungkinkan analisis lebih mendalam mengenai prediksi yang benar dan salah dari setiap model. Dengan memanfaatkan confusion matrix, kita dapat menghitung metrik evaluasi seperti tingkat akurasi, presisi, recall, dan F1-score, yang memberikan gambaran komprehensif mengenai seberapa baik model dalam mengklasifikasikan data dengan tepat. Proses penilaian ini sangat penting untuk mengetahui sejauh mana model mampu menyesuaikan diri dengan data baru yang belum pernah ditemui sebelumnya.



Gambar 2. Confusion matrix decision tree

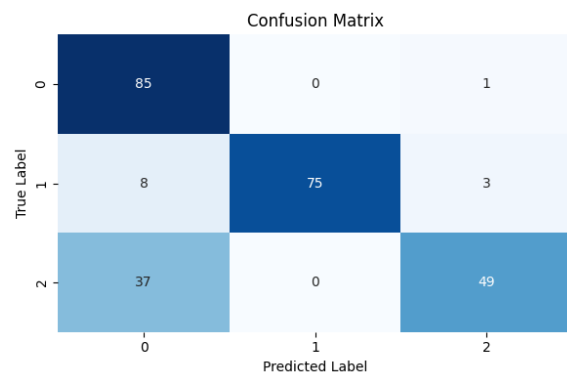
Dari total 258 data tes berdasarkan (85 + 1 + 86 + 86) hasil confusion matrix model decision tree menunjukkan bahwa ada satu kesalahan kecil dan hampir semua prediksinya akurat, yang ditandai oleh angka-angka besar dan warna biru gelap di sepanjang garis diagonal, sementara area lainnya berwarna biru

terang (menunjukkan angka 0 dan 1 yang sangat minim).



Gambar 3. Confusion matrix k-nearest neighbors

Dari total 258 data tes berdasarkan (86 + 86 + 86) hasil confusion matrix model k-nearest neighbors menunjukkan bahwa tidak terdapat kesalahan kecil dan semua prediksinya akurat, yang ditandai oleh angka-angka besar dan warna biru gelap di sepanjang garis diagonal, yang artinya model tidak melakukan kesalahan satu pun saat melakukan prediksi.



Gambar 4. Confusion matrix naive bayes

Dari total 258 data tes berdasarkan (85 + 1 + 8 + 75 + 37 + 49) hasil confusion matrix model naive bayes menunjukkan bahwa terdapat beberapa kesalahan, yang ditandai oleh angka-angka besar dan warna biru gelap di sepanjang garis diagonal, yang artinya model tidak melakukan kesalahan satu pun saat melakukan prediksi.

Tahap evaluasi confusion matrix sebagai alat untuk menilai performa dari setiap model klasifikasi yang telah dibangun. Confusion matrix menyajikan informasi detail mengenai jumlah prediksi benar dan salah dari masing-masing kelas target, sehingga memungkinkan untuk mengevaluasi seberapa akurat model dapat membedakan kualitas susu. Dalam penelitian ini, metrik seperti akurasi, presisi, recall, dan f1 score dihitung berdasarkan nilai-nilai dalam matriks confusion. Keempat metrik ini dipilih karena mereka dapat memberikan gambaran menyeluruh tentang kinerja model, baik dari segi ketepatan

maupun kemampuan model untuk mengidentifikasi masing-masing kelas secara akurat.

Penggunaan confusion matrix juga sangat berperan dalam proses perbandingan berbagai algoritma, karena hasil evaluasi tidak hanya melihat dari nilai akurasi, tetapi juga memperhatikan sebaran kesalahan klasifikasi pada setiap kategori. Sebagai contoh, sebuah model dengan akurasi yang tinggi belum tentu dapat mengenali semua kategori secara seimbang apabila terdapat ketidakseimbangan dalam presisi atau recall. Oleh karena itu, dengan menggunakan confusion matrix, evaluasi yang lebih komprehensif dan adil dapat dilakukan terhadap kemampuan model dalam mengklasifikasikan kualitas susu. Dengan demikian, pemilihan algoritma yang paling tepat tidak hanya didasarkan pada hasil akhir, tetapi juga memperhitungkan konsistensi dan keandalan model.

Dalam studi ini, evaluasi model dilakukan pada ketiga algoritma (Decision Tree, KNN, dan Naive Bayes) dengan menggunakan metrik penting dalam klasifikasi, yaitu untuk mengetahui seberapa banyak prediksi yang tepat dari total prediksi (akurasi), dari semua prediksi positif, berapa yang benar-benar positif (presisi), dari seluruh data yang sebenarnya positif, berapa yang berhasil diprediksi sebagai positif (recall), serta rata-rata harmonis dari Precision dan Recall, yang bermanfaat ketika data tidak seimbang (F1-score).

Tabel 3. Kinerja model decision tree

Report	Milk Quality			Accuracy
	Low	Medium	High	
Precision	1.00	0.99	1.00	1.00
Recall	0.99	1.00	1.00	
F1-Score	0.99	0.99	1.00	
Support	86	86	86	

Tabel 3 menyajikan laporan hasil klasifikasi dari evaluasi model Decision Tree. Kinerja model Decision Tree menunjukkan akurasi 1.00, yang mengindikasikan hasil yang sangat memuaskan.

Tabel 4. Kinerja model k-nearest neighbors

Report	Milk Quality			Accuracy
	Low	Medium	High	
Precision	1.00	1.00	1.00	1.00
Recall	1.00	1.00	1.00	
F1-Score	1.00	1.00	1.00	
Support	86	86	86	

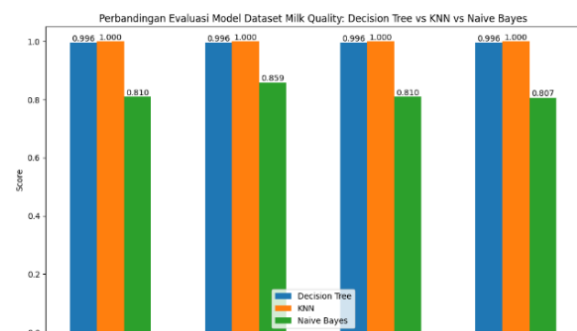
Pada Tabel 4 merupakan merupakan hasil klasifikasi report pengujian model K-Nearest Neighbors. Hasil kinerja model K-Nearest Neighbors menunjukkan accuracy 1.00 menjadikan hasil yang sangat baik. Namun ada sedikit perbedaan hasil precision, recall, dan f1 score yang ditunjukkan jika dibandingkan dengan model decision tree.

Tabel 5. Kinerja model naive bayes

Report	Milk Quality			Accuracy
	Low	Medium	High	
Precision	0.65	1.00	0.92	0.81
Recall	0.99	0.87	0.93	
F1-Score	0.79	0.93	0.71	
Support	86	86	86	

Pada Tabel 5 merupakan hasil klasifikasi report pengujian model Naive Bayes. Hasil kinerja model Naive Bayes menunjukkan accuracy 0.81 menjadikan hasil yang cukup baik. Tetapi hasil pada model ini tergolong yang paling rendah jika dibandingkan dengan Decision Tree dan KNN.

Hasil kinerja masing-masing model yang telah dilakukan pengujian dan evaluasi mendapatkan hasil yang cukup baik terutama pada model k-nearest neighbors yang mendapati hasil paling sempurna meskipun saat evaluasi data train ada beberapa kesalahan kecil namun berbeda saat dilakukan menggunakan data test yang mana model dalam melakukan prediksi pada data test tidak melakukan kesalahan.



Gambar 5. Diagram hasil perbandingan model

Namun demikian, hasil dari pengujian dengan menggunakan metrik evaluasi untuk menilai seberapa baik kinerja model menunjukkan bahwa dengan penerapan metode yang sama pada Decision Tree menghasilkan rata-rata akurasi sebesar 99%, presisi sebesar 99%, recall sebesar 99%, dan f1-score sebesar 99%. Sementara itu, model algoritma K-Nearest Neighbors menunjukkan rata-rata akurasi sebesar 100%, presisi sebesar 100%, recall sebesar 100%, dan f1-score sebesar 100%. Sedangkan pada model algoritma Naive Bayes, rata-rata akurasi yang diperoleh adalah 81%, presisi 85%, recall 81%, dan f1-score 80%.

5. KESIMPULAN DAN SARAN

Berdasarkan hasil yang telah diperoleh sebelumnya dapat dilihat bahwa model KNN menunjukkan hasil paling sempurna dan paling akurat yakni rata-rata sebesar 100%. Model KNN menunjukkan performa terbaik dalam klasifikasi kualitas susu berdasarkan fitur kimia, dengan akurasi 100% pada data uji. Temuan ini menunjukkan potensi

signifikan penerapan algoritma KNN dalam sistem pengawasan mutu susu secara otomatis.

Perbandingan dilakukan dengan 2 model lainnya (Decision Tree dan Naive Bayes) yang masih menunjukkan beberapa kesalahan kecil pada model. Pada decision tree sebenarnya menunjukkan hasil yang sangat baik karena model hanya salah memprediksi sebanyak satu kali pada label 0 sehingga model decision tree mendapatkan akurasi sebesar 99% pada hasil perhitungan evaluasi menggunakan metrik-metrik dalam klasifikasi, selanjutnya pada model naive bayes terlihat sangat lemah pada label 2, model mengalami kesulitan besar untuk mengenali data dan pada evaluasi menggunakan data test model sering menebaknya sebagai label 0 sebanyak 37 kali daripada menebaknya dengan benar sebanyak 49 kali karena hal tersebut model hanya memiliki rata rata akurasi sebesar 81%. Kesalahan-kesalahan yang dibuat oleh model dapat dianggap cukup normal karena terdapat berbagai faktor lain yang turut mempengaruhi hasil dari model tersebut, yang menyebabkan terjadinya kesalahan.

Dalam studi ini, meskipun model naive bayes menunjukkan banyak kesalahan dalam melakukan prediksi baik pada data latihan maupun data pengujian, pasti ada berbagai faktor yang mempengaruhi hal itu. Peneliti yang akan datang diharapkan dapat menerapkan pendekatan lain untuk mencapai hasil yang lebih memuaskan. Selanjutnya, pada model k-nearest neighbors (KNN), meskipun hasil evaluasi dari data pengujian menunjukan performa yang sangat baik, peneliti di masa depan disarankan untuk menerapkan metode lain dalam pengolahan data atau menambah jumlah dataset di setiap kelas untuk menyeimbangkan data yang dimiliki.

DAFTAR PUSTAKA

- [1] N. Wiranti *et al.*, “KUALITAS SUSU SAPI SEGAR PADA PEMERAHAN PAGI DAN SORE Quality of Fresh Cow’s Milk at Morning and Afternoon Milking,” *Jurnal Riset dan Inovasi Peternakan*, vol. 6, no. 2, pp. 2598–3067, 2022, doi: 10.23960/jrip.2022.6.2.123-128.
- [2] M. Pramesti, A. Fadlan, and M. Yasin, “Konsep Industrialisasi Pada Pengembangan Teknologi Di Indonesia,” vol. 2, no. 2, pp. 148–154, 2023, doi: 10.58192/populer.v2i2.
- [3] R. I. Mukhamediev *et al.*, “Review of Artificial Intelligence and Machine Learning Technologies: Classification, Restrictions, Opportunities and Challenges,” Aug. 01, 2022, *MDPI*. doi: 10.3390/math10152552.
- [4] S. C. Huang, A. Pareek, M. Jensen, M. P. Lungren, S. Yeung, and A. S. Chaudhari, “Self-supervised learning for medical image classification: a systematic review and implementation guidelines,” Dec. 01, 2023, *Nature Research*. doi: 10.1038/s41746-023-00811-0.
- [5] I. Irwansyah, A. D. Wiranata, and T. T. M., “KOMPARASI ALGORITMA DECISION TREE, NAIVE BAYES DAN K-NEAREST NEIGHBOR UNTUK MENENTUKAN KUALITAS UDARA DI PROVINSI DKI JAKARTA,” *Infotech: Journal of Technology Information*, vol. 9, no. 2, pp. 193–198, Nov. 2023, doi: 10.37365/jti.v9i2.203.
- [6] K. Kunci, “PENGARUH KUALITAS PRODUK DAN CITRA MEREK TERHADAP KEPUTUSAN PEMBELIAN SUSU INDOMILK THE EFFECT OF PRODUCT QUALITY AND BRAND IMAGE ON PURCHASING DECISIONS.”
- [7] D. Bhavsar, Y. Jobanputra, N. K. Swain, and D. Swain, “Milk Quality Prediction Using Machine Learning,” *EAI Endorsed Transactions on Internet of Things*, vol. 10, 2024, doi: 10.4108/eetiot.4501.
- [8] A. Çelik, “Using Machine Learning Algorithms to Detect Milk Quality,” 2022.
- [9] F. Wardah, A. 1✉, W. Sumarjaya, and E. N. Kencana, “Seleksi Fitur Information Gain Untuk Klasifikasi Kualitas Susu Sapi Menggunakan Metode K-Nearest Neighbor Dan Naïve Bayes,” *INNOVATIVE: Journal Of Social Science Research*, vol. 5, pp. 422–435, 2025.
- [10] N. Suhandi, R. Gustriansyah, A. Destria, M. Amalia, and V. Kris, “Prediksi Kualitas Susu Menggunakan Metode K-Nearest Neighbors Milk Quality Prediction Using The K-Nearest Neighbors Method,” vol. 14, no. 2, 2024, doi: 10.30700/sisfotenika.v14i2.430.
- [11] P. Santoso, H. Abijono, and N. L. Anggreini, “ALGORITMA SUPERVISED LEARNING DAN UNSUPERVISED LEARNING DALAM PENGOLAHAN DATA,” *Unira Malang* |, vol. 4, no. 2, 2021.
- [12] V. Nasteski, “An overview of the supervised machine learning methods,” *HORIZONS.B*, vol. 4, pp. 51–62, Dec. 2017, doi: 10.20544/horizons.b.04.1.17.p05.
- [13] A. Priyam, R. Gupta, A. Rathee, and S. Srivastava, “Comparative Analysis of Decision Tree Classification Algorithms,” 2013. [Online]. Available: <http://inpressco.com/category/ijcet>
- [14] A. Setiawan, “Perbandingan Penggunaan Jarak Manhattan, Jarak Euclid, dan Jarak Minkowski dalam Klasifikasi Menggunakan Metode KNN pada Data Iris,” *Jurnal Sains dan Edukasi Sains*, vol. 5, no. 1, pp. 28–37, May 2022, doi: 10.24246/juses.v5i1p28-37.
- [15] M. Ismail, N. Hassan, and S. S. Bafjaish, “Comparative Analysis of Naive Bayesian Techniques in Health-Related for Classification Task,” *Journal of Soft Computing and Data Mining*, vol. 1, no. 2, pp. 1–10, Dec. 2020, doi: 10.30880/jscdm.2020.01.02.001.
- [16] A. Yunizar, P. Yusuf, and R. Sari, “Implementasi Algoritma Naïve Bayes untuk Klasifikasi

- Pemahaman Program MBKM bagi Mahasiswa,” 2022. [Online]. Available: <https://colab.research.google.com/drive/1rMH8-MzVSaGWYSeFo7QcHiTuJcOfJag2>.
- [17] F. Pradana Rachman, H. Santoso, and A. History, “Jurnal Teknologi dan Manajemen Informatika Perbandingan Model Deep Learning untuk Klasifikasi Sentiment Analysis dengan Teknik Natural Language Processing Article Info ABSTRACT,” vol. 7, no. 2, pp. 103–112, 2021, [Online]. Available: <http://http://jurnal.unmer.ac.id/index.php/jtmi>
- [18] D. R. I. M. Setiadi, K. Nugroho, A. R. Muslikh, S. W. Iriananda, and A. A. Ojugo, “Integrating SMOTE-Tomek and Fusion Learning with XGBoost Meta-Learner for Robust Diabetes Recognition,” *Journal of Future Artificial Intelligence and Technologies*, vol. 1, no. 1, pp. 23–38, May 2024, doi: 10.62411/faith.2024-11