

Analisis Tingkat Kepuasan Mahasiswa Pada Aplikasi Sistem Informasi Terpadu (SITU) di Universitas Labuhanbatu Menggunakan Metode *Naïve Bayes* Dan *Decision Tree*

¹Rista Andini Ritonga, ²Syaiful Zuhri Harahap, ³Irmayanti, ⁴Angga Putra Juledi

^{1,2,3,4}Sistem Informasi, Fakultas Sains dan Teknologi, Universitas Labuhanbatu

Email: [¹ristaandinirtg@gmail.com](mailto:ristaandinirtg@gmail.com), [²syaifulzuhriharahap@gmail.com](mailto:syaifulzuhriharahap@gmail.com),
[³irmayantiritonga2@gmail.com](mailto:irmayantiritonga2@gmail.com), [⁴anggapj19@gmail.com](mailto:anggapj19@gmail.com)

Corresponding Author : ristaandinirtg@gmail.com

Abstract

The Integrated Information System (SITU) is an application used to support various academic services at Universitas Labuhanbatu. The quality of services provided by this application needs to be evaluated to determine the level of student satisfaction as its users. This study aims to analyze student satisfaction with the use of the Integrated Information System (SITU) using the Naïve Bayes and Decision Tree methods. The research applies the Knowledge Discovery in Databases (KDD) process, which consists of Selection, Preprocessing, Transformation, Data Mining, Evaluation, and Interpretation. The research data were collected through questionnaires distributed to 50 students of the Faculty of Science and Technology at Universitas Labuhanbatu. The research variables include System Ease of Use, System Access Speed, Information Accuracy, User Interface, and System Reliability, while the target variable is the student satisfaction level, classified into Satisfied and Dissatisfied categories. The classification process was carried out using Orange Data Mining software and evaluated using a Confusion Matrix. The interpretation results based on 27 testing data showed that the Decision Tree algorithm classified 19 instances as Satisfied and 8 instances as Dissatisfied, while the Naïve Bayes algorithm classified 18 instances as Satisfied and 9 instances as Dissatisfied. Furthermore, the Confusion Matrix evaluation indicated that the Naïve Bayes method achieved a 96.8% prediction accuracy for the Satisfied category, outperforming the Decision Tree method, which achieved 81.6%. Based on these results, the Naïve Bayes method demonstrated superior classification performance in analyzing student satisfaction with the Integrated Information System (SITU). The findings of this study are expected to serve as a reference for Universitas Labuhanbatu in evaluating and improving the quality of services provided through the Integrated Information System (SITU).

Keywords: *Student Satisfaction, Data Mining, Naïve Bayes, Decision Tree, Integrated Information System.*

1. Pendahuluan

Tingkat kepuasan mahasiswa merupakan penilaian mahasiswa terhadap kualitas layanan yang diberikan oleh perguruan tinggi berdasarkan harapan dan pengalaman yang diperoleh (Nopiyanti dkk, 2024). Tingkat kepuasan mahasiswa adalah ukuran keberhasilan institusi dalam memenuhi kebutuhan dan harapan mahasiswa terhadap

layanan akademik maupun nonakademik (Assyahri & Mardaus, 2023). Kepuasan mahasiswa merupakan bentuk evaluasi mahasiswa setelah menggunakan fasilitas atau layanan yang disediakan oleh perguruan tinggi (Maryam, 2024). Dalam sistem informasi, tingkat kepuasan mahasiswa menunjukkan sejauh mana sistem mampu memenuhi kebutuhan pengguna secara efektif dan efisien (Kartini dkk, 2024). Tingkat kepuasan mahasiswa juga menjadi indikator penting dalam mengevaluasi dan meningkatkan kualitas pelayanan di perguruan tinggi (Kurniawan & Farhatuaini, 2024). Aplikasi Sistem Informasi Terpadu (SITU) di Universitas Labuhanbatu digunakan sebagai media utama dalam pelayanan akademik mahasiswa. Namun, masih terdapat berbagai kendala yang memengaruhi tingkat kepuasan mahasiswa, seperti kemudahan penggunaan sistem (X1) yang belum optimal, kecepatan akses sistem (X2) yang terkadang lambat, keakuratan informasi (X3) yang perlu ditingkatkan, tampilan antarmuka atau *user interface* (X4) yang kurang nyaman digunakan, serta keandalan sistem (X5) yang masih mengalami gangguan pada kondisi tertentu. Permasalahan tersebut menyebabkan tingkat kepuasan mahasiswa (Y) menjadi beragam, yaitu berada pada kategori Puas dan Tidak Puas, sehingga diperlukan analisis untuk mengetahui faktor-faktor yang paling berpengaruh terhadap tingkat kepuasan mahasiswa. Untuk mengatasi permasalahan tersebut, penelitian ini menerapkan teknik Data Mining menggunakan metode Naïve Bayes dan Decision Tree. Kedua metode digunakan untuk mengklasifikasikan tingkat kepuasan mahasiswa berdasarkan lima variabel, yaitu kemudahan penggunaan sistem (X1), kecepatan akses sistem (X2), keakuratan informasi (X3), tampilan antarmuka (*user interface*) (X4), dan keandalan sistem (X5). Hasil klasifikasi akan menghasilkan kategori tingkat kepuasan mahasiswa (Y), yaitu Puas atau Tidak Puas. Selain menghasilkan model klasifikasi, penelitian ini juga dapat mengidentifikasi variabel yang paling berpengaruh terhadap kepuasan mahasiswa, sehingga hasilnya dapat dijadikan sebagai dasar evaluasi dan peningkatan kualitas layanan aplikasi Sistem Informasi Terpadu (SITU) di Universitas Labuhanbatu. Data Mining adalah proses menemukan pola atau informasi yang bermanfaat dari sekumpulan data dalam jumlah besar (Mardiansa dkk, 2024). Data Mining merupakan teknik analisis data untuk menghasilkan pengetahuan yang dapat digunakan dalam pengambilan keputusan (Nurani dkk, 2023). Data Mining adalah bagian dari proses Knowledge Discovery in Databases (KDD) yang bertujuan menggali informasi tersembunyi dari data (Rahmadini dkk, 2023). Data Mining memanfaatkan algoritma statistik dan machine learning untuk mengolah data menjadi informasi yang bernilai (Yani dkk, 2023). Data Mining digunakan untuk klasifikasi, prediksi, pengelompokan, dan analisis pola data (Siregar dkk, 2023). Naïve Bayes merupakan algoritma klasifikasi yang menggunakan Teorema Bayes dalam menentukan suatu kelas (Huriah & Nuris, 2023). Naïve Bayes mengasumsikan setiap atribut bersifat independen dalam proses klasifikasi (Adzy dkk, 2023). Metode Naïve Bayes menghitung probabilitas setiap kelas berdasarkan data pelatihan (Sukriadi dkk, 2023). Naïve Bayes dikenal memiliki proses klasifikasi yang sederhana, cepat, dan efisien (Hirwono dkk, 2023). *Decision Tree* merupakan algoritma klasifikasi yang menyajikan proses pengambilan keputusan dalam bentuk pohon (Aktavera & Wijaya, 2023). *Decision Tree* membagi data berdasarkan atribut yang memiliki pengaruh terbesar terhadap hasil klasifikasi (Wirasena dkk, 2024). Metode *Decision Tree* menghasilkan aturan keputusan yang mudah dipahami dan diinterpretasikan (Hasanah & Fatah, 2024). *Decision Tree* banyak digunakan untuk klasifikasi, prediksi, dan analisis pengambilan keputusan (Kirana dkk, 2024). Beberapa penjelasan pada penelitian terdahulu yang digunakan. Penelitian pertama metode Naïve Bayes berhasil mengklasifikasikan kepuasan pengguna Aplikasi UMSU Academy dan

divalidasi melalui perhitungan manual, sehingga dapat menjadi acuan evaluatif bagi pengembang (Simbolon & Riza, 2025). Penelitian kedua perbandingan Naïve Bayes dan SVM menunjukkan kedua metode mampu mengklasifikasikan kepuasan mahasiswa dengan baik, namun masih diperlukan optimasi parameter untuk meningkatkan kinerja pada kelas minoritas (Antika dkk, 2025). Penelitian ketiga model Decision Tree C4.5 efektif memprediksi kepuasan mahasiswa terhadap pelayanan akademik dan mampu menemukan hubungan antar variabel secara baik (Helma dkk, 2025). Model Decision Tree mengidentifikasi dimensi Assurance dan Reliability sebagai faktor dominan kepuasan mahasiswa pada pembelajaran hybrid, serta memberikan kerangka evaluasi berbasis data untuk peningkatan efektivitas pembelajaran (Zaman dkk, 2023).

2. Landasan Teori

Tingkat Kepuasan Mahasiswa

Tingkat kepuasan mahasiswa merupakan tingkat penilaian mahasiswa terhadap kualitas layanan yang diberikan oleh perguruan tinggi berdasarkan kesesuaian antara harapan dengan pengalaman yang diperoleh selama menggunakan layanan tersebut. Kepuasan mahasiswa mencerminkan persepsi mahasiswa terhadap kemampuan institusi dalam memberikan pelayanan akademik maupun nonakademik secara efektif dan efisien. Dalam konteks sistem informasi, tingkat kepuasan mahasiswa menunjukkan sejauh mana suatu sistem mampu memenuhi kebutuhan pengguna melalui kemudahan penggunaan, kecepatan akses, keakuratan informasi, serta keandalan layanan yang diberikan (Zain dkk, 2024).

Data Mining

Data Mining merupakan proses menggali informasi atau pengetahuan yang bermanfaat dari sekumpulan data berukuran besar dengan menggunakan berbagai teknik analisis. Proses ini bertujuan untuk menemukan pola, hubungan, atau kecenderungan yang sebelumnya tidak diketahui sehingga dapat menghasilkan informasi yang berguna dalam pengambilan keputusan. Data Mining merupakan salah satu tahapan dalam proses *Knowledge Discovery in Databases* (KDD) yang melibatkan proses seleksi data, pembersihan data, transformasi data, penambangan data, hingga evaluasi hasil (Qamal dkk, 2023).

Naive Bayes

Naïve Bayes merupakan salah satu algoritma klasifikasi dalam Data Mining yang didasarkan pada Teorema Bayes untuk menghitung probabilitas suatu data terhadap kelas tertentu. Metode ini bekerja dengan asumsi bahwa setiap atribut bersifat independen atau tidak saling memengaruhi dalam proses klasifikasi. Meskipun menggunakan asumsi yang sederhana, Naïve Bayes mampu menghasilkan tingkat akurasi yang baik serta memiliki proses komputasi yang cepat, sehingga banyak digunakan dalam pengolahan data berukuran besar (Olivia dkk, 2023).

$$P(H|X) = \frac{P(X|H)}{\sum_{t=1}^n P(H_t|X)} \cdot P(H) \dots\dots\dots (1)$$

Penjelasan dari persamaan rumus 1, yaitu :

X : Data dengan class yang belum diketahui
 H : Hipotesis data merupakan suatu class spesifik

$P(H|X)$: Probabilitas hipotesis H berdasarkan kondisi X
 $P(H)$: Probabilitas hipotesis H
 $P(X|H)$: Probabilitas X berdasarkan hipotesis H
 $P(X)$: Probabilitas X

Decision Tree

Decision Tree merupakan salah satu algoritma klasifikasi dalam Data Mining yang menyajikan proses pengambilan keputusan dalam bentuk struktur pohon (*tree*). Metode ini bekerja dengan membagi data berdasarkan atribut yang memiliki pengaruh terbesar sehingga menghasilkan aturan-aturan keputusan yang mudah dipahami dan diinterpretasikan. Setiap simpul (*node*) pada pohon menunjukkan atribut yang diuji, setiap cabang (*branch*) menunjukkan hasil pengujian, sedangkan daun (*leaf*) menunjukkan kelas atau keputusan akhir (Ameliatus S & Fatah, 2024).

$$Entropy(s) = \sum_{i=1}^n -p_i * \log_2 p_i \dots\dots\dots (1)$$

Keterangan :

S = Himpunan Kasus

n = Jumlah partisi S

Pi = Proporsi Si, terhadap S

Kemudian hitung nilai gain dengan rumus :

$$Gain(S,A) = Entropy(S) - \sum_{n=1}^c -\frac{S_i}{S} * Entropy(S_i) \dots\dots\dots (2)$$

Keterangan :

S = Himpunan Kasus

A = Fitur

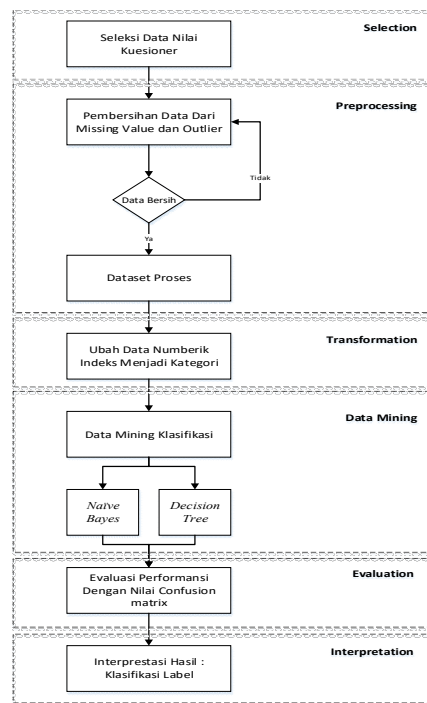
n = Jumlah partisi atribut A

|Si| = Proporsi Si terhadap S

|S| = Jumlah kasus dalam S

3. Metode Penelitian

Tahapan metodologi penelitian merupakan rangkaian langkah sistematis yang dilakukan untuk memastikan bahwa proses penelitian berjalan secara terstruktur, terarah, dan sesuai dengan tujuan yang telah ditetapkan. Dalam penelitian ini, tahapan pelaksanaan disusun untuk mengolah dan menganalisis data kepuasan mahasiswa terhadap penggunaan Sistem Informasi Terpadu (SITU) Universitas Labuhanbatu secara ilmiah dan objektif. Setiap tahapan memiliki peran penting dalam menjamin kualitas data serta keakuratan hasil analisis yang dihasilkan. Oleh karena itu, tahapan penelitian disusun mengikuti alur proses data mining yang mencakup pemilihan data, pembersihan data, transformasi data, penerapan algoritma klasifikasi, hingga evaluasi hasil. Adapun tahapan pelaksanaan penelitian yang digunakan dalam penelitian ini terdiri dari *Selection*, *Preprocessing*, *Transformation*, *Data Mining*, dan *Evaluation*, yang masing-masing dijelaskan secara rinci sebagai berikut.



Gambar 1. Tahapan Penelitian

4. Hasil Dan Pembahasan

Selection

Tahap Selection merupakan tahap awal dalam proses *Knowledge Discovery in Databases* (KDD) yang bertujuan untuk memilih data yang relevan sesuai dengan kebutuhan penelitian. Pada tahap ini, dilakukan proses pemilihan data yang berasal dari hasil penyebaran kuesioner kepada mahasiswa Fakultas Sains dan Teknologi Universitas Labuhanbatu sebagai pengguna Aplikasi Sistem Informasi Terpadu (SITU). Data yang dipilih hanya mencakup atribut-atribut yang memiliki hubungan langsung dengan tujuan penelitian, yaitu untuk menganalisis tingkat kepuasan mahasiswa menggunakan metode Naïve Bayes dan Decision Tree. Adapun data yang digunakan pada tahap *Selection* meliputi variabel independen (X), variabel dependen (Y), serta skala penilaian kuesioner yang digunakan sebagai dasar dalam proses klasifikasi.

Variabel X (Variabel Independen)

Variabel X merupakan variabel bebas yang digunakan sebagai atribut dalam proses klasifikasi tingkat kepuasan mahasiswa terhadap Aplikasi Sistem Informasi Terpadu (SITU). Adapun rincian variabel X disajikan pada tabel dibawah ini.

Tabel 1. Atribut

No.	Atribut	Keterangan	Nilai
1	Kemudahan Penggunaan Sistem	Sangat Baik	5
		Baik	4
		Cukup	3
		Kurang	2
		Sangat Kurang	1
2	Kecepatan Akses Sistem	Sangat Baik	5
		Baik	4

		Cukup	3
		Kurang	2
		Sangat Kurang	1
3	Keakuratan Informasi	Sangat Baik	5
		Baik	4
		Cukup	3
		Kurang	2
		Sangat Kurang	1
4	Tampilan Antarmuka (User Interface)	Sangat Baik	5
		Baik	4
		Cukup	3
		Kurang	2
		Sangat Kurang	1
5	Keandalan Sistem	Sangat Baik	5
		Baik	4
		Cukup	3
		Kurang	2
		Sangat Kurang	1

Variabel Y (Variabel Dependen)

Variabel Y merupakan variabel terikat yang berfungsi sebagai kelas (*class label*) dalam proses klasifikasi. Variabel ini menunjukkan tingkat kepuasan mahasiswa terhadap penggunaan Aplikasi Sistem Informasi Terpadu (SITU).

Tabel 2. Label Kelas

No.	Tingkat Kepuasan	Nilai
1	Tidak Puas	< 0.25
2	Puas	> 0.25

Nilai Kuesioner

Nilai kuesioner merupakan skala penilaian yang digunakan oleh responden untuk memberikan penilaian terhadap setiap indikator penelitian. Skala penilaian ini menggunakan skala Likert yang mengubah persepsi responden menjadi data kuantitatif.

Tabel 3 Nilai Kuesioner

No	Nama Mahasiswa	Kemudahan Penggunaan Sistem	Kecepatan Akses Sistem	Keakuratan Informasi	Tampilan Antarmuka (User Interface)	Keandalan Sistem	Average	Kepuasan
1	Karmila Nasution	5	5	5	5	5	5.00	Puas
2	Elfi Zahra Yuni	5	5	5	5	5	5.00	Puas
3	Indah Febriana	5	5	5	5	5	5.00	Puas
4	Lisa Ariani	3	5	5	3	3	3.80	Puas
5	Yepie Chantika Mysuka	4	4	5	3	4	4.00	Puas
6	Rina Harahap	2	2	4	4	4	3.20	Puas
7	Riska Yugianti	5	5	5	5	5	5.00	Puas
8	Harun	2	3	4	3	4	3.20	Puas
9	Monica Pratiwi	2	2	4	2	1	2.20	Tidak Puas
10	Dita Dwi Cahya	2	3	5	3	2	3.00	Puas
...	...	...	...	...	...	...	...	...
77	Mariani	1	3	3	3	1	2.20	Tidak Puas

Preprocessing

Tahap Preprocessing merupakan proses pembersihan (*data cleaning*) dan penyiapan data sebelum dilakukan proses klasifikasi. Pada tahap ini, data hasil penyebaran kuesioner diperiksa untuk memastikan tidak terdapat data yang kosong

(*missing value*), data ganda (*duplicate*), maupun kesalahan dalam pengisian kuesioner. Setelah proses pembersihan selesai, diperoleh sebanyak 77 data yang layak digunakan dalam penelitian. Selanjutnya, data tersebut dibagi menjadi dua kelompok, yaitu 50 data training yang digunakan untuk membangun model klasifikasi dan 27 data testing yang digunakan untuk menguji performa model yang telah dibentuk. Pembagian data ini bertujuan agar model yang dihasilkan dapat dievaluasi berdasarkan data yang belum pernah digunakan pada saat proses pelatihan.

Data Training

Data training merupakan kumpulan data yang digunakan untuk membentuk model klasifikasi menggunakan metode Naïve Bayes dan Decision Tree. Pada penelitian ini, jumlah data training sebanyak 50 data, yang berisi atribut-atribut penelitian serta kelas tingkat kepuasan mahasiswa.

Tabel 4. Data Training

No	Kemudahan Penggunaan Sistem	Kecepatan Akses Sistem	Keakuratan Informasi	Tampilan Antarmuka (User Interface)	Keandalan Sistem	Average	Kepuasan
1	5	5	5	5	5	5.00	Puas
2	5	5	5	5	5	5.00	Puas
3	5	5	5	5	5	5.00	Puas
4	3	5	5	3	3	3.80	Puas
5	4	4	5	3	4	4.00	Puas
6	2	2	4	4	4	3.20	Puas
7	5	5	5	5	5	5.00	Puas
8	2	3	4	3	4	3.20	Puas
9	2	2	4	2	1	2.20	Tidak Puas
10	2	3	5	3	2	3.00	Puas
11	2	2	4	2	2	2.40	Tidak Puas
12	1	1	3	1	1	1.40	Tidak Puas
13	3	3	5	2	2	3.00	Puas
14	2	3	4	3	1	2.60	Puas
15	2	1	5	2	3	2.60	Puas
16	3	2	5	1	3	2.80	Puas
17	2	1	5	2	1	2.20	Tidak Puas
18	1	2	5	2	1	2.20	Tidak Puas
19	1	3	5	2	3	2.80	Puas
20	1	1	4	2	1	1.80	Tidak Puas
21	1	1	2	1	2	1.40	Tidak Puas
22	1	2	5	2	3	2.60	Puas
23	1	2	3	3	3	2.40	Tidak Puas
24	1	2	5	3	3	2.80	Puas
25	3	2	5	3	3	3.20	Puas
26	2	3	5	2	1	2.60	Puas
27	2	3	2	1	3	2.20	Tidak Puas
28	2	2	5	3	3	3.00	Puas
29	3	1	5	3	3	3.00	Puas
30	2	1	5	3	2	2.60	Puas
31	2	2	4	2	3	2.60	Puas
32	1	3	5	1	3	2.60	Puas
33	2	3	4	2	2	2.60	Puas
34	1	3	5	2	3	2.80	Puas
35	1	2	4	3	1	2.20	Tidak Puas
36	1	1	5	1	2	2.00	Tidak Puas
37	1	3	4	1	1	2.00	Tidak Puas
38	2	2	5	3	3	3.00	Puas
39	2	2	5	2	2	2.60	Puas
40	3	3	5	2	2	3.00	Puas
41	1	2	5	1	1	2.00	Tidak Puas
42	3	3	4	1	3	2.80	Puas
43	3	1	5	3	3	3.00	Puas
44	2	3	5	3	3	3.20	Puas

45	3	2	5	3	1	2.80	Puas
46	2	3	5	2	3	3.00	Puas
47	3	3	4	2	2	2.80	Puas
48	2	1	5	2	1	2.20	Tidak Puas
49	3	3	5	3	1	3.00	Puas
50	2	2	5	2	2	2.60	Puas

Data Testing

Data testing merupakan kumpulan data yang digunakan untuk menguji kemampuan model klasifikasi yang telah dibangun menggunakan data training. Pada penelitian ini, jumlah data testing sebanyak 27 data.

Tabel 5. Data Testing

No	Kemudahan Penggunaan Sistem	Kecepatan Akses Sistem	Keakuratan Informasi	Tampilan Antarmuka (User Interface)	Keandalan Sistem	Average	Kepuasan
51	3	3	4	1	2	2.60	Puas
52	2	1	4	3	3	2.60	Puas
53	3	1	5	1	1	2.20	Tidak Puas
54	2	3	4	2	3	2.80	Puas
55	3	2	5	2	1	2.60	Puas
56	2	3	5	2	2	2.80	Puas
57	1	2	5	2	2	2.40	Tidak Puas
58	2	3	5	2	1	2.60	Puas
59	1	3	4	3	2	2.60	Puas
60	2	2	5	2	3	2.80	Puas
61	2	1	5	2	2	2.40	Tidak Puas
62	3	3	5	1	1	2.60	Puas
63	2	1	4	1	1	1.80	Tidak Puas
64	1	2	5	2	1	2.20	Tidak Puas
65	1	1	4	1	2	1.80	Tidak Puas
66	2	2	5	1	2	2.40	Tidak Puas
67	2	3	4	3	3	3.00	Puas
68	2	2	4	2	3	2.60	Puas
69	2	1	4	2	1	2.00	Tidak Puas
70	1	1	4	3	1	2.00	Tidak Puas
71	1	3	5	3	1	2.60	Puas
72	1	2	5	1	3	2.40	Tidak Puas
73	1	2	5	1	3	2.40	Tidak Puas
74	2	2	5	2	2	2.60	Puas
75	2	3	5	3	1	2.80	Puas
76	1	3	4	3	2	2.60	Puas
77	1	3	3	3	1	2.20	Tidak Puas

Transformation

Tahap Transformation merupakan proses mengubah data yang telah melalui tahap *preprocessing* ke dalam bentuk yang sesuai untuk proses klasifikasi. Pada penelitian ini, data hasil kuesioner yang masih berupa nilai numerik dikonversikan ke dalam kategori berdasarkan masing-masing variabel X, sehingga lebih mudah diproses menggunakan algoritma Naïve Bayes dan Decision Tree. Proses konversi dilakukan secara konsisten pada seluruh data training dan data testing sesuai dengan aturan konversi yang telah ditentukan.

1. Data Training Konversi

Data training konversi merupakan hasil perubahan nilai numerik pada data training ke dalam kategori berdasarkan variabel X1 (Kemudahan Penggunaan Sistem), X2 (Kecepatan Akses Sistem), X3 (Keakuratan Informasi), X4 (Tampilan Antarmuka), dan X5 (Keandalan Sistem).

Tabel 6. Data Training Konversi

No	Kemudahan Penggunaan Sistem	Kecepatan Akses Sistem	Keakuratan Informasi	Tampilan Antarmuka (User Interface)	Keandalan Sistem	Kepuasan
1	Sangat Baik	Sangat Baik	Sangat Baik	Sangat Baik	Sangat Baik	Puas
2	Sangat Baik	Sangat Baik	Sangat Baik	Sangat Baik	Sangat Baik	Puas
3	Sangat Baik	Sangat Baik	Sangat Baik	Sangat Baik	Sangat Baik	Puas
4	Cukup	Sangat Baik	Sangat Baik	Cukup	Cukup	Puas
5	Baik	Baik	Sangat Baik	Cukup	Baik	Puas
6	Kurang	Kurang	Baik	Baik	Baik	Puas
7	Sangat Baik	Sangat Baik	Sangat Baik	Sangat Baik	Sangat Baik	Puas
8	Kurang	Cukup	Baik	Cukup	Baik	Puas
9	Kurang	Kurang	Baik	Kurang	Sangat Kurang	Tidak Puas
10	Kurang	Cukup	Sangat Baik	Cukup	Kurang	Puas
11	Kurang	Kurang	Baik	Kurang	Kurang	Tidak Puas
12	Sangat Kurang	Sangat Kurang	Cukup	Sangat Kurang	Sangat Kurang	Tidak Puas
13	Cukup	Cukup	Sangat Baik	Kurang	Kurang	Puas
14	Kurang	Cukup	Baik	Cukup	Sangat Kurang	Puas
15	Kurang	Sangat Kurang	Sangat Baik	Kurang	Cukup	Puas
16	Cukup	Kurang	Sangat Baik	Sangat Kurang	Cukup	Puas
17	Kurang	Sangat Kurang	Sangat Baik	Kurang	Sangat Kurang	Tidak Puas
18	Sangat Kurang	Kurang	Sangat Baik	Kurang	Sangat Kurang	Tidak Puas
19	Sangat Kurang	Cukup	Sangat Baik	Kurang	Cukup	Puas
20	Sangat Kurang	Sangat Kurang	Baik	Kurang	Sangat Kurang	Tidak Puas
21	Sangat Kurang	Sangat Kurang	Kurang	Sangat Kurang	Kurang	Tidak Puas
22	Sangat Kurang	Kurang	Sangat Baik	Kurang	Cukup	Puas
23	Sangat Kurang	Kurang	Cukup	Cukup	Cukup	Tidak Puas
24	Sangat Kurang	Kurang	Sangat Baik	Cukup	Cukup	Puas
25	Cukup	Kurang	Sangat Baik	Cukup	Cukup	Puas
26	Kurang	Cukup	Sangat Baik	Kurang	Sangat Kurang	Puas
27	Kurang	Cukup	Kurang	Sangat Kurang	Cukup	Tidak Puas
28	Kurang	Kurang	Sangat Baik	Cukup	Cukup	Puas
29	Cukup	Sangat Kurang	Sangat Baik	Cukup	Cukup	Puas
30	Kurang	Sangat Kurang	Sangat Baik	Cukup	Kurang	Puas
31	Kurang	Kurang	Baik	Kurang	Cukup	Puas
32	Sangat Kurang	Cukup	Sangat Baik	Sangat Kurang	Cukup	Puas
33	Kurang	Cukup	Baik	Kurang	Kurang	Puas
34	Sangat Kurang	Cukup	Sangat Baik	Kurang	Cukup	Puas
35	Sangat Kurang	Kurang	Baik	Cukup	Sangat Kurang	Tidak Puas
36	Sangat Kurang	Sangat Kurang	Sangat Baik	Sangat Kurang	Kurang	Tidak Puas
37	Sangat Kurang	Cukup	Baik	Sangat Kurang	Sangat Kurang	Tidak Puas
38	Kurang	Kurang	Sangat Baik	Cukup	Cukup	Puas
39	Kurang	Kurang	Sangat Baik	Kurang	Kurang	Puas
40	Cukup	Cukup	Sangat Baik	Kurang	Kurang	Puas
41	Sangat Kurang	Kurang	Sangat Baik	Sangat Kurang	Sangat Kurang	Tidak Puas

No	Kemudahan Penggunaan Sistem	Kecepatan Akses Sistem	Keakuratan Informasi	Tampilan Antarmuka (User Interface)	Keandalan Sistem	Kepuasan
42	Cukup	Cukup	Baik	Sangat Kurang	Cukup	Puas
43	Cukup	Sangat Kurang	Sangat Baik	Cukup	Cukup	Puas
44	Kurang	Cukup	Sangat Baik	Cukup	Cukup	Puas
45	Cukup	Kurang	Sangat Baik	Cukup	Sangat Kurang	Puas
46	Kurang	Cukup	Sangat Baik	Kurang	Cukup	Puas
47	Cukup	Cukup	Baik	Kurang	Kurang	Puas
48	Kurang	Sangat Kurang	Sangat Baik	Kurang	Sangat Kurang	Tidak Puas
49	Cukup	Cukup	Sangat Baik	Cukup	Sangat Kurang	Puas
50	Kurang	Kurang	Sangat Baik	Kurang	Kurang	Puas

2. Data Testing Konversi

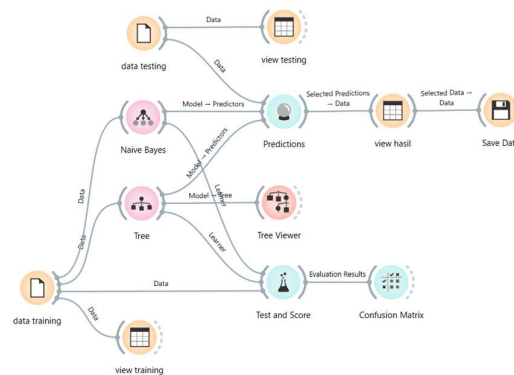
Data testing konversi merupakan hasil perubahan nilai numerik pada data testing ke dalam kategori berdasarkan variabel X1, X2, X3, X4, dan X5 menggunakan aturan konversi yang sama dengan data training.

Tabel 7. Data Testing Konversi

No	Kemudahan Penggunaan Sistem	Kecepatan Akses Sistem	Keakuratan Informasi	Tampilan Antarmuka (User Interface)	Keandalan Sistem	Kepuasan
51	Cukup	Cukup	Baik	Sangat Kurang	Kurang	Puas
52	Kurang	Sangat Kurang	Baik	Cukup	Cukup	Puas
53	Cukup	Sangat Kurang	Sangat Baik	Sangat Kurang	Sangat Kurang	Tidak Puas
54	Kurang	Cukup	Baik	Kurang	Cukup	Puas
55	Cukup	Kurang	Sangat Baik	Kurang	Sangat Kurang	Puas
56	Kurang	Cukup	Sangat Baik	Kurang	Kurang	Puas
57	Sangat Kurang	Kurang	Sangat Baik	Kurang	Kurang	Tidak Puas
58	Kurang	Cukup	Sangat Baik	Kurang	Sangat Kurang	Puas
59	Sangat Kurang	Cukup	Baik	Cukup	Kurang	Puas
60	Kurang	Kurang	Sangat Baik	Kurang	Cukup	Puas
61	Kurang	Sangat Kurang	Sangat Baik	Kurang	Kurang	Tidak Puas
62	Cukup	Cukup	Sangat Baik	Sangat Kurang	Sangat Kurang	Puas
63	Kurang	Sangat Kurang	Baik	Sangat Kurang	Sangat Kurang	Tidak Puas
64	Sangat Kurang	Kurang	Sangat Baik	Kurang	Sangat Kurang	Tidak Puas
65	Sangat Kurang	Sangat Kurang	Baik	Sangat Kurang	Kurang	Tidak Puas
66	Kurang	Kurang	Sangat Baik	Sangat Kurang	Kurang	Tidak Puas
67	Kurang	Cukup	Baik	Cukup	Cukup	Puas
68	Kurang	Kurang	Baik	Kurang	Cukup	Puas
69	Kurang	Sangat Kurang	Baik	Kurang	Sangat Kurang	Tidak Puas
70	Sangat Kurang	Sangat Kurang	Baik	Cukup	Sangat Kurang	Tidak Puas
71	Sangat Kurang	Cukup	Sangat Baik	Cukup	Sangat Kurang	Puas
72	Sangat Kurang	Kurang	Sangat Baik	Sangat Kurang	Cukup	Tidak Puas
73	Sangat Kurang	Kurang	Sangat Baik	Sangat Kurang	Cukup	Tidak Puas
74	Kurang	Kurang	Sangat Baik	Kurang	Kurang	Puas
75	Kurang	Cukup	Sangat Baik	Cukup	Sangat Kurang	Puas
76	Sangat Kurang	Cukup	Baik	Cukup	Kurang	Puas
77	Sangat Kurang	Cukup	Cukup	Cukup	Sangat Kurang	Tidak Puas

Data Mining

Tahap Data Mining merupakan tahap inti dalam penelitian yang bertujuan untuk membangun model klasifikasi berdasarkan data yang telah melalui proses *selection*, *preprocessing*, dan *transformation*. Pada penelitian ini, proses data mining dilakukan menggunakan aplikasi Orange Data Mining karena menyediakan berbagai algoritma klasifikasi serta fitur visualisasi yang memudahkan dalam proses analisis data. Melalui aplikasi Orange, data training digunakan untuk membentuk model klasifikasi menggunakan algoritma Naïve Bayes dan Decision Tree, sedangkan data testing digunakan untuk menguji kemampuan model dalam mengklasifikasikan tingkat kepuasan mahasiswa terhadap Aplikasi Sistem Informasi Terpadu (SITU). Alur proses pembentukan model data mining pada aplikasi Orange dapat dilihat pada gambar dibawah ini.



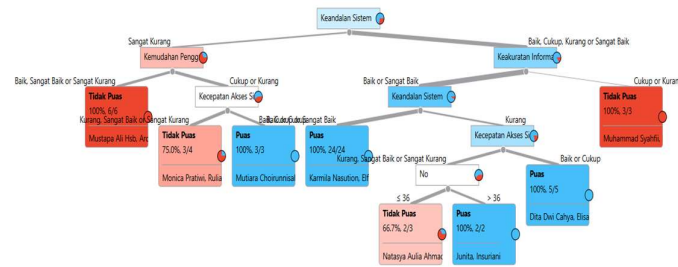
Gambar 2. Model Data Mining

Berdasarkan gambar diatas, proses data mining dimulai dengan memasukkan data training yang telah dipersiapkan pada tahap sebelumnya. Data training kemudian dihubungkan dengan widget Naïve Bayes dan Tree (Decision Tree) untuk membentuk model klasifikasi berdasarkan atribut yang dimiliki. Data training juga ditampilkan melalui widget View Training untuk memastikan data yang digunakan telah sesuai.

Selanjutnya, data testing dimasukkan ke dalam sistem dan ditampilkan melalui widget View Testing sebagai data yang akan digunakan dalam proses pengujian model. Model yang dihasilkan dari algoritma Naïve Bayes dan Decision Tree kemudian dihubungkan ke widget Predictions untuk menghasilkan prediksi tingkat kepuasan mahasiswa berdasarkan data testing. Hasil prediksi tersebut ditampilkan pada widget View Hasil sehingga dapat dianalisis, kemudian disimpan melalui widget Save Data untuk digunakan pada tahap analisis berikutnya.

Selain menghasilkan prediksi, model Decision Tree juga dihubungkan dengan widget Tree Viewer untuk menampilkan struktur pohon keputusan yang terbentuk. Visualisasi ini memudahkan dalam memahami aturan-aturan klasifikasi yang dihasilkan oleh algoritma Decision Tree berdasarkan atribut-atribut penelitian.

Pada tahap evaluasi, data training bersama model klasifikasi dihubungkan ke widget Test and Score untuk menghitung performa masing-masing algoritma. Hasil evaluasi tersebut kemudian diteruskan ke widget Confusion Matrix untuk mengetahui jumlah prediksi yang benar maupun salah, sehingga dapat dihitung nilai Accuracy, Precision, Recall, dan F1-Score. Nilai-nilai tersebut selanjutnya digunakan sebagai dasar dalam menentukan algoritma klasifikasi yang memiliki performa terbaik pada penelitian ini.



Gambar 3. Model Pohon Keputusan

Berdasarkan diatas, terlihat bahwa algoritma Decision Tree membentuk sebuah pohon keputusan untuk mengklasifikasikan tingkat kepuasan mahasiswa terhadap penggunaan Aplikasi Sistem Informasi Terpadu (SITU) ke dalam kategori Puas dan Tidak Puas. Atribut yang menjadi akar (root) pada pohon keputusan adalah Keandalan Sistem (X5), yang menunjukkan bahwa variabel tersebut merupakan atribut dengan pengaruh terbesar dalam proses klasifikasi. Selanjutnya, proses pengambilan keputusan dilakukan secara bertahap melalui atribut Kemudahan Penggunaan Sistem (X1), Keakuratan Informasi (X3), dan Kecepatan Akses Sistem (X2) sesuai dengan nilai masing-masing atribut. Setiap percabangan menghasilkan aturan (*rule*) yang mengarahkan data ke salah satu kelas, yaitu Puas atau Tidak Puas. Pada setiap simpul akhir (*leaf node*), ditampilkan hasil klasifikasi beserta jumlah data yang termasuk pada kelas tersebut, misalnya simpul dengan nilai Puas 100%, 24/24 menunjukkan bahwa seluruh 24 data pada simpul tersebut berhasil diklasifikasikan ke dalam kategori Puas, sedangkan simpul Tidak Puas 100%, 3/3 menunjukkan bahwa seluruh data pada simpul tersebut termasuk kategori Tidak Puas. Dengan demikian, pohon keputusan yang terbentuk mampu memberikan aturan klasifikasi yang mudah dipahami dan menunjukkan bahwa Keandalan Sistem, diikuti oleh Kemudahan Penggunaan Sistem, Keakuratan Informasi, dan Kecepatan Akses Sistem, merupakan variabel yang berperan dalam menentukan tingkat kepuasan mahasiswa terhadap penggunaan Aplikasi Sistem Informasi Terpadu (SITU).

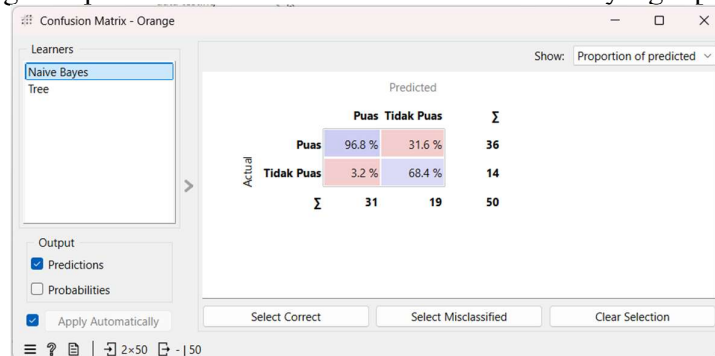
Evaluation

Tahap Evaluation bertujuan untuk mengetahui tingkat kinerja model klasifikasi yang telah dibangun menggunakan algoritma Naïve Bayes dan Decision Tree. Pada tahap ini dilakukan proses evaluasi menggunakan Confusion Matrix pada aplikasi Orange Data Mining. Confusion Matrix digunakan untuk membandingkan hasil prediksi model dengan kelas sebenarnya sehingga dapat diketahui kemampuan masing-masing algoritma dalam mengklasifikasikan tingkat kepuasan mahasiswa ke dalam kategori Puas dan Tidak Puas. Hasil evaluasi ini menjadi dasar dalam menentukan algoritma yang memiliki performa terbaik berdasarkan tingkat ketepatan klasifikasi yang dihasilkan.

1. Evaluasi Menggunakan Confusion Matrix Metode Naïve Bayes

Berdasarkan hasil Confusion Matrix pada metode Naïve Bayes, diperoleh bahwa terdapat 31 data yang diprediksi sebagai kategori Puas dan 19 data yang diprediksi sebagai kategori Tidak Puas, dengan jumlah keseluruhan 50 data. Dari data yang diprediksi Puas, sebesar 96,8% merupakan data yang benar-benar berada pada kategori Puas, sedangkan 3,2% merupakan data yang sebenarnya termasuk kategori Tidak Puas. Sementara itu, dari data yang diprediksi Tidak Puas, sebesar 31,6% merupakan data yang sebenarnya berada pada kategori Puas, sedangkan 68,4% merupakan data yang benar-benar berada pada kategori Tidak Puas. Hasil tersebut menunjukkan bahwa metode Naïve

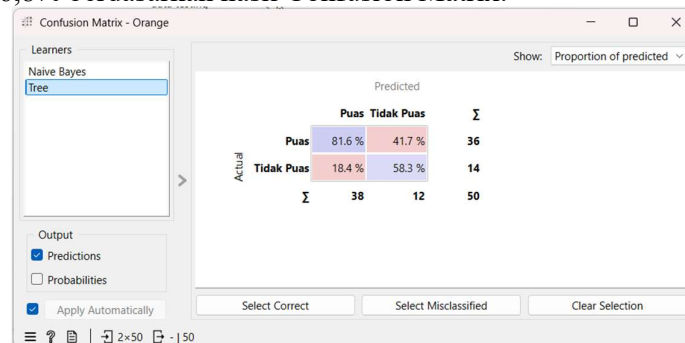
Bayes memiliki kemampuan yang sangat baik dalam mengklasifikasikan mahasiswa pada kategori Puas, serta mampu mengidentifikasi kategori Tidak Puas dengan tingkat ketepatan yang cukup baik berdasarkan hasil Confusion Matrix yang diperoleh.



Gambar 4. Confusion Matrix Metode Naïve Bayes

2. Evaluasi Menggunakan Confusion Matrix Metode Decision Tree

Berdasarkan hasil Confusion Matrix pada metode Decision Tree, diperoleh bahwa terdapat 38 data yang diprediksi sebagai kategori Puas dan 12 data yang diprediksi sebagai kategori Tidak Puas, dengan jumlah keseluruhan 50 data. Dari data yang diprediksi Puas, sebesar 81,6% merupakan data yang benar-benar berada pada kategori Puas, sedangkan 18,4% merupakan data yang sebenarnya termasuk kategori Tidak Puas. Selanjutnya, dari data yang diprediksi Tidak Puas, sebesar 41,7% merupakan data yang sebenarnya berada pada kategori Puas, sedangkan 58,3% merupakan data yang benar-benar berada pada kategori Tidak Puas. Hasil tersebut menunjukkan bahwa metode Decision Tree mampu mengklasifikasikan tingkat kepuasan mahasiswa dengan cukup baik, namun tingkat ketepatan prediksi pada kategori Puas masih lebih rendah dibandingkan metode Naïve Bayes, yang memperoleh tingkat ketepatan prediksi kategori Puas sebesar 96,8% berdasarkan hasil Confusion Matrix.



Gambar 5. Confusion Matrix Metode Decision Tree

Interpretation

Tahap Interpretation merupakan tahap akhir dalam proses *Knowledge Discovery in Databases* (KDD) yang bertujuan untuk menginterpretasikan hasil klasifikasi yang telah diperoleh dari penerapan algoritma Naïve Bayes dan Decision Tree. Pada tahap ini dilakukan analisis terhadap hasil prediksi kedua algoritma menggunakan data testing untuk mengetahui kemampuan masing-masing metode dalam mengklasifikasikan tingkat kepuasan mahasiswa terhadap Aplikasi Sistem Informasi Terpadu (SITU) di Universitas Labuhanbatu. Selain itu, tahap ini juga bertujuan untuk mengidentifikasi kesesuaian hasil prediksi dengan karakteristik setiap variabel penelitian, yaitu Kemudahan Penggunaan

Sistem (X1), Kecepatan Akses Sistem (X2), Keakuratan Informasi (X3), Tampilan Antarmuka (X4), dan Keandalan Sistem (X5). Hasil prediksi kedua algoritma terhadap data testing disajikan pada tabel dibawah ini.

Tabel 8. Hasil Analisis

No	Nama Mahasiswa	Tree	Naive Bayes
51	Julianti	Puas	Puas
52	Deny Alfariy	Puas	Puas
53	Amalia Rizky	Tidak Puas	Tidak Puas
54	Suci Herawati	Puas	Puas
55	Doni Adrian	Tidak Puas	Puas
56	Feni Agniasi	Puas	Puas
57	Risma Ramadhani	Puas	Puas
58	Lidya Annisa	Puas	Puas
59	Novita Nurbaiti	Puas	Puas
60	Saydina Umar	Puas	Puas
61	Renita Hutahuruk	Puas	Puas
62	Nurma Wanda	Puas	Puas
63	Aisyah Harahap	Tidak Puas	Tidak Puas
64	Meiliza Sirait	Tidak Puas	Tidak Puas
65	Sinta Rahmadika	Puas	Tidak Puas
66	Rindiani	Puas	Puas
67	Suci Ramadani	Puas	Puas
68	Mutiara Sapriani	Puas	Puas
69	Agustin Ningsih	Tidak Puas	Tidak Puas
70	Siti Aisyah	Tidak Puas	Tidak Puas
71	Zariah Sakinah	Tidak Puas	Puas
72	Dede Apriani	Puas	Tidak Puas
73	Anggi Virlanda	Puas	Tidak Puas
74	Sri Mualima	Puas	Puas
75	Martauli Pasaribu	Puas	Puas
76	Elyn Muliani	Puas	Puas
77	Mariani	Tidak Puas	Tidak Puas

Berdasarkan hasil prediksi pada tabel diatas, dapat diketahui bahwa algoritma Decision Tree memberikan prediksi kategori Puas pada sebagian besar data testing. Dari 27 data mahasiswa yang diuji, sebanyak 19 data diklasifikasikan ke dalam kategori Puas, sedangkan 8 data lainnya diklasifikasikan ke dalam kategori Tidak Puas, yaitu Amalia Rizky, Doni Adrian, Aisyah Harahap, Meiliza Sirait, Agustin Ningsih, Siti Aisyah, Zariah Sakinah, dan Mariani. Sementara itu, algoritma Naïve Bayes menghasilkan 18 data yang diklasifikasikan ke dalam kategori Puas dan 9 data yang diklasifikasikan ke dalam kategori Tidak Puas, yaitu Amalia Rizky, Aisyah Harahap, Meiliza Sirait, Sinta Rahmadika, Agustin Ningsih, Siti Aisyah, Dede Apriani, Anggi Virlanda, dan Mariani. Hasil tersebut menunjukkan bahwa kedua algoritma memiliki kemampuan yang baik dalam mengklasifikasikan tingkat kepuasan mahasiswa, namun terdapat beberapa perbedaan hasil prediksi pada beberapa data, seperti Doni Adrian dan Zariah Sakinah yang diklasifikasikan sebagai Tidak Puas oleh Decision Tree tetapi Puas oleh Naïve Bayes, serta Sinta Rahmadika, Dede Apriani, dan Anggi Virlanda yang diklasifikasikan sebagai Puas oleh Decision Tree namun Tidak Puas oleh Naïve Bayes. Perbedaan tersebut menunjukkan bahwa masing-masing algoritma memiliki mekanisme klasifikasi yang berbeda dalam menentukan tingkat kepuasan mahasiswa berdasarkan atribut yang digunakan dalam penelitian.

5. Kesimpulan

Berdasarkan hasil penelitian yang telah dilakukan, dapat disimpulkan bahwa penerapan metode Naïve Bayes dan Decision Tree mampu digunakan untuk menganalisis tingkat kepuasan mahasiswa terhadap penggunaan Aplikasi Sistem Informasi Terpadu (SITU) di Universitas Labuhanbatu berdasarkan lima variabel, yaitu Kemudahan

Penggunaan Sistem, Kecepatan Akses Sistem, Keakuratan Informasi, Tampilan Antarmuka (User Interface), dan Keandalan Sistem. Proses analisis dilakukan melalui tahapan Knowledge Discovery in Databases (KDD) yang meliputi Selection, Preprocessing, Transformation, Data Mining, Evaluation, dan Interpretation, sehingga menghasilkan model klasifikasi yang mampu mengelompokkan mahasiswa ke dalam kategori Puas dan Tidak Puas. Berdasarkan hasil interpretasi terhadap 27 data testing, algoritma Decision Tree menghasilkan 19 data pada kategori Puas dan 8 data pada kategori Tidak Puas, sedangkan algoritma Naïve Bayes menghasilkan 18 data pada kategori Puas dan 9 data pada kategori Tidak Puas. Selanjutnya, hasil evaluasi menggunakan Confusion Matrix terhadap 50 data menunjukkan bahwa metode Naïve Bayes menghasilkan 31 data yang diprediksi sebagai kategori Puas dan 19 data sebagai kategori Tidak Puas, dengan tingkat ketepatan prediksi kategori Puas sebesar 96,8%, sedangkan metode Decision Tree menghasilkan 38 data yang diprediksi sebagai kategori Puas dan 12 data sebagai kategori Tidak Puas, dengan tingkat ketepatan prediksi kategori Puas sebesar 81,6%. Berdasarkan hasil tersebut, metode Naïve Bayes menunjukkan performa klasifikasi yang lebih baik dibandingkan Decision Tree dalam menganalisis tingkat kepuasan mahasiswa terhadap penggunaan Aplikasi Sistem Informasi Terpadu (SITU). Oleh karena itu, metode Naïve Bayes dapat dijadikan sebagai model klasifikasi yang lebih efektif dalam penelitian ini serta menjadi dasar bagi Universitas Labuhanbatu dalam melakukan evaluasi dan peningkatan kualitas layanan SITU, terutama pada aspek kemudahan penggunaan sistem, kecepatan akses sistem, keakuratan informasi, tampilan antarmuka, dan keandalan sistem agar tingkat kepuasan mahasiswa dapat terus ditingkatkan.

6. Daftar Pustaka

- Adzy, L. B., Asriyanik, A., & Pambudi, A. (2023). Algoritma naïve bayes untuk klasifikasi kelayakan penerima bantuan iuran jaminan kesehatan pemerintah daerah kabupaten sukabumi. *Jurnal Mnemonic*, 6(1), 1-10.
- Aktavera, B., & Wijaya, H. O. L. (2023). Klasifikasi Produk Menggunakan Algoritma Decision Tree. *Jurnal Teknologi Informasi Mura*, 15(1), 24-29.
- Ameliatus, Q., & Fatah, Z. (2024). Klasifikasi Kelulusan Mahasiswa Menggunakan Metode Decision Tree Menggunakan Aplikasi Rapidminer. *Jurnal Riset Teknik Komputer*, 1(4), 58-64.
- Antika, D., Harahap, S. Z., Ah, R. M., & Juledi, A. P. (2025). Penerapan Data mining Klasifikasi Tingkat Kepuasan Mahasiswa Terhadap Pelayanan Akademik Menggunakan Metode Naïve Bayes Dan Support Vector Machine (Studi Kasus Program Studi Sistem Informasi Universitas Labuhanbatu). *Journal of Computer Science and Information System (JCoInS)*, 6(3), 221-232
- Assyahri, W., & Mardaus, M. (2023). Pengaruh kualitas pelayanan akademik terhadap kepuasan mahasiswa Universitas Negeri Padang. *Jurnal Manajemen Dan Ilmu Administrasi Publik (JMIAP)*, 5(3), 239-247.
- Hasanah, W., & Fatah, Z. (2024). IMPLEMENTASI KELULUSAN MAHASISWA BERDASARKAN DATA NILAI AKADEMIK MENGGUNAKAN ALGORITMA DECISION TREE. *Jurnal Ilmiah Multidisiplin Ilmu*, 1(6), 79-83.
- Hirwono, B., Hermawan, A., & Avianto, D. (2023). Implementasi Metode Naïve Bayes untuk Klasifikasi Penderita Penyakit Jantung. *J. JTIK (Jurnal Teknol. Inf. dan Komunikasi)*, 7(3), 450-457.

- Huriah, D. A., & Nuris, N. D. (2023). Klasifikasi penerima bantuan sosial umkm menggunakan algoritma naïve bayes. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 7(1), 360-365.
- Kartini, A., & Sanmorino, A. (2024). Analisis tingkat kepuasan mahasiswa terhadap sistem informasi akademik STEBIS IGM menggunakan metode PIECES Framework. *AnoaTIK: Jurnal Teknologi Informasi Dan Komputer*, 2(1), 51-59.
- Kirana, A. N., Nurhakim, B., Permana, S. E., Prihartono, W., & Dwilestari, G. (2024). Implementasi Algoritma Naive Bayes Untuk Memprediksi Cuaca Menggunakan Rapidminer. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 8(2), 1637-1642.
- Kurniawan, H. P., & Farhatuaini, L. (2024). Identifikasi Pola Kepuasan Mahasiswa Terhadap Proses Pembelajaran Menggunakan Algoritma K-Means Clustering. *Jurnal Informatika: Jurnal Pengembangan IT*, 9(2), 164-172.
- Mardiansa, M., Sari, H. L., & Prahasti, P. (2023). Penerapan Data Mining Untuk Mengetahui Minat Siswa Pada Pelajaran IPA Menggunakan Metode K-Means Clustering. *Jurnal Multidisiplin Dehasen (MUDE)*, 2(4), 693-702.
- Maryam, N. S. (2024). Pelayanan prima dalam meningkatkan kepuasan mahasiswa terhadap layanan pembelajaran di Universitas Mandiri. *Jurnal Ilmiah Manajemen dan Kewirausahaan (JUMANAGE)*, 3(1), 257-265.
- Nopiyanti, N., Sriati, S., & Imania, K. (2024). Pengaruh kualitas pelayanan akademik terhadap kepuasan mahasiswa di Fakultas Keguruan dan Ilmu Pendidikan Universitas Sriwijaya. *Jurnal Ilmu Administrasi Negara AsIAN (Asosiasi Ilmuwan Administrasi Negara)*, 12(2), 247-253.
- Nur, H., Nurhaeni, N., Ahmad, H., & Muhammad, R. A. P. (2025). ANALISIS PENGARUH METODE PEMBELAJARAN HYBRID TERHADAP TINGKAT KEPUASAN MAHASISWA MENGGUNAKAN ALGORITMA DECISION TREE.
- Nurani, S., Syahra, Y., & Calam, A. (2023). Penerapan Data Mining Dalam Clustering Pencapaian Target Penjualan Menggunakan Algoritma K-Means. *Jurnal Sistem Informasi Triguna Dharma (JURSI TGD)*, 2(3), 355-363.
- Olivia, N., Rizky, F., Angelia, D., & Natasya, D. (2023). Penerapan Metode Naïve Bayes Dalam Klasifikasi Warga Penerima Bantuan Peralatan Tukang Kayu Di Desa Muara Sungai. *Innovative: Journal Of Social Science Research*, 3(2), 9574-9582.
- Qamal, M., Syah, F., & Parapat, A. Z. I. (2023). Implentasi Data Mining Untuk Rekomendasi Paket Menu Makanan Dengan Menggunakan Algoritma Apriori. *TECHSI-Jurnal Teknik Informatika*, 14(1), 42-53.
- Rahmadini, R., LorencisLubis, E. E., Priansyah, A., RWN, Y., & Meutia, T. (2023). Penerapan data mining untuk memprediksi harga bahan pangan di Indonesia menggunakan algoritma K-Nearest Neighbor. *Jurnal Mahasiswa Akuntansi Samudra*, 4(4), 223-235.
- Simbolon, A. I. A., & Riza, F. (2025). Analisis Kepuasan Pengguna Aplikasi Umsu Academy Menggunakan Metode Naïve Bayes. *Portal Riset dan Inovasi Sistem Perangkat Lunak*, 3(3), 198-211.
- Siregar, H. A., Azlan, A., & Gaol, N. Y. L. (2023). Penerapan Data Mining Pada Penjualan Rumah Makan Kasih Ibu Menggunakan Metode K-Means Clustering. *Jurnal Sistem Informasi Triguna Dharma (JURSI TGD)*, 2(5), 750-760.
- Sukriadi, S., Ismail, I., & Andzar, A. M. (2023). Penerapan text mining dalam klasifikasi judul skripsi yang diusulkan mahasiswa menggunakan metode naïve

- Bayes. *Jurnal Ilmiah Sistem Informasi Dan Teknik Informatika (JISTI)*, 6(2), 184-196.
- Wirasena, M. R., Reinaldi, M. R., & Jambak, M. I. (2024). Algoritma Decision Tree ID3 Bagi Lembaga Pemberi Pinjaman Untuk Menentukan Faktor Yang Mempengaruhi Kelayakan Individu Memperoleh Pinjaman. *The Indonesian Journal of Computer Science*, 13(3).
- Yani, A., Azmi, Z., & Suherdi, D. (2023). Implementasi Data Mining Menganalisa Data Penjualan Menggunakan Algoritma K-Means Clustering. *Jurnal Sistem Informasi Triguna Dharma (JURSI TGD)*, 2(2), 315-323.
- Zain, A. N. W., Muafi, M., & Tholib, A. (2024). Klasifikasi Data Mining di Tingkat Kepuasan Mahasiswa Terhadap Pelayanan Sistem Informasi Fakultas Teknik Universitas Nurul Jadid. *SKANIKA: Sistem Komputer Dan Teknik Informatika*, 7(2), 204-213.
- Zaman, B., & Ardima, M. B. (2024). Prediksi Kepuasan Mahasiswa Terhadap Pelayanan Akademik Menggunakan Model Decision Tree. *Jurnal Transformatika*, 21(2), 46-55.