

**ANALISIS KEPUASAN PENGGUNA APLIKASI UMSU
ACADEMY MENGGUNAKAN METODE NAÏVE BAYES**

SKRIPSI

DISUSUN OLEH

ADE IRA AZZAHRA SIMBOLON

NPM. 2109010060



PROGRAM STUDI SISTEM INFORMASI

FAKULTAS ILMU KOMPUTER DAN TEKNOLOGI INFORMASI

UNIVERSITAS MUHAMMADIYAH SUMATERA UTARA

MEDAN

2025

**ANALISIS KEPUASAN PENGGUNA APLIKASI UMSU ACADEMY
MENGUNAKAN METODE NAÏVE BAYES**

SKRIPSI

**Diajukan sebagai salah satu syarat untuk memperoleh gelar Sarjana Komputer
(S.Kom) dalam Program Studi Sistem Informasi pada Fakultas Ilmu Komputer
dan Teknologi Informasi, Universitas Muhammadiyah Sumatera Utara**

ADE IRA AZZAHRA SIMBOLON

NPM. 2109010060

**PROGRAM STUDI SISTEM INFORMASI
FAKULTAS ILMU KOMPUTER DAN TEKNOLOGI INFORMASI
UNIVERSITAS MUHAMMADIYAH SUMATERA UTARA**

MEDAN

2025

LEMBAR PENGESAHAN

Judul Skripsi : ANALISIS KEPUASAN PENGGUNA APLIKASI
UMSU ACADEMY MENGGUNAKAN METODE
NAÏVE BAYES

Nama Mahasiswa : ADE IRA AZZAHRA SIMBOLON

NPM : 2109010060

Program Studi : SISTEM INFORMASI

Menyetujui
Komisi Pembimbing



(FERDY RIZA, S.T., M.Kom)
NIDN. 0103068901

Ketua Program Studi



(Dr. Firahmi Rizky, S.Kom., M.Kom)
NIDN. 0116079201

Dekan



(Dr. Al-Khowarizmi, S.Kom., M.Kom.)
NIDN. 0127099201

PERNYATAAN ORISINALITAS

ANALISIS KEPUASAN PENGGUNA APLIKASI UMSU ACADEMY MENGUNAKAN METODE NAÏVE BAYES

SKRIPSI

Saya menyatakan bahwa karya tulis ini adalah hasil karya sendiri, kecuali beberapa kutipan dan ringkasan yang masing-masing disebutkan sumbernya.

Medan, 12 juli 2025

Yang membuat pernyataan



Ade Ira Azzahra Simbolon

NPM. 2109010060

**PERNYATAAN PERSETUJUAN PUBLIKASI
KARYA ILMIAH UNTUK KEPENTINGAN
AKADEMIS**

Sebagai sivitas akademika Universitas Muhammadiyah Sumatera Utara, saya bertanda tangan dibawah ini:

Nama	: Ade Ira Azzahra Simbolon
NPM	: 2109010060
Program Studi	: Sistem Informasi
Karya Ilmiah	: Skripsi

Demi pengembangan ilmu pengetahuan, menyetujui untuk memberikan kepada Universitas Muhammadiyah Sumatera Utara Hak Bedas Royalti Non-Eksekutif (*Non-Exclusive Royalty free Right*) atas penelitian skripsi saya yang berjudul:

**ANALISIS KEPUASAN PENGGUNA APLIKASI UMSU ACADEMY
MENGUNAKAN METODE NAÏVE BAYES**

Beserta perangkat yang ada (jika diperlukan). Dengan Hak Bebas Royalti Non-Eksekutif ini, Universitas Muhammadiyah Sumatera Utara berhak menyimpan, mengalih media, memformat, mengelola dalam bentuk database, merawat dan mempublikasikan Skripsi saya ini tanpa meminta izin dari saya selama tetap mencantumkan nama saya sebagai penulis dan sebagai pemegang dan atau sebagai pemilik hak cipta.

Demikian pernyataan ini dibuat dengan sebenarnya.

Medan, 12 Juli 2025

Yang membuat pernyataan



Ade Ira Azzahra Simbolon
NPM. 2109010060

RIWAYAT HIDUP

DATA PRIBADI

Nama Lengkap	: Ade Ira Azzahra Simbolon
Tempat dan Tanggal Lahir	: Medan, 28 Januari 2003
Alamat Rumah	: Jl. Tuba II No.52, Medan Denai
Telepon/Faks/HP	: 089630825803
E-mail	: adeiraazzahra28@gmail.com
Instansi Tempat Kerja	: -
Alamat Kantor	: -

DATA PENDIDIKAN

SD	: SD SWASTA SABILINA	TAMAT: 2014
SMP	: MTs NEGERI 2 MEDAN	TAMAT: 2017
SMA	: MAS PLUS AL-ULUM MEDAN	TAMAT: 2020

KATA PENGANTAR



Alhamdulillah, puji syukur kepada Allah SWT atas segala rahmat dan hidayah nya, sehingga penulis dapat menyelesaikan skripsi berjudul “**(ANALISIS KEPUASAN PENGGUNA APLIKASI UMSU ACADEMY MENGGUNAKAN METODE NAÏVE BAYES)**”. Skripsi ini disusun untuk memenuhi syarat mendapatkan gelar sarjana dalam program studi Sistem Informasi Fakultas Ilmu Komputer dan Teknologi Universitas Muhammadiyah Sumatra Utara.

Dalam proses penelitian dan penyusunan skripsi ini, banyak pelajaran dan tantangan yang dihadapi, yang semuanya memberikan manfaat di masa depan. Semua pencapaian ini tidak lepas dari dukungan dan motivasi dari banyak pihak. Oleh karena itu, penulis menyampaikan terima kasih sebesar-besarnya dan penghargaan setinggi tingginya kepada:

1. Bapak Prof. Dr. Agussani, M.AP., Rektor Universitas Muhammadiyah Sumatera Utara (UMSU).
2. Bapak Dr. Al-Khowarizmi, S.Kom., M.Kom. Dekan Fakultas Ilmu Komputer dan Teknologi Informasi (FIKTI) UMSU.
3. Bapak Martiano S.Pd, S.Kom., M.Kom Ketua Program Studi Sistem Informasi.
4. Ibu Yoshida Sary, SE, S.Kom., M.Kom. Sekretaris Program Studi Sistem Informasi.

5. Pembimbing sekaligus mentor peneliti saya yaitu bapak Ferdy Riza, S.T., M.Kom. terimakasih sudah membimbing saya dan memberikan saya kemudahan untuk bimbingan dengan sangat baik dan penuh kesabaran, terimakasih juga atas ilmu yang bapak berikan kepada saya sehingga bisa sampai ke tahap ini.
6. Kedua orang tua, ibu dan ayah saya yang telah memberikan semangat dan motivasi serta dukungan dalam pengerjaan tugas akhir saya.
7. Kakak saya Nisya Nainita Simbolon dan adik saya Raihan Asri Simbolon yang telah memberikan saya semangat, nasihat dan memotivasi saya dalam pengerjaan tugas akhir.
8. Keluarga besar MARPAREBAN dan persepupuan opung bayo yang telah memberikan saya semangat, saran, motivasi dan membantu saya dalam pengerjaan tugas akhir.

ANALISIS KEPUASAN PENGGUNA APLIKASI UMSU ACADEMY MENGUNAKAN METODE NAÏVE BAYES

ABSTRAK

Transformasi digital dalam dunia pendidikan mendorong peningkatan kualitas sistem informasi akademik, khususnya melalui pemanfaatan aplikasi mobile. Universitas Muhammadiyah Sumatera Utara mengembangkan aplikasi *UMSU Academy* sebagai sarana pendukung administrasi akademik mahasiswa. Penelitian ini bertujuan untuk mengukur dan mengklasifikasikan tingkat kepuasan pengguna terhadap aplikasi tersebut menggunakan metode *Naïve Bayes*. Data diperoleh melalui penyebaran kuesioner yang terdiri dari 25 indikator, dikelompokkan dalam lima variabel utama yaitu content, format, accuracy, timeliness, dan ease of use. Rata-rata dari indikator tersebut digunakan untuk menentukan kelas kepuasan pengguna. Sistem klasifikasi dibangun berbasis web menggunakan bahasa pemrograman Python dan framework Flask, serta dilengkapi dengan fitur input data, prediksi otomatis, evaluasi model, dan validasi perhitungan manual. Evaluasi performa dilakukan menggunakan metrik akurasi, precision, recall, dan F1-score. Hasil pengujian menunjukkan bahwa algoritma *Naïve Bayes* berhasil mengklasifikasikan tingkat kepuasan pengguna dengan tingkat akurasi sebesar 95,98%, serta divalidasi melalui perhitungan probabilitas posterior secara manual. Temuan ini diharapkan dapat menjadi acuan evaluatif bagi pengembang dalam meningkatkan kualitas layanan akademik digital di lingkungan Universitas Muhammadiyah Sumatera Utara.

Kata kunci: UMSU Academy, Kepuasan Pengguna, Naïve Bayes, Klasifikasi, Sistem Informasi Akademik.

ANALISIS KEPUASAN PENGGUNA APLIKASI UMSU ACADEMY MENGUNAKAN METODE NAÏVE BAYES

ABSTRACT

Digital transformation in higher education has driven the improvement of academic information systems, particularly through the use of mobile applications. Universitas Muhammadiyah Sumatera Utara developed the *UMSU Academy* application as a tool to support students' academic administrative processes. This study aims to measure and classify user satisfaction with the application using the *Naïve Bayes* method. Data were collected through a questionnaire consisting of 25 indicators grouped into five main variables are content, format, accuracy, timeliness, and ease of use. The average score of these indicators was used to determine the satisfaction level. A web-based classification system was developed using the Python programming language and Flask framework, featuring data input, automatic prediction, model evaluation, and manual probability validation. The model's performance was evaluated using accuracy, precision, recall, and F1-score metrics. The results showed that the *Naïve Bayes* algorithm successfully classified user satisfaction levels with an accuracy of 95.98%, which was further validated through manual posterior probability calculations. These findings are expected to serve as an evaluative reference for application developers in improving the quality of academic digital services at Universitas Muhammadiyah Sumatera Utara.

Keywords: UMSU Academy, User Satisfaction, Naïve Bayes, Classification, Academic Information System.

DAFTAR ISI

ANALISIS KEPUASAN PENGGUNA APLIKASI UMSU ACADEMY MENGUNAKAN METODE NAÏVE BAYES.....	i
ANALISIS KEPUASAN PENGGUNA APLIKASI UMSU ACADEMY MENGUNAKAN METODE NAÏVE BAYES.....	ii
DAFTAR ISI.....	iii
DAFTAR TABEL.....	v
DAFTAR GAMBAR.....	vi
BAB I.....	1
PENDAHULUAN.....	1
1.1 Latar Belakang Masalah	1
1.2 Rumusan Masalah	4
1.3 Batasan Masalah.....	4
1.4 Tujuan Penelitian.....	5
1.5 Manfaat Penelitian.....	6
BAB II	7
LANDASAN TEORI.....	7
2.1 Aplikasi.....	7
2.2 UMSU <i>Academy</i>	10
2.3 Analisis Data	12
2.4 Kepuasan Pengguna.....	15
2.5 Data Mining.....	19
2.6 Naïve Bayes.....	31
2.7 Penelitian Terkait.....	35
BAB III.....	37
METODOLOGI PENELITIAN	37
3.1 Obyek Penelitian	37
3.1.1 Lokasi Penelitian.....	37
3.1.2 Atribut atau Variabel Penelitian.....	37
3.1.3 Kelas (Label).....	37
3.2 Alat dan Bahan	38
3.2.1 Kuesioner	38
3.2.2 Microsoft Excel.....	38
3.2.3 Kebutuhan Perangkat Keras.....	38
3.2.4 Literatur dan Referensi.....	39
3.3 Tahap Desain.....	39
3.3.1 Use Case Diagram	39
3.3.2 Activity Diagram	40

3.3.3	Desain Interface	40
3.4	Algoritma Naïve Bayes	41
3.4.1	Rumus Pencarian Tingkat Akurasi	43
3.5	Jadwal Penelitian	44
BAB IV		45
HASIL DAN PEMBAHASAN		45
4.1	Hasil Dan Pembahasan	45
4.2	Implementasi Sistem Naïve Bayes	46
4.2.1	Tampilan Dashboard	46
4.2.2	Tampilan Menu Dataset	47
4.2.3	Tampilan Menu Initial Process	48
4.2.4	Tampilan Menu Performance.....	49
4.2.5	Tampilan Menu Predict.....	55
4.3	Perhitungan Naïve Bayes	57
4.3.1	Perhitungan Probabilitas Class	61
4.3.2	Perhitungan Probabilitas Kategori	62
4.3.3	Perhitungan Manual Naïve Bayes.....	70
4.3.4	Menghitung Prior Probabilitas dan <i>Likelihood</i>	71
4.3.5	Menghitung Posterior Dan Menentukan Hasil Prediksi	77
4.3.6	Perhitungan Precision, Recall, F1-Score Dan Akurasi	80
BAB V		83
PENUTUP		83
5.1	Kesimpulan.....	83
5.2	Saran.....	85
DAFTAR PUSTAKA		87

DAFTAR TABEL

	HALAMAN
Tabel 3. 1 Skala Dan Bobot.	38
Tabel 3. 2 Confusion Matrix	43
Tabel 3. 3 Jadwal Penelitian.....	44
Tabel 4. 1 Kuesioner Kepuasan Pengguna UMSU Academy.....	57
Tabel 4. 2 Interval Kategori kepuasan	61
Tabel 4. 3 Perhitungan Probabilitas Class.....	62
Tabel 4. 4 Probabilitas Content (C1).....	63
Tabel 4. 5 Probabilitas Content (C2).....	64
Tabel 4. 6 Probabilitas Content (C3).....	65
Tabel 4. 7 Probabilitas Content (C4).....	66
Tabel 4. 8 Probabilitas Content (C5).....	67
Tabel 4. 9 Probabilitas Ease of use (E4)	69
Tabel 4. 10 Probabilitas Ease of use (E5)	70
Tabel 4. 11 Potongan Data Testing	70
Tabel 4. 12 Menghitung Prior Probabilitas	71
Tabel 4. 13 Fitur C1	72
Tabel 4. 14 Fitur C2	73
Tabel 4. 15 Fitur C3	73
Tabel 4. 16 Fitur C4	74
Tabel 4. 17 Fitur C5	75
Tabel 4. 18 Fitur E4	76
Tabel 4. 19 Fitur E5	77
Tabel 4. 20 Hasil Posterior Tanpa Normalisasi	78
Tabel 4. 21 Hasil Posterior Naïve Bayes	80
Tabel 4. 22 Perhitungan Precision, Recall, F1-Score Dan Akurasi	81

DAFTAR GAMBAR

	HALAMAN
Gambar 2. 1 Logo Aplikasi UMSU Academy	7
Gambar 2. 2 Tampilan Isi Aplikasi UMSU Academy	12
Gambar 2. 3 Teknik Analisis Data	14
Gambar 2. 4 Tingkat Kepuasan Pengguna	19
Gambar 2. 5 Teknik Data Mining	20
Gambar 2. 6 Proses Data Mining	29
Gambar 3. 1 Use Case Diagram	40
Gambar 3. 2 Activity Diagram	40
Gambar 3. 3 Desain Interface	41
Gambar 4. 1 Tampilan Utama Sistem	46
Gambar 4. 2 Tampilan Halaman Dashboard	47
Gambar 4. 3 Tampilan Halaman Menu Dataset	48
Gambar 4. 4 Tampilan Menu Initial Process	49
Gambar 4. 5 Tampilan Menu performance Persentase Training	50
Gambar 4. 6 Data Training	50
Gambar 4. 7 Data Testing	51
Gambar 4. 8 Tampilan Menu Performance Evaluasi Model	51
Gambar 4. 9 Perhitungan Manual Precision, Recall, F1-Score, Akurasi	53
Gambar 4. 10 Tampilan Classification Report	53
Gambar 4. 11 Tampilan Input Data	56
Gambar 4. 12 Tampilan Hasil Prediksi	56

BAB I

PENDAHULUAN

1.1 Latar Belakang Masalah

Dalam bidang pendidikan penggunaan teknologi sangat membantu dalam proses pembelajaran (Criollo-C et al., 2021), salah satu contoh penggunaan teknologi dalam bidang pendidikan adalah menjadikan *e-learning* sebagai media untuk proses pembelajaran (Sajiatmojo et al., 2021). Penggunaan e-learning dalam proses pembelajaran membantu mahasiswa dalam memahami materi ataupun pesan pembelajaran (Fakhruddin et al., 2022), tidak jarang beberapa universitas menggunakan *e-learning* sebagai sarana pembelajaran mahasiswa baik dalam jaringan atau luar jaringan yang disebut sebagai *Hybrid Learning System* (Li et al., 2021). Dalam suatu universitas memiliki sistem informasi yang bertujuan untuk meningkatkan efisiensi dan efektivitas, dimulai dengan fungsional dan operasi administrasi (Prasetyo Utomo & Mariana, 2023). Sistem informasi sangat membantu dalam cakupan ilmu teknologi informasi dalam melakukan pekerjaan di berbagai bidang salah satu diantaranya seperti jadwal perkuliahan mahasiswa (Lukman Santoso & Juni Amanullah, 2022).

Universitas Muhammadiyah Sumatera Utara memiliki sebuah aplikasi sistem informasi akademik yang bernama *UMSU Academy*, sistem informasi akademik bertujuan untuk membantu perkuliahan mahasiswa (Sobral, 2020). Sebelum adanya *UMSU Academy*, Universitas Muhammadiyah Sumatera Utara memiliki aplikasi yang serupa dengan *UMSU Academy* yaitu *UMSU Mobile*. *UMSU mobile* melakukan pembaharuan dan evaluasi mengenai isi dan tampilan dari aplikasi tersebut, kini *UMSU Mobile* resmi berganti menjadi *UMSU*

Academy. Aplikasi seperti sistem informasi akademik akan digunakan oleh mahasiswa dalam proses administrasi ataupun sebagai salah satu informasi (Parras-Burgos et al., 2020). Pada dasarnya setiap aplikasi yang berjalan pasti memiliki kendala ataupun permasalahan yang biasa terjadi pada isi aplikasi ataupun tampilan aplikasi (Wijayanto & Susetyo, 2022). Error merupakan salah satu kendala yang sering kali terjadi pada sebuah aplikasi (Razak & Nasution, 2022), seperti yang terdapat pada UMSU *Academy* dalam melakukan input data serta dalam melakukan download data berupa file pada fitur kartu rencana studi, kartu hasil studi dan transkrip nilai.

Penerapan penggunaan aplikasi sebagai proses akademik maupun administrasi berkaitan dengan kecepatan dalam mengakses fitur – fitur yang terdapat dalam aplikasi (Afrianto et al., 2021). Jika layanan aplikasi yang digunakan sesuai dengan harapan pengguna, dapat disimpulkan bahwa aplikasi tersebut memiliki kualitas yang baik (Chan et al., 2022). *Error* yang terdapat pada aplikasi merupakan hal yang sering terjadi maka dari itu akan dilakukan penelitian mengenai kepuasan pengguna (Lattu, A., Sihabuddin, & Jatmika, 2022), yang bertujuan sebagai tolak ukur bahan evaluasi bagi developer yang nantinya dapat meningkatkan kualitas aplikasi (Nurohman & Nurhayati, 2021). Kepuasan pengguna juga merupakan respon atau umpan balik pengguna aplikasi dalam menggunakan isi dari aplikasi tersebut (Kinanti et al., 2021). Akibat dari kendala jaringan (*Error*) maupun akses yang sulit dituju pada aplikasi, menentukan kepuasan pengguna pada penelitian ini menjadi solusi terbaik (S Willermark, N Pantic, 2021).

Untuk mengukur tingkat kepuasan pengguna aplikasi *UMSU Academy*, penelitian ini menggunakan metode *Naïve Bayes* sebagai alat untuk mengukur kepuasan pengguna (Ramadhani et al., 2024). Algoritma *Naïve Bayes* merupakan salah satu metode pengklasifikasi dengan menggunakan probabilitas serta menggunakan metode statistik (Putry, 2022).

Pada penelitian ini algoritma *Naïve Bayes* digunakan dalam menentukan kepuasan pengguna aplikasi *UMSU Academy* dengan melakukan klasifikasi tingkat kepuasan. Klasifikasi merupakan bentuk dari analisis data dengan menggunakan teknik untuk menentukan anggota kelompok berdasarkan data (Azis et al., 2020). Klasifikasi dengan menggunakan metode *Naïve Bayes* bertujuan untuk menentukan ataupun melihat akurasi data (Perez & Perez, 2021), untuk menentukan mahasiswa Universitas Muhammadiyah Sumatera Utara merasa sangat tidak puas, tidak puas, cukup puas, puas dan sangat puas dalam penggunaan aplikasi *UMSU Academy* baik itu dari isi aplikasi maupun tampilan dari aplikasi tersebut.

Penggunaan *Naïve bayes* dapat memberikan efesiensi dan fleksibilitas yang signifikan dalam melakukan klasifikasi, serta memerlukan dataset pelatihan yang terbatas untuk menerapkan nilai parameter dalam memfasilitasi proses klasifikasi kepuasan pengguna aplikasi (Rais et al., 2022). Algoritma *Naïve Bayes* digunakan dalam menentukan kepuasan pengguna aplikasi *UMSU Academy* yang memiliki kendala dalam penerapan sistem atau terdapat permasalahan *error* pada laman yang di kunjungi, serta dapat di selesaikan dengan menentukan kepuasan pengguna yang apabila nantinya terdapat kendala pada aplikasi akan menjadi evaluasi untuk pihak *developer*. Dengan hasil akhir nilai akurasi kepuasan

pengguna aplikasi *UMSU Academy* untuk mengetahui atau menentukan kepuasan mahasiswa terhadap penerapan sistem informasi akademik Universitas Muhammadiyah Sumatera Utara berbasis mobile yang terdapat pada *UMSU Academy*.

Diharapkan dengan adanya penelitian ini dapat menjadi evaluasi maupun acuan dalam meningkatkan kualitas penerapan sistem informasi akademik Universitas Muhammadiyah Sumatera Utara demi kepuasan pengguna yaitu kepada seluruh mahasiswa yang menggunakan *UMSU Academy*.

1.2 Rumusan Masalah

Berdasarkan uraian latar belakang masalah yang ada, maka penulis dapat merumuskan permasalahan yang diangkat dalam penelitian ini adalah sebagai berikut:

1. Bagaimana mengukur kepuasan pengguna aplikasi *UMSU Academy* menggunakan metode *Naïve Bayes*?
2. Apakah algoritma *Naïve Bayes* efektif dalam mengevaluasi kepuasan pengguna terhadap aplikasi *UMSU Academy* di Universitas Muhammadiyah Sumatera Utara?
3. Bagaimana tingkat kepuasan pengguna terhadap aplikasi *UMSU Academy* berdasarkan data yang diperoleh dengan menggunakan metode *Naïve Bayes*?

1.3 Batasan Masalah

Berdasarkan uraian latar belakang yang ada, untuk menghindari pembahasan yang lebih luas terkait Analisis Kepuasan Pengguna Aplikasi *UMSU*

Academy Menggunakan Metode *Naive Bayes*. maka penulis melakukan pembatasan masalah sebagai berikut:

1. Penelitian ini hanya membahas kepuasan pengguna aplikasi UMSU *Academy* pada mahasiswa Universitas Muhammadiyah Sumatera Utara tahun ajaran 2021.
2. Penelitian ini menggunakan algoritma *Naïve Bayes* dalam melakukan analisis kepuasan pengguna.
3. Hanya menggunakan lima variabel yaitu *content* (isi), *format* (tampilan), *accuracy* (keakuratan), *timelines* (ketepatan waktu) dan *ease of use* (kemudahan pengguna). Kriteria parameter kepuasan pengguna dijabarkan menjadi lima kategori yaitu sangat tidak puas, tidak puas, cukup puas, puas, dan sangat puas.

1.4 Tujuan Penelitian

Berdasarkan latar belakang dan perumusan masalah yang telah dikemukakan, maka tujuan yang ingin dicapai dalam penelitian ini adalah:

1. Untuk mengetahui sejauh mana tingkat kepuasan pengguna aplikasi UMSU *Academy*.
2. Penelitian ini bertujuan untuk mengevaluasi efektivitas algoritma *Naïve bayes* dalam mengklasifikasi kepuasan pengguna aplikasi UMSU *Academy*.
3. Penelitian ini dilakukan untuk mengukur tingkat akurasi kepuasan pengguna aplikasi UMSU *Academy* menggunakan metode *Naïve Bayes*.

1.5 Manfaat Penelitian

Adapun manfaat dari penelitian ini dibagi menjadi dua bagian yaitu:

1. Aspek Teoritis

Penelitian ini memiliki manfaat teoritis yaitu seperti bagi mahasiswa dapat memberikan pengetahuan maupun wawasan dan menjadikan penelitian ini sebagai referensi pada penelitian lainnya serta bagi pembaca dapat menambah pengetahuan dan dapat dijadikan sebagai bahan acuan untuk melakukan penelitian yang berkaitan dengan kepuasan penggunaan dengan menggunakan metode *Naïve Bayes*.

2. Aspek Praktis

Penelitian ini memiliki manfaat praktis yaitu bagi pihak biro sistem informasi Universitas Muhammadiyah Sumatera Utara, penelitian ini dapat membantu dalam mengukur tingkat kepuasan pengguna aplikasi UMSU *Academy* yaitu kepuasan mahasiswa terhadap UMSU *Academy*. Dengan adanya penelitian ini dapat diidentifikasi aspek-aspek yang sudah baik serta yang memerlukan perbaikan dan bagi penulis penelitian ini dapat menambang pengetahuan implementasi algoritma *Naïve Bayes* dalam melakukan analisis kepuasan pengguna aplikasi.

BAB II

LANDASAN TEORI

2.1 Aplikasi

Aplikasi merupakan perangkat lunak yang terdapat di dalam komputer dan memiliki fungsi-fungsi tertentu. Aplikasi sangat membantu dalam melayani kebutuhan manusia dan beberapa aktivitas lainnya seperti sistem penjualan, game, pelayanan masyarakat, pada bidang pendidikan, dan hampir semua kegiatan menggunakan aplikasi. Perangkat lunak seperti aplikasi memanfaatkan kemampuan komputer untuk melakukan tugas yang diinginkan pengguna, salah satunya adalah pengelolaan kata, lembar kerja dan pemutar media (Oman Sumantri, 2015).

Aplikasi terbagi menjadi tiga platform yaitu *web*, *desktop*, dan *mobile*. Aplikasi *mobile* merupakan aplikasi yang sering digunakan dalam kegiatan sehari-hari, salah satunya pada bidang pendidikan. Perkembangan aplikasi *mobile* ini didukung karena kecanggihan penerapan teknologi di zaman yang semakin maju dan berkembang (Larasati et al., 2021).



Gambar 2. 1 Logo Aplikasi UMSU Academy

Aplikasi *mobile* memiliki platform mobile yang mengulas sistem operasi populer seperti *Android*, *IOS* serta *tools* atau *framework* seperti *flutter*, *react native*, *kotlin*, serta *swift*. Adapun contoh dari aplikasi mobile seperti aplikasi UMSU *Academy* yang digunakan oleh mahasiswa Universitas Muhammadiyah Sumatera Utara dalam menerima informasi mengenai kartu rencana studi *online*, kartu hasil studi *Online*, serta transkrip nilai.

Aplikasi mobile memiliki jenis-jenis seperti:

1. *Native*

Aplikasi *native* adalah jenis aplikasi yang dirancang dan dikembangkan untuk berjalan pada sistem operasi tertentu (seperti *Android*, *iOS*, atau *Windows*) dengan menggunakan bahasa pemrograman dan *tools* khusus untuk *platform* tersebut. Aplikasi ini diinstal langsung di perangkat dan dapat memanfaatkan semua fitur dan kemampuan perangkat keras (*hardware*) serta perangkat lunak (*software*) yang ada pada sistem operasi tersebut. Kelebihan dari aplikasi *native* seperti kinerja optimal, akses penuh ke fitur perangkat keras dan *API* sistem operasi, pengalaman pengguna yang lebih baik, serta memiliki pengembangan untuk pengamanan. Adapun kekurangan Aplikasi *native* seperti pengembangan ganda, biaya pengembangan yang lebih tinggi serta waktu pengembangan yang lebih lama.

2. *Hybrid*

Aplikasi *hybrid* adalah jenis aplikasi yang dikembangkan dengan menggunakan teknologi *web*, seperti *HTML*, *CSS*, dan *JavaScript*, tetapi dapat dijalankan di perangkat seluler seperti aplikasi *native*.

Aplikasi *hybrid* menggabungkan elemen-elemen dari aplikasi *native* dan aplikasi berbasis *web*. Mereka dikemas dalam kontainer *native*, yang memungkinkan aplikasi tersebut dijalankan di perangkat *mobile* sambil tetap bisa mengakses sebagian besar fitur perangkat keras melalui *plugin* atau *API* yang disediakan oleh platform pengembangan *hybrid*. Kelebihan dari aplikasi *hybrid* seperti pengembangan satu kode untuk banyak *platform*, biaya pengembangan yang lebih rendah, waktu pengembangan lebih cepat, kemudahan dalam pemeliharaan, serta akses ke *plugin* dan *API native*. Adapun kekurangan aplikasi *hybrid* kinerja yang lebih rendah dibandingkan dengan aplikasi *native*, pengalaman pengguna (*UX*) yang kurang optimal, terbatas pada fitur tertentu, serta keterbatasan dalam penggunaan *plugin*.

3. *Web*

Aplikasi *web* adalah aplikasi yang dijalankan melalui *browser web*, yang dapat diakses oleh pengguna melalui koneksi internet tanpa perlu diunduh atau dipasang di perangkat pengguna. Aplikasi *web* berfungsi dengan menggunakan teknologi *web* standar seperti *HTML*, *CSS*, *JavaScript*, dan *PHP* (atau bahasa *server-side* lainnya) untuk memberikan interaksi yang dinamis kepada pengguna. Kelebihan dari aplikasi *web* seperti tidak perlu instalasi, dapat diakses dari berbagai perangkat, pemeliharaan yang lebih mudah, pengurangan biaya pengembangan, serta akses terpusat. Adapun kekurangan dari aplikasi *web* keterbatasan kinerja, ketergantungan pada koneksi internet, pengalaman pengguna (*UX*) yang terbatas dan masalah keamanan.

2.2 UMSU Academy

UMSU Academy merupakan aplikasi sistem informasi akademik berbasis *mobile* yang digunakan oleh mahasiswa dan mahasiswi Universitas Muhammadiyah Sumatera Utara dalam melakukan proses yang berkaitan dengan akademik dan informasi seputar perkuliahan. Universitas Muhammadiyah Sumatera Utara (UMSU) merupakan salah satu universitas swasta yang didirikan pada tanggal 27 Februari 1957 di kota Medan, Sumatera Utara, Indonesia. Aplikasi UMSU Academy baru saja digunakan oleh seluruh mahasiswa Universitas Muhammadiyah Sumatera Utara dalam pengisian kartu rencana studi, melihat informasi mengenai kartu hasil studi dan transkrip nilai serta mendownload data yang terdapat pada aplikasi tersebut. Aplikasi UMSU Academy dirilis pada tanggal 14 juli 2023 serta terakhir di perbaharui pada tanggal 2 November 2023.

Dalam aplikasi UMSU Academy terdapat beberapa fitur seperti:

1. Fakultas

Pada fitur fakultas terdapat beberapa nama-nama serta informasi mengenai fakultas dan jurusan kuliah seperti Fakultas Agama Islam, Fakultas Keguruan dan Ilmu Pendidikan, Fakultas Ilmu Sosial dan Ilmu Politik, Fakultas Pertanian, Fakultas Ekonomi dan Bisnis, Fakultas teknik, Fakultas Ilmu komputer dan Teknologi Informasi, Fakultas Kedokteran serta Pascasarjana Universitas Muhammadiyah Sumatera Utara.

2. *Live Chat*

Pada fitur *live chat* berfungsi sebagai fitur untuk melakukan pengaduan ataupun keluhan serta saran terhadap kendala yang terdapat pada sistem informasi akademik Universitas Muhammadiyah Sumatera Utara maupun mengenai perkuliahan.

3. Kalender

Fitur kalender berfungsi sebagai alat pemberitahuan kepada seluruh mahasiswa mengenai informasi akademik.

4. Kontak

Fitur kontak berfungsi sebagai informasi mengenai kampus ataupun kontak yang bisa dihubungkan untuk mengetahui informasi lebih mendalam mengenai Universitas Muhammadiyah Sumatera Utara.

5. *E-learning*

Fitur *E-learning* berfungsi sebagai media pembelajaran *online* yang digunakan mahasiswa dalam mengakses mata kuliah setiap jurusan mahasiswa Universitas Muhammadiyah Sumatera Utara.

6. SKPI

Fitur SKPI berfungsi sebagai informasi mengenai pendaftaran Surat Keterangan Pendamping Ijazah mahasiswa secara online.



Gambar 2. 2 Tampilan Isi Aplikasi UMSU Academy

2.3 Analisis Data

Analisis data merupakan proses menentukan kesimpulan dengan melakukan pengorganisasian, pemahaman, interpretasi dari data yang dikumpulkan (Suwandi, 2022). Analisis data digunakan untuk mengelolah maupun menemukan data secara sistematis baik dengan melakukan wawancara, observasi dan lainnya untuk meningkatkan pengetahuan (Ahmad & Muslimah, 2021). Analisis data menggunakan berbagai metode dan statistik, matematika dan komputasi dalam mengelolah serta menginterpretasikan data. Proses analisis data melakukan penyaringan data, pembersihan data (data cleansing), transformasi data, visualisasi data, pemodelan statistic serta penarikan kesimpulan (Rizky Fadilla & Ayu Wulandari, 2023).

Analisis data sering digunakan dalam melakukan ataupun menganalisis penelitian terhadap sebuah informasi untuk memecahkan masalah yang diteliti agar dapat menghasilkan kesimpulan yang akurat. Teknik analisis dibagi menjadi dua yaitu analisis kuantitatif dan analisis kualitatif, yang membedakan dari kedua

jenis analisis tersebut adalah sifat data dari masing masing ketentuan data yang ingin diolah.

Langkah-langkah dalam Analisis Data:

1. Pengumpulan Data

Mengumpulkan data yang relevan dari berbagai sumber, baik dari data yang sudah ada, survei, eksperimen, atau pengamatan. Data ini bisa berupa data primer (misalnya, hasil survei) atau data sekunder (misalnya, laporan perusahaan atau data publik). Contoh seperti mengumpulkan data penjualan harian dan data biaya iklan selama beberapa bulan terakhir.

2. Pembersian Data

Data yang dikumpulkan dan dibersihkan dan dipersiapkan sebelum analisis, seperti dilakukan pembersihan data untuk menangani data yang hilang, duplikasi, atau kesalahan pencatatan adapun transformasi data untuk menyesuaikan format data atau melakukan agregasi data serta melakukan normalisasi data untuk mengatur skala data agar konsisten. Contoh menghapus entri duplikat, mengganti nilai yang hilang dengan nilai rata-rata atau menggunakan interpolasi.

3. Analisis Data

Pada analisis data menggunakan teknik yang sesuai untuk mendapatkan wawasan dari data seperti analisis statistik yang melakukan uji hipotesis, regresi, dan analisis varians (*ANOVA*). Adapun model prediktif yang menggunakan algoritma *machine learning* untuk memprediksi atau mengklasifikasikan data serta

pencarian yang mengidentifikasi pola atau tren dalam data. Contoh menggunakan regresi linear untuk melihat bagaimana biaya iklan memengaruhi penjualan, atau menggunakan analisis korelasi untuk menemukan hubungan antar variabel.

4. Interpretasi Data

Setelah melakukan analisis, hasil yang diperoleh perlu diinterpretasikan untuk menjawab pertanyaan awal atau masalah yang ingin diselesaikan. Ini juga melibatkan penilaian apakah hasil tersebut signifikan secara statistik atau tidak. Contoh menginterpretasikan hasil regresi untuk menentukan apakah biaya iklan memang berhubungan positif dengan peningkatan penjualan, dan sejauh mana pengaruh tersebut signifikan.

5. Visualisasi Data

Visualisasi data menggunakan statistik deskriptif dan teknik visualisasi (seperti grafik, histogram, *box plot*, dan diagram sebar) untuk mendapatkan gambaran umum tentang data. *EDA* membantu untuk memahami distribusi data, hubungan antar variabel, dan potensi masalah yang perlu diperbaiki.



Gambar 2. 3 Teknik Analisis Data

2.4 Kepuasan Pengguna

Kepuasan adalah suatu keadaan yang diputuskan setelah melakukan pengalaman yang didapatkan. Salah satu contoh kepuasan adalah kepuasan pengguna, Kepuasan pengguna adalah faktor yang dapat memengaruhi loyalitas pengguna terhadap sebuah aplikasi (Pidie, 2023). Kepuasan pengguna sering kali digunakan untuk menjadi bahan evaluasi developer dalam mengembangkan sebuah sistem ataupun aplikasi. Kepuasan pengguna dapat disimpulkan sebagaimana pengguna merasa puas ketika menggunakan aplikasi atau suatu sistem informasi yang dapat memenuhi keinginan atau ekspektasi mereka pada suatu sistem atau aplikasi tersebut.

Salah satu contoh kepuasan pengguna seperti kepuasan mahasiswa terhadap kualitas pendidikan yang diberikan oleh sebuah institusi, baik itu dari segi sistem informasi akademik maupun informasi seputar pendidikan lainnya. Pengguna dalam suatu aplikasi akan merasa tidak puas jika ekspektasi yang diharapkan tidak sesuai, jika pengguna merasa puas terhadap aplikasi maka kinerja yang dicapai telah melebihi ekspektasi pengguna (Setiawan & Novita, 2021).

Faktor- faktor yang mempengaruhi kepuasan pengguna diantaranya:

1. Kemudahan Penggunaan (*Usability*)

Kemudahan penggunaan merujuk pada sejauh mana aplikasi tersebut mudah dipahami dan digunakan oleh penggunanya. Aplikasi yang memiliki antarmuka yang intuitif, navigasi yang sederhana, dan instruksi yang jelas lebih memuaskan pengguna. Pengguna lebih cenderung merasa puas dengan aplikasi yang tidak membingungkan,

yang memungkinkan mereka untuk dengan cepat memahami dan menggunakannya tanpa banyak waktu yang dihabiskan untuk belajar mengetahui alur kerjanya.

2. Kinerja Aplikasi (*Performance*)

Kinerja aplikasi melibatkan seberapa cepat aplikasi berjalan dan seberapa stabil aplikasinya. Ini termasuk waktu muat (*load time*), responsivitas, dan sejauh mana aplikasi berjalan tanpa mengalami crash atau masalah teknis. Aplikasi yang berjalan lambat atau sering crash sangat menurunkan kepuasan pengguna.

3. Desain Antarmuka Pengguna (*UI Design*)

Desain antarmuka pengguna (*UI Design*) adalah tampilan dan interaksi yang terjadi antara pengguna dan aplikasi. Desain *UI* yang menarik dan nyaman digunakan untuk meningkatkan kepuasan pengguna. Aplikasi dengan desain yang bersih, responsif, dan estetis biasanya memberikan pengalaman yang lebih menyenangkan bagi pengguna. Desain yang buruk dapat membuat pengguna merasa frustrasi.

4. Fungsionalitas Aplikasi

Fungsionalitas mengacu pada kemampuan aplikasi untuk menjalankan tugas atau fungsi yang dijanjikan dengan efektif. Aplikasi harus memenuhi kebutuhan pengguna dengan menyediakan fitur yang berguna dan mudah diakses. Aplikasi yang memiliki fungsi yang sesuai dengan kebutuhan pengguna dan berjalan sesuai dengan harapan mereka dan lebih memuaskan. Fungsi yang tidak berfungsi dengan

baik atau tidak sesuai dengan kebutuhan pengguna bisa menyebabkan ketidakpuasan.

5. Pengalaman Pengguna (*User Experience - UX*)

Pengalaman pengguna mencakup keseluruhan perasaan atau emosi pengguna saat berinteraksi dengan aplikasi, mulai dari antarmuka hingga interaksi sehari-hari. Ini mencakup kemudahan penggunaan, desain, dan kepuasan secara keseluruhan. Aplikasi yang memberikan pengalaman positif, menyenangkan, dan memenuhi kebutuhan pengguna dan menghasilkan tingkat kepuasan yang tinggi. Misalnya, aplikasi yang mudah digunakan dengan desain yang menyenangkan dan memberikan pengalaman pengguna yang lebih baik.

6. Keamanan dan Privasi

Keamanan dan privasi merujuk pada seberapa baik aplikasi melindungi data pribadi dan informasi pengguna. Pengguna sangat peduli dengan bagaimana data mereka dikumpulkan dan digunakan oleh aplikasi. Aplikasi yang memastikan keamanan dan privasi pengguna dengan fitur seperti enkripsi data, kebijakan privasi yang jelas, dan kontrol atas data pribadi dan meningkatkan rasa aman pengguna, yang berkontribusi pada kepuasan mereka.

7. Dukungan Pelanggan

Dukungan pelanggan mencakup berbagai fungsi aplikasi membantu penggunanya saat menghadapi masalah atau pertanyaan, baik melalui *FAQ*, *chat support*, atau *email*. Aplikasi yang menyediakan dukungan pelanggan yang cepat, ramah, dan efektif dapat membuat pengguna

merasa dihargai dan lebih puas. Sebaliknya, jika pengguna kesulitan mendapatkan bantuan, ini dapat menurunkan tingkat kepuasan mereka.

8. Keandalan dan Stabilitas

Keandalan mengacu pada konsistensi dan kemampuan aplikasi untuk berfungsi dengan baik dalam berbagai kondisi. Stabilitas adalah seberapa jarang aplikasi mengalami gangguan atau masalah teknis. Aplikasi yang sering mengalami crash atau masalah teknis dapat membuat pengguna frustrasi dan menurunkan kepuasan mereka.

9. Pembaruan dan Perbaikan

Aplikasi yang secara teratur diperbarui dan diperbaiki, baik untuk menambahkan fitur baru maupun untuk memperbaiki bug atau masalah teknis, cenderung lebih disukai oleh pengguna. Pengguna lebih cenderung merasa puas dengan aplikasi yang terus diperbarui dan ditingkatkan. Pembaruan yang sering menunjukkan bahwa pengembang peduli terhadap pengalaman pengguna dan siap untuk memperbaiki masalah yang ada.

10. Harga atau Biaya

Harga atau biaya aplikasi (terutama untuk aplikasi berbayar atau aplikasi *premium*) juga berperan penting dalam kepuasan pengguna. Apakah aplikasi tersebut menawarkan nilai yang sesuai dengan harga yang dibayar oleh pengguna. Pengguna cenderung merasa lebih puas jika mereka merasa bahwa aplikasi memberikan nilai yang baik untuk uang yang mereka keluarkan. Misalnya, aplikasi yang memiliki banyak

fitur bermanfaat atau memberikan pengalaman yang menyenangkan dengan harga yang wajar.

Faktor-faktor yang mempengaruhi kepuasan pengguna menurut (Erlich & Zviran, 2003):

1. Hubungan antara pengguna dengan sistem informasi
2. Hubungan antara manajemen organisasi dengan sistem informasi
3. Informasi yang diterima dari sistem
4. Penyediaan layanan dari sistem informasi
5. Fitur-fitur yang dimiliki oleh sistem informasi



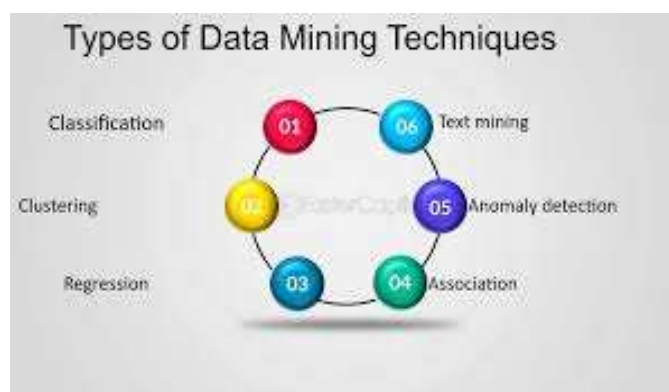
Gambar 2. 4 Tingkat Kepuasan Pengguna

2.5 Data Mining

Data mining adalah proses semi otomatis yang menggunakan teknik seperti matematika, statistik, *artificial intelligence* serta *machine learning* untuk mengidentifikasi informasi pengetahuan potensial yang tersimpan di dalam database (Zai, 2022). Data mining merupakan suatu alat yang memungkinkan para pengguna untuk mengakses sebuah data secara cepat dengan jumlah data yang besar (Wu et al., 2021). Data mining merupakan pencarian pengetahuan dalam

database dalam proses identifikasi pola-pola yang valid, dapat dipahami secara muda, dan berpotensi bermanfaat bagi pengguna.

Data mining adalah proses untuk menemukan pengetahuan yang diambil dari sekumpulan data yang memiliki volume yang besar. Pada dasarnya data mining memiliki sebuah aplikasi yang berfungsi untuk mengelolah data di bidang pendidikan, bisnis, pengendalian produksi serta analisis pasar seperti memperoleh pola dan hubungan yang dapat dimanfaatkan untuk meningkatkan cara penjualan, pembelajaran, serta pengelolaan sumber daya yang lebih baik (Firdaus, 2017).



Gambar 2. 5 Teknik Data Mining

Data mining memiliki enam jenis teknik data, yaitu:

1. Klasifikasi (*classification*)

Klasifikasi (*classification*) melakukan generalisasi struktur yang untuk diaplikasikan pada data-data baru. Klasifikasi adalah pengelompokan objek kedalam kelas tertentu yang berdasarkan kelompoknya dan disebut sebagai kelas (*class*) (Putro et al., 2020). Klasifikasi merupakan suatu kegiatan yang melakukan pelatihan dan pembelajaran terhadap fungsi target f yang memetakan setiap vektor (set fitur) x kedalam satu dari jumlah label kelas y yang tersedia (Febriani, 2022). Klasifikasi memiliki dua tahapan dalam melakukan

proses pengklasifikasian. Tahapan pertama adalah fase training atau tahapan learning yang menggunakan beberapa data dengan kelas datanya digunakan untuk membentuk model perkiraan, tahapan kedua adalah *fase testing* atau proses *test* bertujuan menguji model yang dibangun untuk mengetahui tingkat keakuratan dari model tersebut. Dapat disimpulkan apabila akurasi memenuhi persyaratan, model ini digunakan untuk memprediksi kelas data yang belum diketahui kelasnya. Proses klasifikasi adalah menentukan sebuah model untuk memprediksi kelas data yang belum diketahui dengan mengelompokkan dan memisahkan kelas data tersebut (Widyadara & Irawan, 2019). Metode yang sering menggunakan klasifikasi diantaranya seperti *Naïve Bayes*, *Decision Tree*, *Neural Network*, *Genetic Algorithm*, *Fuzzy* dan *K-Nearest Neighbour* (Hendrian, 2018). Klasifikasi mengelompokkan informasi yang digunakan untuk mengklasifikasi informasi yang belum memiliki kelas kedalam kelas yang telah ada. Proses klasifikasi terdapat dua tahapan awal, tahapan *learning (fase training)* memiliki sebagian informasi yang dikenal kelas informasinya yang digunakan untuk membentuk model perkiraan. Tahapan kedua, proses *test (fase testing)* menguji model yang dibentuk agar diketahui tingkatan keakuratan dari model tersebut. Bila akurasi memenuhi persyaratan, model ini dapat digunakan untuk memprediksi kelas data yang belum diketahui kelasnya.

2. Klasterisasi (*clustering*)

Klasterisasi (*clustering*) melakukan pengelompokan data yang tidak diketahui label kelasnya ke dalam sejumlah kelompok tertentu sesuai dengan kemiripan datanya (Aulia, 2021). Clustering adalah teknik untuk mengelompokkan data ke dalam grup yang homogen (*cluster*) berdasarkan kesamaan fitur. Data dalam satu cluster harus lebih mirip dengan yang lain dibandingkan dengan data di cluster lain. Klasterisasi (*clustering*) juga merupakan salah satu teknik dalam analisis data yang digunakan untuk mengelompokkan objek-objek yang memiliki kesamaan berdasarkan fitur atau atribut tertentu. Dalam klasterisasi, data dibagi ke dalam kelompok-kelompok atau *cluster* di mana setiap objek dalam satu kelompok memiliki kesamaan yang lebih tinggi dengan objek lainnya di dalam kelompok tersebut daripada dengan objek di kelompok lain. Klasterisasi termasuk dalam kategori pembelajaran tanpa pengawasan (*unsupervised learning*), dapat diartikan bahwa algoritma tidak memerlukan label atau kategori yang sudah diketahui sebelumnya untuk melakukan pengelompokan. Adapun beberapa algoritma klasterisasi seperti *K-Means clustering* yang merupakan Salah satu algoritma klasterisasi yang paling populer (Wahyu & Rushenda, 2022). *K-Means* membagi data ke dalam k jumlah klaster, di mana k adalah jumlah yang ditentukan oleh pengguna. Algoritma ini berusaha meminimalkan jarak antara titik data dengan pusat klaster. Kelebihan *K-Means* cepat dan mudah diimplementasikan serta Efektif pada dataset besar jika jumlah klaster

yang tepat diketahui sedangkan kekurangan dari *K-Means* adalah memerlukan pemilihan jumlah klaster k yang tepat dan Tidak cocok untuk data dengan distribusi yang sangat tidak teratur atau memiliki bentuk klaster yang kompleks. Selain *K-Means* adapun *Hierarchical Clustering* yang merupakan teknik membangun hierarki klaster, yang bisa berupa pohon dendrogram. Ada dua pendekatan utama dalam *hierarchical clustering* yaitu *agglomerative* (penggabungan klaster) dan *divisive* (pembagian klaster). *Hierarchical Clustering* membangun hierarki klaster secara bertahap, baik dengan tahapan *agglomerative (bottom-up)* atau *divisive (top-down)*. Dalam *agglomerative clustering*, setiap data dimulai sebagai klaster tunggal, dan pada setiap langkah, dua klaster yang paling mirip digabungkan. Kelebihan pada *Hierarchical Clustering* adalah tidak memerlukan penentuan jumlah klaster terlebih dahulu serta memberikan pemahaman lebih baik tentang struktur data sedangkan kekurangan *Hierarchical Clustering* adalah lebih lambat dibandingkan dengan algoritma seperti *K-Means*, terutama untuk dataset besar dan tidak efektif pada data yang sangat besar. Terdapat *DBSCAN (Density-Based Spatial Clustering of Applications with Noise)* yaitu algoritma ini berbasis pada kepadatan data, yang mengelompokkan titik data berdasarkan area yang lebih padat. *DBSCAN* dapat menangani data dengan bentuk klaster yang lebih kompleks dan juga dapat mengidentifikasi *noise* (data yang tidak termasuk dalam klaster). Algoritma klasterisasi yang berbasis pada kepadatan. *DBSCAN* mendeteksi klaster yang lebih padat dengan

tahapan mengelompokkan data yang memiliki banyak tetangga dalam radius tertentu dan memisahkan data yang jaraknya jauh sebagai *noise* atau anomali. Kelebihan *DBSCAN* adalah tidak perlu menentukan jumlah kluster sebelumnya, Baik untuk data dengan bentuk kluster yang tidak teratur serta mendeteksi anomali secara otomatis sedangkan kekurangan *DBSCAN* adalah memilih parameter yang tepat (terutama *eps*) bisa sulit dan tidak cocok untuk dataset dengan variasi kepadatan yang tinggi. Terdapat *Gaussian Mixture Models (GMM)* merupakan algoritma yang mengasumsikan bahwa data berasal dari beberapa distribusi *Gaussian* (normal). *GMM* adalah model probabilistik yang dapat memberikan informasi lebih lanjut tentang keanggotaan data dalam kluster tertentu. *Gaussian Mixture Models (GMM)* adalah model probabilistik yang menganggap bahwa data berasal dari distribusi *Gaussian* (normal) yang berbeda. Setiap kluster dianggap sebagai distribusi *Gaussian*, dan data dikelompokkan berdasarkan probabilitas keanggotaan pada masing-masing distribusi. Kelebihan dari *Gaussian Mixture Models* adalah dapat menangani data dengan bentuk kluster yang lebih kompleks daripada *K-Means* serta menghasilkan probabilitas keanggotaan untuk setiap titik data sedangkan kekurangan *Gaussian Mixture Models* adalah lebih rumit dan lebih mahal secara komputasi dibandingkan algoritma lainnya dan membutuhkan asumsi bahwa data terdistribusi secara normal. Serta terdapat *Agglomerative Clustering* yang merupakan sebuah jenis

klasterisasi hierarkis yang bekerja dengan menggabungkan klaster yang paling mirip satu dengan yang lain pada setiap langkah.

3. Regresi (*regression*)

Regresi (*regression*) menemukan suatu fungsi yang memodelkan data dengan galat (kesalahan prediksi) seminimal mungkin. Regresi digunakan untuk memprediksi nilai numerik berdasarkan hubungan antara variabel independen dan dependen. Teknik ini sering digunakan dalam prediksi nilai seperti harga rumah, penjualan, atau suhu. Metode yang digunakan untuk melakukan regresi seperti, *linier regression* yang berfungsi untuk memprediksi hubungan linier antara variabel independen dan dependen adapun *logistik regression* yang digunakan untuk klasifikasi biner seperti memprediksi sesuatu yang pasti ataupun tidak pasti (Sholeh et al., 2023). Adapun regresi berganda yang merupakan model perpanjangan dari regresi linear sederhana, dengan menemukan hubungan linier antara satu variabel dependen dan beberapa variabel independen (Panggabean et al., 2020). *Regresi polinomial* juga merupakan metode yang sering digunakan jika data menunjukkan pola yang tidak bisa diwakili oleh garis lurus, selain itu regresi ini juga merupakan bentuk *regresi non-linear*, yang Dimana terdapat hubungan antara variabel independen dan dependen digambarkan dengan fungsi polinomial (misalnya, kuadrat, kubus, dll). Selain itu terdapat metode *regresi ridge* dan *lasso*, merupakan teknik regresi yang digunakan untuk mengatasi *overfitting* (penyesuaian model yang terlalu ketat dengan data latih).

Keduanya menggunakan *penalty* pada koefisien model untuk mengurangi kompleksitas. *Ridge regression* menambahkan penalti berupa kuadrat dari nilai koefisien dan *lasso regression* menambahkan penalti berupa nilai mutlak dari koefisien. Adapun regresi kuantil digunakan untuk memodelkan berbagai kuantil (misalnya, median atau kuantil lainnya) dari distribusi variabel dependen.

4. Deteksi anomali (*anomaly detection*)

Deteksi anomali (*anomaly detection*) melakukan identifikasi data yang tidak umum seperti *outlier*, deviasi atau perubahan yang sangat penting serta perlu dilakukan investigasi lanjutan. Deteksi anomali digunakan untuk mengidentifikasi data yang tidak biasa atau menyimpang dari pola umum dalam dataset. Ini penting dalam mendeteksi penipuan, kerusakan sistem, atau kesalahan. Deteksi anomali merupakan analisis data yang digunakan untuk mengidentifikasikan data yang tidak biasa atau berbeda dari pola umum yang ada dalam dataset. Anomali atau *outlier* ini bisa menunjukkan kesalahan, kegagalan sistem, penipuan, atau kejadian yang menarik dan layak untuk dianalisis lebih lanjut. Deteksi anomali digunakan dalam berbagai aplikasi, seperti deteksi penipuan kartu kredit, analisis keamanan jaringan, pemantauan kesehatan perangkat. *Isolation Forest* merupakan metode yang digunakan untuk mengisolasi anomali dalam data dengan tahapan mengidentifikasi titik data yang berbeda dari mayoritas dan *One-Class SVM* yang digunakan untuk mendeteksi data yang berbeda atau *outlier* yang tidak sesuai

dengan pola normal. Jenis anomali salah satunya anomali titik yaitu data tertentu dianggap anomali jika nilai atau fitur individu dari titik data tersebut sangat berbeda dari yang lainnya. Misalnya, jika ada transaksi kartu kredit yang jauh lebih besar dari transaksi lainnya. Adapun anomali kontekstual merupakan anomali yang hanya bisa dipahami dalam konteks tertentu. Sebagai contoh, suhu yang sangat tinggi mungkin tidak aneh pada musim panas dan menjadi anomali pada musim dingin. Anomali kolektif merupakan jenis anomali yang terjadi ketika sekelompok titik data yang saling berhubungan menunjukkan perilaku yang tidak biasa, meskipun setiap titik individu mungkin tidak terlihat anomali. Sebagai contoh, sekumpulan transaksi yang terjadi dalam waktu yang sangat singkat dapat menunjukkan penipuan meskipun setiap transaksi mungkin terlihat normal.

5. Pembelajaran aturan asosiasi (*association rule mining*)

Pembelajaran aturan asosiasi (*association rule mining*) merupakan pencarian relasi antar variabel. Teknik ini digunakan untuk menemukan hubungan antara item dalam dataset yang besar. Seperti, jika seseorang membeli roti, mereka cenderung membeli mentega juga. Metode yang digunakan seperti algoritma *Apriori* untuk menentukan aturan asosiasi dalam transaksi data. *Apriori* adalah algoritma yang paling terkenal untuk penambangan aturan asosiasi. Kelebihan algoritma *Apriori* adalah mudah dipahami dan diimplementasikan serta dapat mengatasi dataset besar jika parameter yang tepat digunakan (seperti support minimum) sedangkan

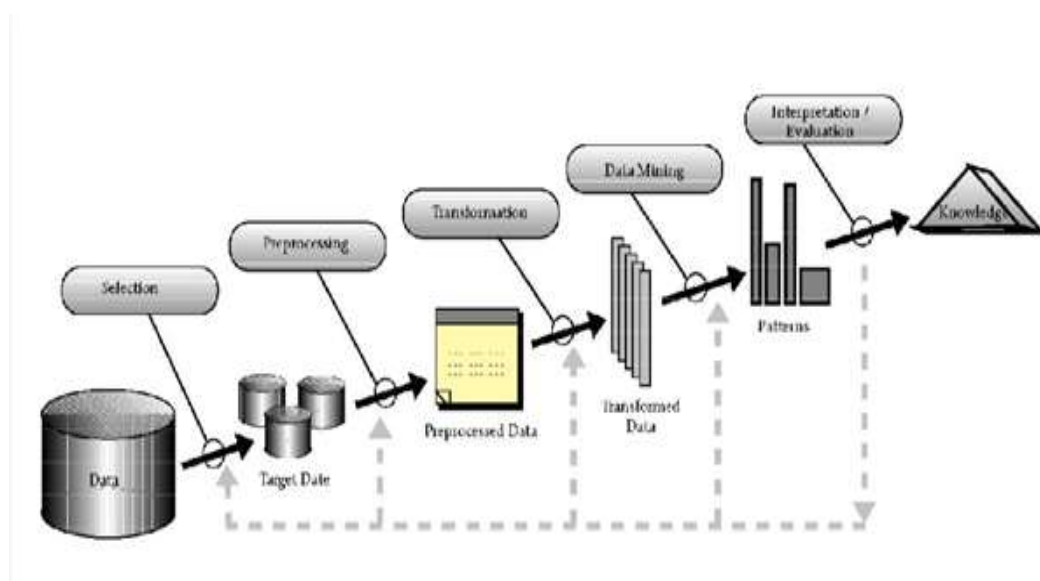
kekurangan algoritma *Apriori* adalah waktu komputasi bisa tinggi untuk dataset yang sangat besar dan memerlukan banyak perhitungan untuk itemset yang lebih besar (kebanyakan kombinasi itemset). Algoritma ini menggunakan prinsip "*Apriorite*" yaitu "Jika suatu itemset sering muncul dalam transaksi, maka semua subset itemset tersebut harus sering muncul juga." Contoh aturan asosiasi adalah "Jika X terjadi, maka Y akan terjadi". Selain *apriori* adapun metode *FP-Growth* (*Frequent Pattern Growth*) yang merupakan alternatif yang lebih efisien dari algoritma *Apriori* untuk menemukan pola yang sering terjadi dalam dataset. *FP-Growth* adalah algoritma yang lebih efisien dibandingkan *Apriori*, terutama pada dataset yang besar. *FP-Growth* tidak membutuhkan pembuatan banyak kandidat item set, seperti yang dilakukan oleh *Apriori*. Sebaliknya, dapat membangun struktur data yang disebut *FP-tree* (*Frequent Pattern Tree*) untuk mewakili itemset yang sering muncul. Kelebihan *FP-Growth* adalah Lebih efisien dibandingkan dengan *Apriori*, terutama untuk dataset besar serta Tidak memerlukan banyak perhitungan kandidat item set sedangkan kekurangan dari *FP-Growth* adalah Lebih kompleks dalam implementasi dibandingkan dengan *Apriori*.

6. Perangkuman (*summarization*)

Perangkuman (*summarization*) menyediakan repretasi data yang lebih sederhana seperti visualisasi serta pembuatan laporan. Teknik Perangkuman (*Summarization*) dapat dibagi menjadi dua pendekatan utama seperti perangkuman ekstraktif (*extractive summarization*)

Perangkuman ekstraktif bertujuan untuk memilih kalimat-kalimat penting atau bagian-bagian dari teks asli dan menggabungkannya untuk membentuk ringkasan. Dapat disimpulkan kalimat dalam ringkasan diambil langsung dari teks asli tanpa ada perubahan dalam strukturnya dan perangkuman abstraktif (*abstractive summarization*) merupakan perangkuman abstraktif yang berfokus pada pembuatan ringkasan baru yang lebih singkat dengan menghasilkan kalimat-kalimat baru yang merangkum inti dari teks asli. Ini merupakan model memahami teks secara lebih mendalam, kemudian menyusun kembali informasi tersebut menggunakan kata-kata dan struktur kalimat yang berbeda. Masing-masing teknik ini memiliki tahapan yang berbeda dalam merangkum teks, dan keduanya dapat diterapkan dalam berbagai konteks, baik itu untuk artikel berita, penelitian, atau dokumen panjang lainnya.

Proses tahapan yang dilalui data *mining* adalah:



Gambar 2. 6 Proses Data Mining

1. Tahapan Seleksi Data

Pada tahapan seleksi data digunakan untuk memilih dan menyeleksi data dari banyaknya data. Data yang diseleksi di proses menggunakan data mining (Randi Rian Putra¹, 2018).

2. Tahapan Pembersian Data

Tahapan pembersian data dilakukan untuk menghindari duplikat data, memeriksa data dan data yang tidak sesuai serta tidak memiliki nilai yang utuh.

3. Tahapan Transformasi Data

Pada tahapan ini dilakukan pembersian data yang dimana data tersebut dibersihkan dan dirubah menjadi data untuk di proses pada data mining (Kamil & Cholil, 2020).

4. Tahapan Data Mining

Pada tahapan ini data mining merupakan proses utama penerapan metode untuk menemukan pengetahuan tersembunyi dalam data, kemudian menjadi informasi serta pada tahapan ini menghasilkan hubungan, pola maupun tren baru pada data (Rahmi & Mikola, 2021).

5. Knowledge

Tahapan ini data yang telah diproses dipresentasikan untuk memperoleh informasi berdasarkan metode yang telah diterapkan. Pada tahapan ini keputusan dirumuskan berdasarkan hasil analisis yang telah ditemukan dan penyajian pada tahapan ini berbentuk informasi yang mudah dipahami.

2.6 Naïve Bayes

Menurut Thomas Bayes, *Naïve Bayes* adalah pengklasifikasian statistik dan probabilitas untuk memprediksi peluang berdasarkan dari pengalaman yang telah terjadi. *Naïve Bayes* adalah algoritma yang merupakan salah satu metode machine learning yang menggunakan perhitungan probabilitas (Sinaga et al., 2022). Algoritma *Naïve Bayes* dapat diterapkan pada data yang berskala ordinal dan *Naïve bayes* juga banyak digunakan dalam menganalisis data yang besar. penggunaan *Naïve Bayes* memiliki struktur algoritma yang sederhana dan cepat (Chen et al., 2021). *Naïve Bayes* sering kali digunakan dalam proses klasifikasi data dengan cara pengelompokan data berdasarkan model yang dimiliki oleh objek klasifikasi (Rinanda et al., 2022). *Naïve Bayes* memerlukan sedikit data latih untuk memastikan parameter yang akan digunakan pada klasifikasi, baik dari variabel independen yang dipertimbangkan maupun hanya dari variable kelas yang diperlukan untuk dapat menentukan kategorisasinya yang bukan dari seluruh variabel kelas (Normah et al., 2022).

Jenis-jenis *Naïve Bayes*:

1. Multinomial *Naïve Bayes*

Multinomial Naïve Bayes digunakan terutama untuk data kategori (misalnya, kata-kata dalam teks untuk analisis sentimen atau spam). Dapat menghitung frekuensi kemunculan fitur (kata-kata) dalam setiap kelas. Banyak digunakan untuk aplikasi *Natural Language Processing* (NLP).

2. Gaussian Naive Bayes

Gaussian Naive Bayes digunakan untuk data kontinu, dengan asumsi bahwa setiap fitur mengikuti distribusi normal (*Gaussian*) misalnya, data numerik seperti tinggi badan, berat badan, atau skor ujian.

3. Bernoulli Naive Bayes

Bernoulli Naive Bayes digunakan untuk data biner, yaitu data dengan dua kategori (misalnya, ada/tidak ada, 1/0). serta biasanya digunakan dalam aplikasi di mana fitur hanya ada atau tidak ada (misalnya, kehadiran kata tertentu dalam sebuah dokumen).

Kelebihan dari metode *Naïve Bayes* sebagai berikut:

1. Menggunakan data *training* dengan jumlah data sedikit dapat memperkirakan parameter (rata-rata serta variasi yang terdapat pada variabel) untuk dilakukan klasifikasi.
2. Skala nilai yang hilang disebabkan oleh pengabaian kejadian selama estimasi memungkinkan terjadinya peluang.
3. Kuat terhadap karakteristik yang tidak relevan.

Kekurangan dari metode *Naïve Bayes* sebagai berikut:

1. Jika probabilitas bersyarat memiliki nilai 0, maka probabilitas prediksi dapat menjadi 0 juga.
2. Memperikarakan bahwa karakteristiknya bebas.

Persamaan *teorema Bayes*:

$$P(H|X) = \frac{P(X|H)P(H)}{P(Y)} \dots \dots \dots (2.1)$$

Dimana,

X = Data dengan kelas yang belum diketahui

H = Hipotesis data X adalah suatu kelas spesifik

$P(H|X)$ = Probabilitas hipotesis data H berdasarkan kondisi X

$P(H)$ = Probabilitas hipotesis H

$P(X|H)$ = Probabilitas X berdasarkan kondisi pada hipotesis H

$P(X)$ = Probabilitas X

Penjabaran lebih lanjut rumus *Bayes* dilakukan dengan menjabarkan $(C|X_1.X_n)$ memakai aturan perkalian salaku berikut:

$$\begin{aligned}
 P(C|x_1, \dots, x_n) &= P(C) P(x_1, \dots, x_n|C) \\
 &= P(C)P(X_1|C)P(X_2, \dots, X_n|C, X_1) \\
 &= (C)P(X_1|C)P(X_2|C, X_1)P(X_3, \dots, X_n|C, X_1, X_2) \\
 &= (C)P(X_1|C)P(X_2|C, X_1) \\
 &\quad P(X_3|C, X_1, X_2)P(X_4, \dots, X_n|C, X_1, X_2, X_3) \\
 &= P(C)P(X_1|C)P(X_2|C, X_1)P(X_3|C, X_1, X_2) \\
 &\quad P(X_n|C, X_1, X_2, X_3, \dots, X_{n-1} \dots \dots \dots \dots \dots (2.2)
 \end{aligned}$$

Dapat disimpulkan bahwa banyak faktor-faktor yang kompleks mempengaruhi nilai probabilitas sehingga mustahil untuk menghitung nilai tersebut satu persatu. Dampaknya perhitungan semakin sulit untuk dilakukan, sehingga digunakan asumsi independensi yang sangat besar maupun tinggi, untuk masing-masing atribut agar saling bebas. Dengan asumsi tersebut, dibutuhkan persamaan (3).

$$P(X_i | X_j) = \frac{P(X_i \cap X_j)P(H)}{P(X_j)} = \frac{P(X_i)P(X_j)}{P(X_j)} = P(X_i)$$

Untuk $i \neq j$, sehingga:

$$P(X_i|C, X_j) = P(X_i|C) \dots \dots \dots \dots \dots (2.3)$$

Dapat disimpulkan bahwa dari persamaan di atas asumsi independensi membuat ketentuan perhitungan menjadi sederhana. Selanjutnya penjabaran

$(C|X_1, \dots, X_n)$ dapat disederhanakan menjadi persamaan (2.4):

$$P(X_2|C)P(X_3|C)\dots(C|X_1, \dots, X_n) = P(X_i|C) = \prod_{i=1}^n P(X_i|C) \dots \dots \dots (2.4)$$

$\prod_{i=1}^n P(X_i|C)$ = Perkalian ranting antar atribut Persamaan (2.4) ialah teorema Bayes yang digunakan untuk melakukan perhitungan klasifikasi. Klasifikasi data *continue* ataupun data angka menggunakan rumus distribusi *Gaussian* dengan 2 parameter yaitu mean μ dan varian σ :

$$P(X_i = X_i|C = C) = P(X_i|C_j) = \frac{1}{\sigma\sqrt{2\pi}} \exp \frac{(X_i - \mu_{ij})^2}{2\sigma^2_{ij}} \dots \dots \dots (2.5)$$

Keterangan:

P: Peluang

X_i : Atribut ke i

X_j : Nilai atribut ke i

C: Kelas yang dicari

C_i : Sub kelas Y yang dicari

μ : Menyatakan rata-rata dari seluruh atribut

σ : Deviasi standar, menyatakan varian dari seluruh atribut

Metode *Naïve Bayes* membutuhkan data latih dan data uji untuk melakukan klasifikasi, semakin banyak data yang digunakan maka semakin baik hasil prediksi yang didapatkan. $P(C_i)$ merupakan probabilitas untuk setiap sub kelas c yang dapat dihasilkan menggunakan persamaan:

$$P(c_i) = \frac{S_i}{S} \dots \dots \dots (2.6)$$

Si merupakan jumlah data latih atau training dari kategori C_i , dan S adalah jumlah total data *training*. Menghitung $P(X_i|C_i)$ merupakan probabilitas posterior X_i dengan syarat C menggunakan persamaan.

2.7 Penelitian Terkait

Penelitian terkait kepuasan pengguna aplikasi banyak digunakan dalam menentukan apakah aplikasi tersebut dapat digunakan dengan baik dan bermanfaat bagi pengguna aplikasi walaupun dalam penelitian menggunakan topik yang berbeda-beda.

Berikut contoh penelitian terkait analisis kepuasan pengguna:

1. Analisis Tingkat Kepuasan Pengguna Aplikasi Ojek Online Dengan Metode *Naïve Bayes*, (Dari & Elen Tania Hanayah, 2023). Memiliki kelebihan pada jurnal “Analisis Tingkat Kepuasan Pengguna Aplikasi Ojek Online Dengan Metode *Naïve Bayes*” yaitu pada penggunaan metode *Naïve Bayes* dalam penelitian ini memiliki empat variabel dengan hasil akhir yang memuaskan serta memiliki kekurangan berupa data yang digunakan pada penelitian untuk mengklasifikasi menggunakan metode *Naïve Bayes* terlalu sedikit serta nilai akursi pada penelitian ini tidak tersedia.
2. Klasifikasi Tingkat Kepuasan Pengguna *Google Classroom* Dalam Pembelajaran *Online* Menggunakan Algoritma *Naïve Bayes*, (Ramadhani et al., 2024). Memiliki kelebihan pada jurnal “Klasifikasi Tingkat Kepuasan Pengguna *Google Classroom* Dalam Pembelajaran *Online* Menggunakan Algoritma *Naïve Bayes*” yaitu penelitian ini menggunakan 500 responden dengan jumlah data latih sekitar 400 dan

100 untuk data uji. Akurasi data sebesar 89% dengan klasifikasi 88 responden dan 12 pengguna tidak puas, sehingga dapat disimpulkan bahwa *Naïve Bayes* dalam penelitian ini memiliki kinerja yang baik dalam menentukan kepuasan pengguna serta memiliki kekurangan berupa pada penelitian sebelumnya metode *Naïve Bayes* kurang optimal dalam menghadapi ketidakseimbangan akibat dari keterbatasan penelitian pada kelas ketidakpuasan pengguna.

3. Analisis data kepuasan pengguna layanan *E-wallet* Gopay Menggunakan *Naïve Bayes Classifier Algorithm*, (Sudipa et al., 2023). Memiliki kelebihan pada jurnal “Analisis data kepuasan pengguna layanan *E-wallet* Gopay Menggunakan *Naïve Bayes Classifier Algorithm*” yaitu terdapat 100 responden dan 79% pengguna *E-wallet* Gopay merasa puas dan 21% lainnya merasa tidak puas. Dapat disimpulkan bahwa penggunaan metode *Naïve bayes* pada penelitian ini memiliki performa serta memiliki kekurangan seperti pada penelitian ini terdapat 100 responden dan diantaranya 91 puas dan 9 tidak puas. Tetapi nilai akurasi 79% dimana data yang digunakan terlalu sedikit dalam klasifikasi sehingga kepuasan tidak mencapai 80% - 90%.

BAB III

METODOLOGI PENELITIAN

3.1 Obyek Penelitian

Obyek pada penelitian ini adalah mahasiswa Universitas Muhammadiyah Sumatera Utara yang menggunakan aplikasi *UMSU Academy*. Fokus utama adalah mengumpulkan data dari para responden mengenai berbagai aspek yang berkaitan dengan kepuasan pengguna pada aplikasi *UMSU Academy*.

3.1.1 Lokasi Penelitian

Penelitian ini dilaksanakan pada Universitas Muhammadiyah Sumatera Utara, Kecamatan Medan Kota, Kota Medan, Sumatera Utara. Pada penelitian ini dilakukan penyebaran kuesioner kepada mahasiswa Universitas Muhammadiyah Sumatera Utara dengan lokasi dan rentan waktu pengumpulan pada kuesioner ini dipilih dengan memperhatikan aksesibilitas responden serta kebutuhan dalam memperoleh representasi yang baik untuk populasi yang diteliti.

3.1.2 Atribut atau Variabel Penelitian

Atribut ataupun variabel pada penelitian ini meliputi beberapa faktor yang mempengaruhi kepuasan pengguna aplikasi *UMSU Academy*. Atribut atau variabel yang digunakan seperti Jenis Kelamin, Jurusan, *Content* (isi aplikasi), *Format* (tampilan aplikasi), *Accuracy* (keakuratan aplikasi), *Timelines* (ketepatan waktu), *Ease of use* (kemudahan pengguna aplikasi).

3.1.3 Kelas (Label)

Kelas atau label dalam penelitian ini adalah tingkat kepuasan pengguna

aplikasi UMSU Academy, dibagi menjadi beberapa kategori seperti sangat puas, puas, cukup puas, tidak puas dan sangat tidak puas.

3.2 Alat dan Bahan

Alat dan bahan yang digunakan untuk menganalisis data dalam penelitian ini yaitu:

3.2.1 Kuesioner

Kuesioner dalam penelitian ini digunakan sebagai alat utama dalam pengumpulan data melalui google formulir, serta kuesioner ini berisi mengenai pertanyaan terkait dengan kepuasan pengguna aplikasi UMSU Academy.

Nilai numerik pada kuesioner mewakili keterangan pada Tabel 3.1.

Tabel 3. 1 Skala Dan Bobot.

Keterangan	Bobot
Sangat tidak setuju	1
Tidak setuju	2
Netral	3
Setuju	4
Sangat setuju	5

3.2.2 Microsoft Excel

Microsoft excel pada penelitian ini berguna untuk mengatur dan membersihkan data kepuasan pengguna secara efisien, baik itu pemformatan data, penghapusan duplikat data, dan pengelompokan data berdasarkan kriteria tertentu.

3.2.3 Kebutuhan Perangkat Keras

Pada penelitian ini digunakan perangkat keras untuk mempermudah penelitian dan untuk merancang kuesioner, menganalisis data, serta menyusun laporan. Perangkat keras yang digunakan seperti laptop. Adapun pada penelitian

ini menggunakan perangkat lunak seperti Microsoft windows 10, Visual Studi Code dan Python.

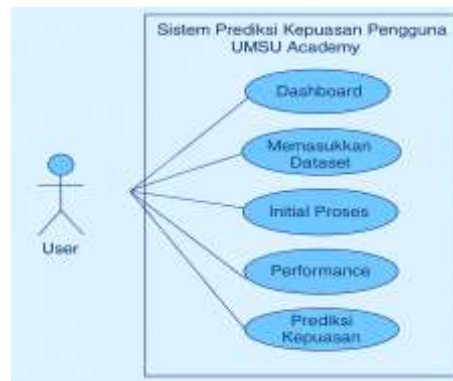
3.2.4 Literatur dan Referensi

Literatur dan referensi yang digunakan dalam penelitian ini berfungsi sebagai analisis data, perancangan penelitian serta interpretasi hasil. Literatur dan referensi yang digunakan bisa berupa buku, jurnal ilmiah, artikel, serta dokumen – dokumen yang berkaitan dengan kepuasan pengguna aplikasi terkhusus mengenai aplikasi UMSU *Academy*.

3.3 Tahap Desain

3.3.1 Use Case Diagram

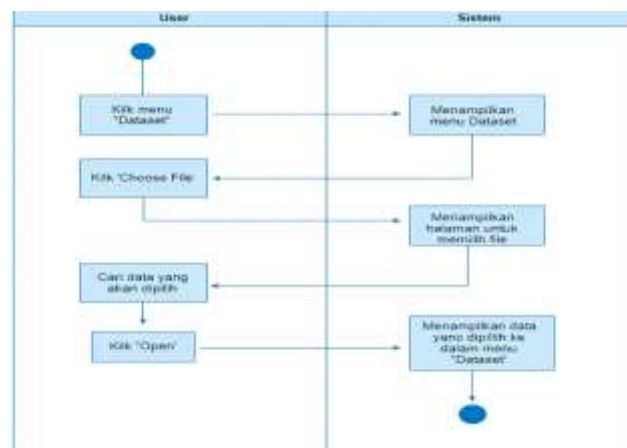
Use case diagram merupakan suatu bentuk diagram yang menggambarkan fungsionalitas yang diharapkan dari sebuah *system* yang dilihat dari prespektif pengguna di luar *system* yang dibuat. *Use case diagram* juga digunakan untuk mempresentasikan interaksi yang terjadi antara actor atau user dengan proses *system* yang dibuat. Diagram ini sangat berguna pada tahap awal pengembangan sistem karena dapat membantu pengembang dan pemangku kepentingan lainnya untuk memahami kebutuhan fungsional sistem secara visual dan intuitif. *Use case diagram* juga menjadi dasar dalam menyusun dokumen analisis kebutuhan serta mendesain sistem yang sesuai dengan harapan pengguna.



Gambar 3. 1 Use Case Diagram

3.3.2 Activity Diagram

Activity diagram menggambarkan aliran kerja atau aktivitas dari sebuah sistem, tetapi bukan aktivitas aktor. *Activity diagram* juga menggambarkan bagaimana alur sistem berawal, pilihan (decision) yang mungkin terjadi, dan bagaimana akhir alur sistem tersebut.



Gambar 3. 2 Activity Diagram

3.3.3 Desain Interface

Desain Interface ditujukan pada penataan elemen-elemen yang akan ditampilkan kepada pengguna pada layar, termasuk bagaimana pengguna

berinteraksi dengan sistem melalui tombol, *input field*, menu, ikon, dan sebagainya.



Gambar 3. 3 Desain *Interface*

3.4 Algoritma Naïve Bayes

Penelitian ini menggunakan metode algoritma *Naïve Bayes* untuk menganalisis data yang telah dikumpulkan. Algoritma *Naïve Bayes* digunakan untuk memprediksi kepuasan responden berdasarkan atribut-atribut yang ada. *Naïve Bayes* merupakan pengklasifikasian dengan bentuk model probabilistik dan statistik yang disederhanakan berdasarkan pada *Teorema Bayes* dengan asumsi bahwa setiap atribut bersifat *Independent* (bebas). Dalam penelitian dapat melatih model *Naïve Bayes* menggunakan data yang ada dan kemudian menguji model untuk memprediksi apakah pengguna merasa sangat tidak puas, tidak puas, cukup puas, puas serta sangat puas dengan aplikasi *UMSU Academy* berdasarkan atribut-atribut yang telah diberikan.

Adapun rumus metode *Naïve Bayes* yang digunakan dalam penelitian ini:

Menghitung probabilitas prior dari setiap kelas dalam data latih:

$$P(C_k) = \frac{\text{jumlah sampel dalam kelas } C_k}{\text{total jumlah sampel}} \dots \dots \dots (3.1)$$

Keterangan:

$P(C_k)$ = Probabilitas prior dari kelas C_k

Menghitung likelihood, yaitu menghitung probabilitas kondisi dari setiap fitur diberikan kelas tertentu.

$$P = (X_i|C_k) \dots \dots \dots (3.2)$$

Menghitung probabilitas posterior, menggunakan *teorema bayes* untuk menghitung probabilitas posterior untuk setiap kelas diberikan instance dari fitur.

$$P(C_k|X) = \frac{P(C_k) \cdot \prod_{i=1}^n P(X_i|C_k)}{p(X)} \dots \dots \dots (3.3)$$

Keterangan:

$P(C_k|X)$: Probabilitas posterior dari kelas C_k diberikan fitur X

$P(C_k)$: Probabilitas prior dari kelas C_k

$P(X_i|C_k)$: Likelihood dari fitur X_i diberikan kelas C_k

$P(X)$: Probabilitas total dari fitur X

Untuk menghitung nilai akhir kelas menggunakan rumus:

$$C_{MAP} = \operatorname{argmax}_{C \in C} P(X|C) \dots \dots \dots (3.4)$$

Keterangan:

C_{MAP} : Hipotesa nilai tertinggi

$\operatorname{argmax}_{C \in C}$: Nilai rata-rata dari setiap kelas

Bayesian Classification terbukti memiliki tingkat akurasi dan kecepatan yang tinggi. Berikut merupakan *Teorema Bayesian* dalam bentuk umum:

$$p(A|B) = \frac{(p(B|A) \cdot p(A))}{p(B)} \dots \dots \dots (3.5)$$

Keterangan:

B : Data dengan kelas yang belum diketahui

A : Hipotesa data B merupakan suatu kelas spesifik

$P(A)$: Probabilitas hipotesa A (prior probability/probabilitas awal)

$P(B)$: Probabilitas B

$P(B|A)$: Probabilitas hipotesa B berdasarkan kondisi A

$P(A|B)$: Probabilitas hipotesa A berdasarkan kondisi B (Posterior probability/probabilitas akhir)

3.4.1 Rumus Pencarian Tingkat Akurasi

Akurasi adalah seberapa jauh nilai sebenarnya yang didapatkan dan prediksi berbeda. Untuk mengukur akurasi model, digunakan *matrix confusion* yang menitik beratkan pada kelasnya. Sebuah *array* yang digunakan untuk mencatat hasil kerja klasifikasi disebut confusion matrix. Pada langkah ini, *matrix confusion* menguji data dan menggunakan model untuk menemukan tingkat akurasi terbaik. *Confusion Matrix* mengacu pada Tabel 3.2.

Tabel 3. 2 *Confusion Matrix*

Klasifikasi Naïve Bayes	Puas (+)	Tidak Puas (-)
Puas (+)	<i>True Positive (TP)</i>	<i>False Negative (FN)</i>
Tidak Puas (-)	<i>False Positive (FP)</i>	<i>True Positive (TP)</i>

Rumus confusion matrix yang digunakan adalah:

$$Precision = \frac{TP}{TP + FP} \dots \dots \dots (3.6)$$

$$Recall = \frac{TP}{TP + FN} \dots \dots \dots (3.7)$$

$$F1 - Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \dots \dots \dots (3.8)$$

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \dots \dots \dots (3.9)$$

Keterangan:

1. *Precision* mengukur ketepatan prediksi positif.
2. *Recall* mengukur kemampuan model menangkap semua data positif.
3. *F1-Score* menghitung rata-rata harmonis antara *precision* dan *recall*.
4. *Accuracy* adalah jumlah prediksi yang benar yang menyatakan puas.
5. TP (*True positive*) adalah jumlah pengguna yang puas dan diprediksi puas oleh model.
6. TN (*True negative*) adalah jumlah pengguna yang tidak puas tetapi diprediksi puas oleh model.
7. FP (*False positive*) adalah jumlah pengguna yang puas tetapi diprediksi tidak puas oleh model.
8. FN (*False negative*) adalah jumlah pelanggan yang tidak puas dan diprediksi tidak puas oleh model.

3.5 Jadwal Penelitian

Tabel 3. 3 Jadwal Penelitian

No	Kegiatan	Bulan (Tahun 2024 – 2025)				
		Desember	Januari	Februari	Maret	April Mei
1.	Persiapan Penelitian					
	1. Pengajuan Judul					
	2. Pembuatan Proposal					
2.	Tahapan Pelaksanaan					
	1. Pengumpulan Data					
	2. Pengelolaan Data					
	3. Penyusunan Skripsi					

BAB IV

HASIL DAN PEMBAHASAN

4.1 Hasil Dan Pembahasan

Analisis kepuasan pengguna terhadap aplikasi UMSU *Academy* menggunakan metode *Naïve Bayes* dibangun berdasarkan hasil analisis kebutuhan dan perancangan sistem yang telah dijelaskan pada Bab III. Sistem ini dirancang sebagai aplikasi berbasis web yang memungkinkan pengguna untuk memproses data kuesioner, melatih model klasifikasi, melakukan evaluasi performa model, serta melakukan prediksi terhadap tingkat kepuasan berdasarkan masukan data baru. Hasil implementasi sistem ditampilkan dalam lima menu utama, yaitu: *Dashboard*, *Dataset*, *Initial Process*, *Performance*, dan *Predict*. Menu *Dashboard* menampilkan halaman awal dan identitas aplikasi. Menu *Dataset* memungkinkan pengguna untuk mengunggah file kuesioner dalam format .csv atau .xlsx dan melihat isi data secara langsung dalam bentuk tabel. Menu *Initial Process* digunakan untuk melakukan pra-pemrosesan awal sebelum data dilatih dengan algoritma klasifikasi mencakup pembersihan data seperti menghapus baris kosong, duplikat, dan karakter tidak valid.

Menu *Performance* merupakan bagian penting dari sistem yang menampilkan proses pelatihan model *Naïve Bayes* dan hasil evaluasi performa model berdasarkan data uji. Pengguna dapat mengatur parameter seperti persentase data latih dan nilai *random state* untuk menjaga konsistensi hasil. Setelah model dilatih, sistem akan menampilkan hasil evaluasi dalam bentuk akurasi, *precision*, *recall*, dan *F1-score*. Pada menu *Predict*, pengguna dapat melakukan prediksi secara manual dengan mengisi 25 nilai indikator kepuasan,

mulai dari C1 hingga E5. Sistem kemudian menghitung skor rata-rata, jumlah total nilai, dan memproses prediksi kelas kepuasan menggunakan model yang telah dilatih. Hasil prediksi ditampilkan secara langsung dalam bentuk label seperti “Sangat Puas”, “Puas”, dan sebagainya. Sebagai validasi sistem, dilakukan pula perhitungan manual terhadap salah satu baris data uji. Perhitungan dilakukan dengan menghitung prior probability setiap kelas, menghitung likelihood dari setiap nilai indikator, lalu dikalikan untuk mendapatkan nilai posterior. Setelah dilakukan normalisasi, kelas dengan nilai posterior tertinggi ditetapkan sebagai hasil prediksi.

Sistem atau program ini dapat di dijalankan dengan terlebih dahulu ketika membuka atau menjalankan program seperti gambar 4.1



Gambar 4. 1 Tampilan Utama Sistem

4.2 Implementasi Sistem Naïve Bayes

4.2.1 Tampilan Dashboard

Pada halaman *dashboard* berfungsi sebagai beranda utama atau *landing page*, tempat pengguna pertama kali masuk dan melihat tampilan umum sistem.



Gambar 4. 2 Tampilan Halaman *Dashboard*

4.2.2 Tampilan Menu Dataset

Menu dataset merupakan salah satu fitur utama yang terdapat pada aplikasi *Analisis Kepuasan Pengguna Aplikasi UMSU Academy*. Fitur ini berfungsi sebagai tempat awal untuk memuat data kuesioner yang akan dianalisis menggunakan metode *Naive Bayes*. Melalui halaman ini, pengguna diberikan kemudahan untuk mengunggah dataset dalam format .csv atau .xlsx, yang berisi jawaban responden terhadap sejumlah pernyataan pada skala Likert. Data yang diunggah umumnya terdiri dari beberapa indikator, seperti C1 hingga E5, serta kolom target berupa kelas kepuasan pengguna, misalnya "Sangat Puas", "Puas", "Cukup Puas", "Tidak Puas", dan "Sangat Tidak Puas".

Setelah proses unggah berhasil dilakukan, sistem akan secara otomatis menampilkan isi dataset dalam bentuk tabel interaktif. Tabel ini memudahkan pengguna untuk melakukan verifikasi data secara langsung, seperti memastikan kesesuaian nama kolom, kelengkapan data, serta nilai-nilai yang diinput oleh responden. Menu ini sangat penting karena berfungsi sebagai dasar dari seluruh

proses pengolahan data selanjutnya, mulai dari proses pra-pemrosesan (initial process), pelatihan model, hingga evaluasi performa model klasifikasi.



Gambar 4. 3 Tampilan Halaman Menu *Dataset*

4.2.3 Tampilan Menu Initial Process

Menu *Initial Process* digunakan untuk melakukan tahap pra-pemrosesan terhadap data yang telah diunggah sebelumnya melalui menu dataset. Tahapan ini meliputi proses konversi data kategorikal menjadi numerik, pengecekan kelengkapan data. Pra-pemrosesan data merupakan langkah penting dalam penerapan algoritma *Naive Bayes*, karena algoritma ini membutuhkan data dalam format numerik dan bebas dari nilai kosong atau anomali. Dengan adanya menu *Initial Process*, pengguna dapat memastikan bahwa data telah siap untuk digunakan dalam proses pelatihan model klasifikasi, serta dapat melihat hasil pembagian data secara langsung melalui tampilan sistem. Menu ini membantu meminimalkan kesalahan pemrosesan data dan menjamin kualitas model yang dihasilkan.

Gambar 4. 4 Tampilan Menu *Initial Process*

4.2.4 Tampilan Menu Performance

Menu *Performance* berfungsi untuk menampilkan proses pelatihan (*training*) model klasifikasi *Naive Bayes* serta melakukan evaluasi terhadap performa model tersebut. Pada halaman ini, pengguna dapat menentukan parameter penting yang memengaruhi proses pelatihan model, yaitu persentase data pelatihan dan nilai random state. Gambar tersebut menunjukkan bahwa pengguna dapat memasukkan nilai persentase pelatihan (dalam kisaran 1–100) melalui sebuah kotak input. Sebagai contoh, nilai “80” menunjukkan bahwa sebanyak 80% data akan digunakan untuk pelatihan model, sementara sisanya (20%) akan digunakan untuk pengujian. Selain itu, terdapat pula kolom input untuk memasukkan nilai *random state*, yang berfungsi sebagai pengontrol pembagian data agar hasil pembagian dapat direproduksi kembali secara konsisten. Dalam gambar, nilai random state yang digunakan adalah “42”, yang

merupakan nilai default yang umum digunakan dalam praktik pembelajaran mesin.

Analisis Kepuasan Pengguna Aplikasi UMSU Academy

Dashboard Dataset Initial Process Performance Predict

Performance

Persentase Training (1-100):

80

Random State:

42

Tampilkan Hasil

Uyuhui Model

Data Training:

	C1	C2	C3	C4	C5	F1	F2	F3	F4	F5	A1	A2	A3	A4	A5	T1	T2	T3	T4	T5	E1	E2	E3	E4	E5	Jumlah	Rata-Rata	Class Kepuasan
6	3	3	4	4	3	3	4	4	3	3	4	4	2	4	3	3	4	3	2	4	2	2	3	5	92	3.82	Puas	
977	4	4	5	5	4	4	3	4	5	5	5	2	3	5	3	4	4	4	5	4	4	4	5	4	105	4.20	Sangat Puas	
773	2	4	3	4	5	4	4	4	2	3	5	4	1	3	4	4	4	3	5	4	4	4	2	5	2	89	3.56	Puas
853	4	4	5	5	4	4	3	4	5	5	5	2	3	5	3	4	4	4	5	4	4	4	5	4	103	4.20	Sangat Puas	
660	4	4	4	4	3	3	4	3	4	3	2	3	5	3	4	4	2	4	4	4	5	3	2	5	3	89	3.58	Puas
332	2	1	1	1	1	1	1	1	1	1	1	1	2	1	1	2	1	1	1	1	2	1	2	1	2	31	1.24	Sangat Tidak Puas
528	4	4	5	5	4	4	3	4	5	5	5	2	5	5	3	4	4	4	5	4	4	4	5	4	105	4.20	Sangat Puas	
861	1	1	1	1	1	1	1	1	1	1	2	1	1	1	1	2	2	1	1	1	1	2	1	1	1	29	1.18	Sangat Tidak Puas
215	4	4	4	4	4	4	5	3	5	2	1	5	5	4	4	5	5	5	4	2	4	4	5	4	4	94	3.76	Puas
968	4	3	4	4	2	5	3	4	4	3	3	3	4	4	2	2	2	4	4	2	4	2	5	5	4	86	3.44	Puas
495	1	1	2	1	1	1	2	2	1	2	1	1	3	2	3	1	3	2	2	2	1	3	2	1	1	43	1.72	Tidak Puas

Gambar 4. 5 Tampilan Menu *performance Persentase Training*

Setelah parameter dimasukkan, pengguna dapat mengklik tombol “Tampilkan Hasil” untuk menampilkan data yang telah dibagi menjadi data pelatihan yaitu data *training* dan data *testing*.

Data Training:

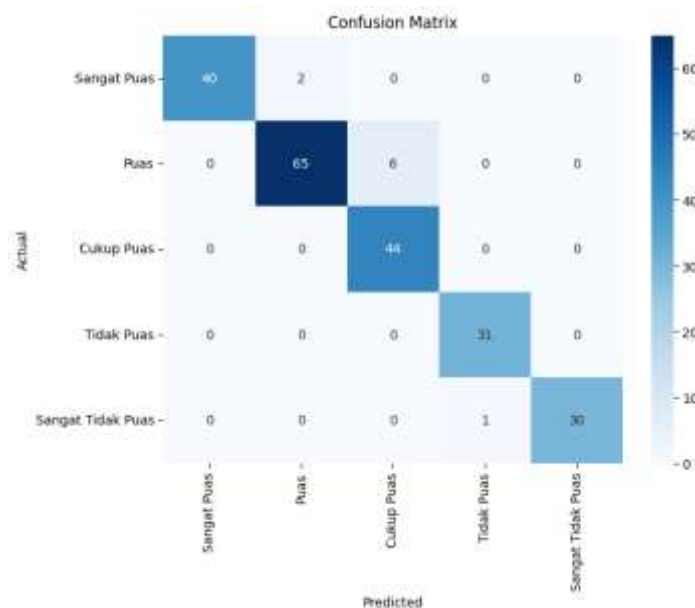
	C1	C2	C3	C4	C5	F1	F2	F3	F4	F5	A1	A2	A3	A4	A5	T1	T2	T3	T4	T5	E1	E2	E3	E4	E5	Jumlah	Rata-Rata	Class Kepuasan
6	3	3	4	4	3	3	4	4	3	3	3	4	4	2	4	3	3	4	3	2	4	2	2	3	5	92	3.68	Puas
977	4	4	5	5	4	4	3	4	5	5	5	2	5	5	3	4	4	4	5	4	4	4	4	5	4	105	4.20	Sangat Puas
773	2	4	3	4	5	4	4	4	2	3	5	4	1	3	4	4	4	3	5	4	4	4	2	5	2	89	3.56	Puas
853	4	4	5	5	4	4	3	4	5	5	5	2	5	5	3	4	4	4	5	4	4	4	4	5	4	103	4.20	Sangat Puas
660	4	4	4	4	3	3	4	3	4	3	2	3	5	3	4	4	2	4	4	4	5	3	2	5	3	89	3.58	Puas
332	2	1	1	1	1	1	1	1	1	1	1	1	2	1	1	2	1	1	1	1	2	1	2	1	2	31	1.24	Sangat Tidak Puas
528	4	4	5	5	4	4	3	4	5	5	5	2	5	5	3	4	4	4	5	4	4	4	4	5	4	105	4.20	Sangat Puas
861	1	1	1	1	1	1	1	1	1	1	2	1	1	1	1	2	2	1	1	1	1	2	1	1	1	29	1.18	Sangat Tidak Puas
215	4	4	4	4	4	4	5	3	5	2	1	5	5	4	4	5	5	5	4	2	4	4	3	4	4	94	3.76	Puas
968	4	3	4	4	2	5	3	4	4	3	3	3	4	4	2	2	2	4	4	2	4	2	5	5	4	86	3.44	Puas
495	1	1	2	1	1	1	2	2	1	2	2	1	3	2	3	1	3	2	2	2	1	3	2	1	1	43	1.72	Tidak Puas

Gambar 4. 6 Data *Training*

	C1	C2	C3	C4	C5	F1	F2	F3	F4	F5	A1	A2	A3	A4	A5	T1	T2	T3	T4	T5	E1	E2	E3	E4	E5	Jumlah	Rata-Rata	Class	Keputusan	
487	4	4	3	4	4	2	2	4	3	3	4	4	3	4	3	1	3	4	3	1	4	4	3	4	1	79	3.09	Cukup	Puas	
130	3	3	3	3	3	4	3	4	3	4	3	3	3	1	3	3	3	3	3	3	3	3	4	3	3	87	3.05	Puas		
86	4	3	3	4	4	3	3	4	4	4	4	4	4	1	4	4	4	3	4	1	4	3	4	3	81	3.03	Puas			
101	2	4	4	3	3	4	4	3	3	4	3	3	4	3	4	4	4	4	4	4	4	4	4	2	2	87	3.46	Puas		
794	4	4	3	3	4	4	3	4	3	3	3	2	3	1	3	4	4	4	3	4	4	4	4	3	4	103	4.20	Sangat	Puas	
104	1	4	4	3	4	4	4	3	4	3	4	3	2	1	1	4	4	4	2	3	3	3	4	4	3	81	3.72	Puas		
794	1	3	1	4	4	3	4	4	3	4	3	1	4	1	4	4	2	4	3	1	4	4	4	1	4	88	3.44	Puas		
446	3	4	2	3	4	4	2	4	3	4	1	4	3	4	3	4	3	4	2	3	2	3	4	1	3	61	3.28	Cukup	Puas	
1004	4	4	3	3	4	4	3	4	3	3	3	3	3	3	1	4	4	4	3	4	4	4	4	3	4	103	4.20	Sangat	Puas	
923	1	3	2	3	1	3	1	1	3	1	2	1	1	1	1	1	1	2	1	2	1	3	1	1	1	96	1.20	Sangat	Tidak	Puas

Gambar 4. 7 Data Testing

Kemudian, terdapat tombol tambahan yaitu “Evaluasi Model” yang akan memproses evaluasi model terhadap data uji menggunakan metrik evaluasi seperti akurasi, *precision*, *recall*, *F1-score*, dan *confusion matrix*. Menu *Performance* memungkinkan pengguna untuk mengatur, melatih, serta meninjau ulang performa model klasifikasi sebelum digunakan dalam prediksi nyata. Hal ini sangat penting dalam memastikan bahwa model yang digunakan telah melewati proses evaluasi yang objektif dan valid.

Gambar 4. 8 Tampilan Menu *Performance* Evaluasi Model

Gambar di atas menunjukkan visualisasi confusion matrix dari hasil evaluasi model klasifikasi Naïve Bayes terhadap data uji pada sistem analisis kepuasan pengguna aplikasi UMSU *Academy*. *Confusion matrix* ini digunakan untuk mengevaluasi performa model dalam memetakan hasil prediksi terhadap label aktual dari setiap kelas. Pada sumbu vertikal ditampilkan kelas aktual, sementara sumbu horizontal menunjukkan hasil prediksi model. Setiap sel dalam matriks menggambarkan jumlah data yang diklasifikasikan ke dalam suatu kelas, baik secara benar maupun salah.

Berdasarkan hasil yang ditampilkan, model mampu mengklasifikasikan sebagian besar data dengan cukup baik. Sebanyak 40 data dari kelas “Sangat Puas” berhasil diprediksi dengan benar, terdapat 2 data yang salah diklasifikasikan ke dalam kelas “Puas”. Untuk kelas “Puas”, sebanyak 65 data berhasil diprediksi dengan tepat, dan terdapat 6 data yang salah diklasifikasikan sebagai “Cukup Puas”. Kelas “Cukup Puas” menunjukkan hasil yang sangat baik dengan seluruh 44 data berhasil diklasifikasikan dengan benar tanpa kesalahan. Hal yang sama juga terlihat pada kelas “Tidak Puas”, di mana seluruh 31 data berhasil diprediksi dengan benar. Sementara itu, pada kelas “Sangat Tidak Puas”, sebanyak 30 data diklasifikasikan dengan benar, dan hanya 1 data yang salah diprediksi sebagai “Tidak Puas”.

Secara keseluruhan, visualisasi confusion matrix menunjukkan bahwa model *Naïve Bayes* memiliki performa yang sangat baik dalam mengklasifikasikan data ke dalam lima kategori kepuasan, dengan tingkat akurasi yang tinggi dan kesalahan klasifikasi yang relatif rendah.

Perhitungan Manual Tiap Kelas

```

Perhitungan Manual:
Kelas: Sangat Puas
TP = 40, FP = 0, FN = 2
Precision = 1.0000
Recall = 0.9524
F1-Score = 0.9756

Kelas: Puas
TP = 65, FP = 2, FN = 6
Precision = 0.9701
Recall = 0.9155
F1-Score = 0.9420

Kelas: Cukup Puas
TP = 44, FP = 6, FN = 0
Precision = 0.8800
Recall = 1.0000
F1-Score = 0.9362

Kelas: Tidak Puas
TP = 31, FP = 1, FN = 0
Precision = 0.9688
Recall = 1.0000
F1-Score = 0.9841

Kelas: Sangat Tidak Puas
TP = 30, FP = 0, FN = 1
Precision = 1.0000
Recall = 0.9677
F1-Score = 0.9838

Akurasi Total = 0.9580

```

Gambar 4. 9 Perhitungan Manual *Precision*, *Recall*, *F1-Score*, Akurasi

Gambar di atas menunjukkan perhitungan manual terhadap performa klasifikasi model pada masing-masing kelas untuk menilai sejauh mana akurasi model dalam mengenali label kepuasan pengguna aplikasi UMSU *Academy*. Perhitungan ini didasarkan pada nilai-nilai dasar dari *confusion matrix*, yaitu *True Positive (TP)*, *False Positive (FP)*, dan *False Negative (FN)*, yang kemudian digunakan untuk menghitung nilai *precision*, *recall*, dan *F1-score* pada tiap kelas.

Untuk kelas “Sangat Puas”, model memperoleh nilai $TP = 40$, $FP = 0$, dan $FN = 2$. Nilai *precision* yang dihasilkan adalah 1.000, *recall* sebesar 0.9524, dan *F1-score* sebesar 0.9756. Pada kelas “Puas”, diperoleh $TP = 65$, $FP = 2$, dan $FN = 6$, yang menghasilkan *precision* sebesar 0.9701, *recall* sebesar 0.9155, dan *F1-score* sebesar 0.9420. Untuk kelas “Cukup Puas”, $TP = 44$, $FP = 6$, dan $FN = 0$. Dari nilai ini, diperoleh *precision* sebesar 0.8800, *recall* sebesar 1.0000, dan *F1-score* sebesar 0.9362. Pada kelas “Tidak Puas”, model menghasilkan $TP = 31$, $FP = 1$, dan $FN = 0$, sehingga *precision* sebesar 0.9688, *recall* sebesar 1.0000, dan *F1-score* sebesar 0.9841. Sementara itu, pada kelas “Sangat Tidak Puas”,

diperoleh $TP = 30$, $FP = 0$, dan $FN = 1$, dengan *precision* sebesar 1.0000, *recall* sebesar 0.9677, dan *F1-score* sebesar 0.9836.

Secara keseluruhan, akurasi total model mencapai 0.9589 atau 95,89%, yang menunjukkan bahwa model klasifikasi *Naïve Bayes* memiliki performa yang sangat baik dalam mengenali berbagai tingkat kepuasan pengguna secara konsisten dan akurat.

Classification Report				
	precision	recall	f1-score	support
Sangat Puas	1.000000	0.952381	0.975610	42.000000
Puas	0.970149	0.915493	0.942029	71.000000
Cukup Puas	0.880000	1.000000	0.936170	44.000000
Tidak Puas	0.988750	1.000000	0.984127	31.000000
Sangat Tidak Puas	1.000000	0.967742	0.983607	31.000000
accuracy	0.958904	0.958904	0.958904	0.958904
macro avg	0.963780	0.967123	0.964308	219.000000
weighted avg	0.961789	0.958904	0.959137	219.000000

Gambar 4. 10 Tampilan Classification Report

Gambar di atas menampilkan metrik evaluasi model klasifikasi yang meliputi *precision*, *recall*, *F1-score*, dan *support* untuk masing-masing kelas. Laporan ini memberikan gambaran menyeluruh mengenai performa model terhadap setiap label kepuasan pengguna, baik pada kelas dengan jumlah data besar maupun kecil.

Untuk kelas “Sangat Puas”, model memperoleh *precision* sebesar 1.0000, *recall* sebesar 0.9523, dan *F1-score* sebesar 0.9756, dengan jumlah data (*support*) sebanyak 42. Pada kelas “Puas”, *precision* tercatat sebesar 0.9701, *recall* 0.9155,

dan *F1-score* sebesar 0.9420, dengan support sebanyak 71. Pada kelas “Cukup Puas” memiliki *precision* sebesar 0.8800, *recall* 1.0000, dan *F1-score* sebesar 0.93617 (0.9362), dari total 44 data. Untuk kelas “Tidak Puas”, nilai *precision* sebesar 0.96875 (0.9688), *recall* 1.0000, dan *F1-score* sebesar 0.9841 diperoleh dari 31 data. Sedangkan pada kelas “Sangat Tidak Puas”, *precision* mencapai 1.0000, *recall* 0.9677, dan *F1-score* sebesar 0.9836, dengan total data sebanyak 31.

Secara keseluruhan, model mencapai akurasi sebesar 95,89%, yang menunjukkan bahwa prediksi terhadap data uji dilakukan dengan sangat baik dan konsisten. Nilai macro average untuk *precision*, *recall*, dan *F1-score* masing-masing adalah 0.9637, 0.9671, dan 0.9643, yang mencerminkan rata-rata performa antar kelas tanpa mempertimbangkan jumlah data di masing-masing kelas. Sedangkan nilai weighted average yang mempertimbangkan distribusi data tiap kelas menunjukkan *precision* sebesar 0.9617, *recall* 0.9589, dan *F1-score* 0.9591. Hal ini mengindikasikan bahwa model mampu mempertahankan performa tinggi secara menyeluruh dalam mengklasifikasikan tingkat kepuasan pengguna aplikasi UMSU Academy.

4.2.5 Tampilan Menu Predict

Menu *Predict* merupakan bagian dari sistem yang berfungsi untuk melakukan prediksi tingkat kepuasan pengguna berdasarkan data input baru yang dimasukkan secara manual oleh pengguna. Pada halaman ini, pengguna diberikan form input untuk mengisi 25 indikator pertanyaan yang terbagi dalam lima dimensi utama, yaitu: *content* (C1–C5), *Format* (F1–F5), *Accuracy* (A1–A5), *Timelines* (T1–T5), dan *Ease of use* (E1–E5). Masing-masing indikator dinilai

4.3 Perhitungan Naïve Bayes

Tahapan pertama dilakukan persiapan data yang merupakan tahapan penting untuk memastikan data yang akan digunakan dalam model analisis atau *machine learning*. Dalam melakukan penelitian untuk menganalisis kepuasan pengguna menggunakan *Naïve Bayes* pada aplikasi *UMSU Academy* dilakukan proses *preprocessing* data. Data dikumpulkan dari kuesioner yang telah diisi oleh responden, yaitu pengguna aplikasi *UMSU Academy*. Pada penelitian ini Dataset yang akan digunakan sebagai data *training* sebanyak 1092 data dan untuk data *testing* yang digunakan yaitu 20% dari data *training*. Data mencakup berbagai atribut seperti:

Tabel 4. 1 Kuesioner Kepuasan Pengguna *UMSU Academy*

Variabel	Keterangan	Kategori
Content		
C1	Isi informasi pada aplikasi sesuai dengan kebutuhan akademik mahasiswa	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
C2	Isi informasi pada aplikasi mudah dipahami	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
C3	Isi informasi pada aplikasi sudah lengkap	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
C4	Isi informasi pada aplikasi sangat jelas	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju

C5	Isi informasi pada aplikasi selalu diperbaharui	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
Format		
F1	Desain tampilan aplikasi memiliki warna yang menarik	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
F2	Desain tampilan aplikasi memiliki tata letak yang memudahkan pengguna	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
F3	Desain tampilan aplikasi memiliki struktur menu dan fitur yang mudah dipahami	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
F4	Desain tampilan aplikasi menggunakan font yang jelas dan mudah dibaca	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
F5	Desain tampilan aplikasi menggunakan bahasa yang mudah dimengerti	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
Accuracy		
A1	Aplikasi menampilkan informasi yang benar dan akurat	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
A2	Setiap menu atau fitur yang diakses menampilkan halaman yang sesuai	1: Sangat Tidak Setuju 2: Tidak Setuju

		3: Netral 4: Setuju 5: Sangat Setuju
A3	Aplikasi memiliki keamanan yang baik	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
A4	Aplikasi jarang mengalami error saat digunakan	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
A5	Aplikasi berjalan sesuai dengan standar yang sudah ditentukan	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
Timelines		
T1	Informasi akademik cepat diperoleh melalui aplikasi	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
T2	Aplikasi selalu menampilkan informasi akademik terbaru	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
T3	Aplikasi menyajikan informasi atau data tepat waktu	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
T4	Aplikasi menginput data akademik tepat waktu	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju

T5	Download data akademik tepat waktu tanpa kendala	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
Ease of Use		
E1	Aplikasi sangat nyaman dan mudah digunakan	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
E2	Aplikasi mudah diakses kapan saja	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
E3	Aplikasi memberikan peringatan jika ada kendala jaringan/error	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
E4	Aplikasi tidak memerlukan waktu lama untuk menampilkan informasi	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju
E5	Informasi yang disampaikan pada aplikasi terstruktur dan mudah dipahami	1: Sangat Tidak Setuju 2: Tidak Setuju 3: Netral 4: Setuju 5: Sangat Setuju

Dari setiap atribut pada kuesioner yang telah di sebarakan kepada pengguna atau responden, dilakukan perhitungan interval untuk menentukan nilai rata-rata dari responden agar mengetahui kelas kepuasannya. Pada kuesioner memiliki 25 pertanyaan dan setiap pertanyaan memiliki skor 1-5, yaitu skala likert. Dimana

apabila sangat setuju = 5 skor, setuju = 4 skor, netral = 3 skor, tidak setuju = 2 skor, dan sangat tidak setuju = 1 skor.

perhitungan interval kelasnya sebagai berikut:

Skor minimum: 25 (pertanyaan) x 1 (nilai minimum dari skala linkert) = 25

Rata-rata minimum: $25/25 = 1.00$

Skor maksimum: 25 (pertanyaan) x 5 (nilai maksimum dari skala linkert) = 125

Rata-rata maksimum: $125/25 = 5.00$

Rentang rata-rata: $5.00 - 1.00 = 4.00$

Dikarnakan menggunakan 5 kategori kepuasan, maka interval $4.00/5 = 0.80$

Tabel 4. 2 Interval Kategori kepuasan

NO	Rata-rata nilai	Kategori kelas kepuasan
1	1,00 - 1,80	Sangat Tidak Puas
2	1,81 - 2,60	Tidak Puas
3	2,61 - 3,40	Cukup Puas
4	3,41 - 4,20	Puas
5	4,21 - 5,00	Sangat Puas

4.3.1 Perhitungan Probabilitas Class

Pada penelitian ini Dataset terdiri dari 1092 data latih yang digunakan, untuk menentukan perhitungan probabilitas kelas dapat dilihat pada perhitungan dengan rumus berikut:

$$P(C_k) = \frac{\text{jumlah sampel dalam kelas } C_k}{\text{total jumlah sampel}} \dots \dots \dots (4.1)$$

1. Sangat Puas

$$P = \frac{238}{1092} = 0.2179487179$$

2. Puas

$$P = \frac{409}{1092} = 0.3745421245$$

3. Cukup Puas

$$P = \frac{160}{1092} = 0.1465201465$$

4. Tidak Puas

$$P = \frac{144}{1092} = 0.1318681318$$

5. Sangat Tidak Puas

$$P = \frac{141}{1092} = 0.1291208791$$

Tabel 4. 3 Perhitungan Probabilitas *Class*

Label	Jumlah data	Jumlah seluruh data	Hasil
Sangat puas	238	1092	0.2179487179
Puas	409	1092	0.3745421245
Cukup puas	160	1092	0.1465201465
Tidak puas	144	1092	0.1318681318
Sangat Tidak Puas	141	1092	0.1291208791

4.3.2 Perhitungan Probabilitas Kategori

Pada tahapan ini dilakukan untuk menghitung kemungkinan kemunculan suatu nilai fitur (kategori) dalam setiap kelas berdasarkan data latih. dapat dijabarkan menggunakan rumus berikut:

$$P(X_i = x|C_k) = \frac{n_{x,k}}{n_k} \dots \dots \dots (4.2)$$

Keterangan:

$P(X_i = x|C_k)$ = Fitur ke- i (X_i) memiliki nilai x , data tersebut kelas C_k .

X_i = Fitur ke- i dari data

x = Nilai tertentu yang bisa dimiliki oleh fitur X_i

C_k = Kelas/kategori tertentu

$n_{x,k}$ = Jumlah data yang memenuhi dua kondisi sekaligus, fitur X_i

bernilai x , dan data tersebut termasuk dalam kelas C_k

n_k = jumlah total data yang termasuk kelas C_k

Perhitungan manual untuk probabilitas kategori fitur C1:

1. Baris: “Sangat Setuju”, kelas: “Sangat Puas”

$$P = (C1 = \text{"Sangat Setuju"} | \text{"Sangat Puas"}) = \frac{40}{238} = 0.168$$

2. Baris: “Setuju”, kelas: “Puas”

$$P = (C1 = \text{"Setuju"} | \text{"Puas"}) = \frac{195}{409} = 0.477$$

3. Baris: “Cukup Setuju”, “Cukup Puas”

$$P = (C1 = \text{"Cukup Setuju"} | \text{"Cukup Puas"}) = \frac{53}{160} = 0.331$$

4. Baris: “Tidak Setuju”, “Tidak Puas”

$$P = (C1 = \text{"Tidak Setuju"} | \text{"Tidak Puas"}) = \frac{41}{144} = 0.285$$

5. Baris: “Sangat Tidak Setuju”, “Sangat Tidak Puas”

$$P = (C1 = \text{"Sangat Tidak Setuju"} | \text{"Sangat Tidak Puas"}) = \frac{117}{141} = 0.83$$

Tabel 4. 4 Probabilitas Content (C1)

C1	Sangat Puas (J)	Prob	Puas (J)	Prob	Cukup Puas (J)	Prob	Tidak Puas (J)	Prob	Sangat Tidak Puas (J)	Prob
Sangat Setuju	40	0.168	80	0.196	21	0.131	0	0	0	0
Setuju	196	0.824	195	0.477	58	0.362	0	0	0	0
Cukup Setuju	2	0.008	99	0.242	53	0.331	22	0.153	2	0.014

C1	Sangat Puas (J)	Prob	Puas (J)	Prob	Cukup Puas (J)	Prob	Tidak Puas (J)	Prob	Sangat Tidak Puas (J)	Prob
Tidak Setuju	0	0	28	0.068	11	0.069	41	0.285	22	0.156
Sangat Tidak Setuju	0	0	7	0.017	17	0.106	81	0.562	117	0.83
Jumlah	238		409		160		144		141	

Perhitungan manual untuk probabilitas kategori fitur C2:

1. Baris: "Sangat Setuju", kelas: "Sangat Puas"

$$P = (C2 = "Sangat Setuju" | "Sangat Puas") = \frac{39}{238} = 0.164$$

2. Baris: "Setuju", kelas: "Puas"

$$P = (C2 = "Setuju" | "Puas") = \frac{192}{409} = 0.469$$

3. Baris: "Cukup Setuju", "Cukup Puas"

$$P = (C2 = "Cukup Setuju" | "Cukup Puas") = \frac{43}{160} = 0.269$$

4. Baris: "Tidak Setuju", "Tidak Puas"

$$P = (C2 = "Tidak Setuju" | "Tidak Puas") = \frac{52}{144} = 0.361$$

5. Baris: "Sangat Tidak Setuju", "Sangat Tidak Puas"

$$P = (C2 = \text{Sangat Tidak Setuju} | \text{Sangat Tidak Puas}) = \frac{115}{141} = 0.816$$

Tabel 4. 5 Probabilitas Content (C2)

C2	Sangat Puas	Prob	Puas	Prob	Cukup Puas	Prob	Tidak Puas	Prob	Sangat Tidak Puas	Prob
----	-------------	------	------	------	------------	------	------------	------	-------------------	------

C2	Sangat Puas	Prob	Puas	Prob	Cukup Puas	Prob	Tidak Puas	Prob	Sangat Tidak Puas	Prob
Sangat Setuju	39	0.164	73	0.178	10	0.062	0	0	0	0
Setuju	198	0.832	192	0.469	71	0.444	0	0	0	0
Cukup Setuju	1	0.004	95	0.232	43	0.269	12	0.083	1	0.007
Tidak Setuju	0	0	33	0.081	26	0.162	52	0.361	25	0.177
Sangat Tidak Setuju	0	0	16	0.039	10	0.062	80	0.556	115	0.816
Jumlah	238		409		160		144		141	

Perhitungan manual untuk probabilitas kategori fitur C3:

1. Baris: "Sangat Setuju", kelas: "Sangat Puas"

$$P = (C3 = "Sangat Setuju" | "Sangat Puas") = \frac{220}{238} = 0.924$$

2. Baris: "Setuju", kelas: "Puas"

$$P = (C3 = "Setuju" | "Puas") = \frac{192}{409} = 0.469$$

3. Baris: "Cukup Setuju", "Cukup Puas"

$$P = (C3 = "Cukup Setuju" | "Cukup Puas") = \frac{50}{160} = 0.312$$

4. Baris: "Tidak Setuju", "Tidak Puas"

$$P = (C3 = "Tidak Setuju" | "Tidak Puas") = \frac{78}{144} = 0.542$$

5. Baris: "Sangat Tidak Setuju", "Sangat Tidak Puas"

$$P = (C3 = "Sangat Tidak Setuju" | "Sangat Tidak Puas") = \frac{102}{141} = 0.723$$

Tabel 4. 6 Probabilitas Content (C3)

C3	Sangat Puas	Prob	Puas	Prob	Cukup Puas	Prob	Tidak Puas	Prob	Sangat Tidak Puas	Prob
Sangat Setuju	220	0.924	80	0.196	19	0.119	0	0	0	0
Setuju	15	0.063	192	0.469	55	0.344	0	0	0	0
Cukup Setuju	3	0.013	84	0.205	50	0.312	15	0.104	0	0
Tidak Setuju	0	0	40	0.098	20	0.125	78	0.542	39	0.277
Sangat Tidak Setuju	0	0	13	0.032	16	0.1	51	0.354	102	0.723
Jumlah	238		409		160		144		141	

Perhitungan manual untuk probabilitas kategori fitur C4:

1. Baris: “Sangat Setuju”, kelas: “Sangat Puas”

$$P = (C4 = "Sangat Setuju" | "Sangat Puas") = \frac{225}{238} = 0.945$$

2. Baris: “Setuju”, kelas: “Puas”

$$P = (C4 = "Setuju" | "Puas") = \frac{206}{409} = 0.504$$

3. Baris: “Cukup Setuju”, “Cukup Puas”

$$P = (C4 = "Cukup Setuju" | "Cukup Puas") = \frac{54}{160} = 0.338$$

4. Baris: “Tidak Setuju”, “Tidak Puas”

$$P = (C4 = "Tidak Setuju" | "Tidak Puas") = \frac{78}{144} = 0.542$$

5. Baris: “Sangat Tidak Setuju”, “Sangat Tidak Puas”

$$P = (C4 = "Sangat Tidak Setuju" | "Sangat Tidak Puas") = \frac{122}{141} = 0.865$$

Tabel 4. 7 Probabilitas Content (C4)

C4	Sangat Puas	Prob	Puas	Prob	Cukup Puas	Prob	Tidak Puas	Prob	Sangat Tidak Puas	Prob
Sangat Setuju	225	0.945	84	0.205	12	0.075	0	0	0	0
Setuju	11	0.046	206	0.504	59	0.369	1	0.007	0	0
Cukup Setuju	2	0.008	83	0.203	54	0.338	44	0.306	0	0
Tidak Setuju	0	0	22	0.054	23	0.144	78	0.542	19	0.135
Sangat Tidak Setuju	0	0	14	0.034	12	0.075	21	0.146	122	0.865
Jumlah	238		409		160		144		141	

Perhitungan manual untuk probabilitas kategori fitur C5:

1. Baris: “Sangat Setuju”, kelas: “Sangat Puas”

$$P = (C5 = \text{"Sangat Setuju"} | \text{"Sangat Puas"}) = \frac{39}{238} = 0.164$$

2. Baris: “Setuju”, kelas: “Puas”

$$P = (C5 = \text{"Setuju"} | \text{"Puas"}) = \frac{191}{409} = 0.467$$

3. Baris: “Tidak Setuju”, “Cukup Puas”

$$P = (C5 = \text{"Tidak Setuju"} | \text{"Cukup Puas"}) = \frac{61}{160} = 0.381$$

4. Baris: “Tidak Setuju”, “Tidak Puas”

$$P = (C5 = \text{"Tidak Setuju"} | \text{"Tidak Puas"}) = \frac{55}{144} = 0.382$$

5. Baris: “Sangat Tidak Setuju”, “Sangat Tidak Puas”

$$P = (C5 = \text{"Sangat Tidak Setuju"} | \text{"Sangat Tidak Puas"}) = \frac{124}{141} = 0.879$$

Tabel 4. 8 Probabilitas Content (C5)

C5	Sangat Puas	Prob	Puas	Prob	Cukup Puas	Prob	Tidak Puas	Prob	Sangat Tidak Puas	Prob
Sangat Setuju	39	0.164	73	0.178	14	0.088	0	0	0	0
Setuju	198	0.832	191	0.467	55	0.344	1	0.007	0	0
Cukup Setuju	1	0.004	96	0.235	61	0.381	23	0.16	0	0
Tidak Setuju	0	0	34	0.083	22	0.138	55	0.382	17	0.121
Sangat Tidak Setuju	0	0	15	0.037	8	0.05	65	0.451	124	0.879
Jumlah	238		409		160		144		141	

Tabel format (F1) hingga tabel ease of use (E3) tidak ditampilkan dalam dokumen ini guna efisiensi, namun seluruh data tetap digunakan dalam proses pengolahan dan analisis.

Perhitungan manual untuk probabilitas kategori fitur E4:

1. Baris: “Sangat Setuju”, kelas: “Sangat Puas”

$$P = (E4 = "Sangat Setuju" | "Sangat Puas") = \frac{217}{238} = 0.912$$

2. Baris: “Setuju”, kelas: “Puas”

$$P = (E4 = "Setuju" | "Puas") = \frac{186}{409} = 0.455$$

3. Baris: “Cukup Setuju”, “Cukup Puas”

$$P = (E4 = "Cukup Setuju" | "Cukup Puas") = \frac{49}{160} = 0.306$$

4. Baris: “Tidak Setuju”, “Tidak Puas”

$$P = (E4 = "Tidak Setuju" | "Tidak Puas") = \frac{68}{144} = 0.472$$

5. Baris: “Sangat Tidak Setuju”, “Sangat Tidak Puas”

$$P = (E4 = \text{"Sangat Tidak Setuju"} \mid \text{"Sangat Tidak Puas"}) = \frac{114}{141} = 0.809$$

Tabel 4. 9 Probabilitas *Ease of use* (E4)

E4	Sangat Puas	Prob	Puas	Prob	Cukup Puas	Prob	Tidak Puas	Prob	Sangat Tidak Puas	Prob
Sangat Setuju	217	0.912	75	0.183	14	0.088	0	0	0	0
Setuju	13	0.055	186	0.455	62	0.388	0	0	0	0
Cukup Setuju	7	0.029	107	0.262	49	0.306	32	0.222	0	0
Tidak Setuju	1	0.004	32	0.078	23	0.144	68	0.472	27	0.191
Sangat Tidak Setuju	0	0	9	0.022	12	0.075	44	0.306	114	0.809
Jumlah	238		409		160		144		141	

Perhitungan manual untuk probabilitas kategori fitur E5:

1. Baris: "Sangat Setuju", kelas: "Sangat Puas"

$$P = (E5 = \text{"Sangat Setuju"} \mid \text{"Sangat Puas"}) = \frac{34}{238} = 0.143$$

2. Baris: "Setuju", kelas: "Puas"

$$P = (E5 = \text{"Cukup Setuju"} \mid \text{"Puas"}) = \frac{166}{409} = 0.406$$

3. Baris: "Cukup Setuju", "Cukup Puas"

$$P = (E5 = \text{"Setuju"} \mid \text{"Cukup Puas"}) = \frac{46}{160} = 0.288$$

4. Baris: "Setuju", "Tidak Puas"

$$P = (E5 = \text{"Tidak Setuju"} \mid \text{"Tidak Puas"}) = \frac{47}{144} = 0.326$$

5. Baris: "Sangat Tidak Setuju", "Sangat Tidak Puas"

$$P = (E5 = \text{"Sangat Tidak Setuju"} \mid \text{"Sangat Tidak Puas"}) = \frac{131}{141} = 0.929$$

Tabel 4. 10 Probabilitas *Ease of use* (E5)

E5	Sangat Puas	Prob	Puas	Prob	Cukup Puas	Prob	Tidak Puas	Prob	Sangat Tidak Puas	Prob
Sangat Setuju	34	0.143	82	0.2	17	0.106	0	0	0	0
Setuju	200	0.84	166	0.406	50	0.312	1	0.007	0	0
Cukup Setuju	4	0.017	108	0.264	46	0.288	27	0.188	0	0
Tidak Setuju	0	0	35	0.086	30	0.188	47	0.326	10	0.071
Sangat Tidak Setuju	0	0	18	0.044	17	0.106	69	0.479	131	0.929
Jumlah	238		409		160		144		141	

4.3.3 Perhitungan Manual Naïve Bayes

Pada perhitungan manual *Naïve bayes*, menggunakan potongan data dari data uji atau testing yang telah tersedia, berikut data yang digunakan:

Tabel 4. 11 Potongan Data *Testing*

C	C	C	C	C	F	F	F	F	F	A	A	A	A	A	T	T	T	T	T	E	E	E	E	E
1	2	3	4	5	1	2	3	4	5	1	2	3	4	5	1	2	3	4	5	1	2	3	4	5
4	5	5	5	4	4	4	5	5	5	4	3	3	2	4	4	4	4	4	5	5	5	5	5	4

Pada tabel 4.11 data yang digunakan merupakan potongan data dari data testing pada baris 96 dalam dataset. Dalam perhitungan manual ini akan menggunakan rumus seperti menghitung prior probabilitas, menghitung *likelihood* untuk setiap fitur terhadap setiap kelas, menghitung posterior dan menentukan hasil prediksi berdasarkan nilai tertinggi.

4.3.4 Menghitung Prior Probabilitas dan *Likelihood*

Rumus Prior probabilitas:

$$P(C_k) = \frac{n_k}{n} \dots \dots \dots (4.3)$$

Keterangan:

n_k = Jumlah data pada kelas C_k dalam data training

n = Total data training

1. Sangat Puas

$$P = \frac{196}{873} = 0.225$$

2. Puas

$$P = \frac{338}{873} = 0.387$$

3. Cukup Puas

$$P = \frac{116}{873} = 0.133$$

4. Tidak Puas

$$P = \frac{113}{873} = 0.129$$

5. Sangat Tidak Puas

$$P = \frac{110}{873} = 0.126$$

Tabel 4. 12 Menghitung Prior Probabilitas

Class Kepuasan	Jumlah	Probabilitas
Sangat Puas	196	0.225
Puas	338	0.387
Cukup Puas	116	0.133
Tidak Puas	113	0.129
Sangat Tidak Puas	110	0.126

Rumus Likelihood:

$$P(X_i = x|C_k) = \frac{n_{x,k}}{n_k} \dots \dots \dots (4.4)$$

Perhitungan manual untuk kategori fitur C1:

1. Sangat Puas

$$P = (C1 = 4 | Sangat Puas) = \frac{162}{196} = 0.8265$$

2. Puas

$$P = (C1 = 4 | Puas) = \frac{163}{338} = 0.4822$$

3. Cukup Puas

$$P = (C1 = 4 | Cukup Puas) = \frac{44}{116} = 0.3793$$

4. Tidak Puas

$$P = (C1 = 4 | Tidak Puas) = \frac{0}{113} = 0.000$$

5. Sangat Tidak Puas

$$P = (C1 = 4 | Sangat Tidak Puas) = \frac{0}{110} = 0.000$$

Tabel 4. 13 Fitur C1

Class	Jumlah C1 = 4	Total Kelas	Likelihood
Sangat Puas	162	196	0.8265
Puas	163	338	0.4822
Cukup Puas	44	116	0.3793
Tidak Puas	0	113	0.0000
Sangat Tidak Puas	0	110	0.0000

Perhitungan manual untuk kategori fitur C2:

1. Sangat Puas

$$P = (C2 = 5 | Sangat Puas) = \frac{30}{196} = 0.1530$$

2. Puas

$$P = (C2 = 5 | Puas) = \frac{61}{338} = 0.1804$$

3. Cukup Puas

$$P = (C2 = 5 | Cukup Puas) = \frac{9}{116} = 0.0775$$

4. Tidak Puas

$$P = (C2 = 5 | Tidak Puas) = \frac{0}{113} = 0.000$$

5. Sangat Tidak Puas

$$P = (C2 = 5 | Sangat Tidak Puas) = \frac{0}{110} = 0.000$$

Tabel 4. 14 Fitur C2

<i>Class</i>	Jumlah C2 = 5	Total Kelas	Likelihood
Sangat Puas	30	196	0.7333
Puas	61	338	0.3699
Cukup Puas	9	116	0.0526
Tidak Puas	0	113	0.0000
Sangat Tidak Puas	0	110	0.0000

Perhitungan manual untuk kategori fitur C3:

1. Sangat Puas

$$P = (C3 = 5 | Sangat Puas) = \frac{181}{196} = 0.9234$$

2. Puas

$$P = (C3 = 5 | Puas) = \frac{62}{338} = 0.1834$$

3. Cukup Puas

$$P = (C3 = "5" | Cukup Puas) = \frac{14}{116} = 0.1206$$

4. Tidak Puas

$$P = (C3 = 5 | Tidak Puas) = \frac{0}{113} = 0.000$$

5. Sangat Tidak Puas

$$P = (C3 = 5 | Sangat Tidak Puas) = \frac{0}{110} = 0.000$$

Tabel 4. 15 Fitur C3

<i>Class</i>	Jumlah C3 = 5	Total Kelas	Likelihood
Sangat Puas	181	196	0.923469
Puas	62	338	0.183432
Cukup Puas	14	116	0.12069
Tidak Puas	0	113	0
Sangat Tidak Puas	0	110	0

Perhitungan manual untuk kategori fitur C4:

1. Sangat Puas

$$P = (C4 = 5 | Sangat Puas) = \frac{185}{196} = 0.9438$$

2. Puas

$$P = (C4 = 5 | Puas) = \frac{71}{338} = 0.2100$$

3. Cukup Puas

$$P = (C4 = "5" | Cukup Puas) = \frac{7}{116} = 0.0603$$

4. Tidak Puas

$$P = (C4 = 5 | Tidak Puas) = \frac{0}{113} = 0.000$$

5. Sangat Tidak Puas

$$P = (C4 = 5 | Sangat Tidak Puas) = \frac{0}{110} = 0.000$$

Tabel 4. 16 Fitur C4

<i>Class</i>	Jumlah C4 = 5	Total Kelas	Likelihood
Sangat Puas	185	196	0.045918
Puas	71	338	0.210059
Cukup Puas	7	116	0.060345
Tidak Puas	0	113	0
Sangat Tidak Puas	0	110	0

Perhitungan manual untuk kategori fitur C5:

1. Sangat Puas

$$P = (C5 = 4 | Sangat Puas) = \frac{162}{196} = 0.8265$$

2. Puas

$$P = (C5 = 4 | Puas) = \frac{151}{338} = 0.4467$$

3. Cukup Puas

$$P = (C5 = "4" | Cukup Puas) = \frac{42}{116} = 0.3620$$

4. Tidak Puas

$$P = (C5 = 4 | Tidak Puas) = \frac{1}{113} = 0.0088$$

5. Sangat Tidak Puas

$$P = (C5 = 4 | Sangat Tidak Puas) = \frac{0}{110} = 0.000$$

Tabel 4. 17 Fitur C5

<i>Class</i>	Jumlah C5 = 4	Total Kelas	Likelihood
Sangat Puas	162	196	0.826531
Puas	151	338	0.446746
Cukup Puas	42	116	0.362069
Tidak Puas	1	113	0.00885
Sangat Tidak Puas	0	110	0

Tabel fitur (F1) hingga tabel fitur (E3) tidak ditampilkan dalam dokumen ini guna efisiensi, namun seluruh data tetap digunakan dalam pengolahan dan analisis.

Perhitungan manual untuk kategori fitur E4:

1. Sangat Puas

$$P = (E4 = 5 | Sangat Puas) = \frac{178}{196} = 0.9081$$

2. Puas

$$P = (E4 = 5 | Puas) = \frac{66}{338} = 0.1952$$

3. Cukup Puas

$$P = (E4 = "5" | Cukup Puas) = \frac{13}{116} = 0.1120$$

4. Tidak Puas

$$P = (E4 = 5 | Tidak Puas) = \frac{0}{113} = 0.000$$

5. Sangat Tidak Puas

$$P = (E4 = 5 | Sangat Tidak Puas) = \frac{0}{110} = 0.000$$

Tabel 4. 18 Fitur E4

<i>Class</i>	Jumlah E4 = 3	Total Kelas	Likelihood
Sangat Puas	178	196	0.908163
Puas	66	338	0.195266

<i>Class</i>	Jumlah E4 = 3	Total Kelas	Likelihood
Cukup Puas	13	116	0.112069
Tidak Puas	0	113	0
Sangat Tidak Puas	0	110	0

Perhitungan manual untuk kategori fitur E5:

1. Sangat Puas

$$P = (E5 = 4 | Sangat Puas) = \frac{164}{196} = 0.8367$$

2. Puas

$$P = (E5 = 4 | Puas) = \frac{138}{338} = 0.4082$$

3. Cukup Puas

$$P = (E5 = 4 | "Cukup Puas") = \frac{38}{116} = 0.3275$$

4. Tidak Puas

$$P = (E5 = 4 | Tidak Puas) = \frac{1}{113} = 0.0088$$

5. Sangat Tidak Puas

$$P = (E5 = 4 | Sangat Tidak Puas) = \frac{0}{110} = 0.000$$

Tabel 4. 19 Fitur E5

<i>Class</i>	Jumlah E5 = 4	Total Kelas	Likelihood
Sangat Puas	164	196	0.836735
Puas	138	338	0.408284
Cukup Puas	38	116	0.327586
Tidak Puas	1	113	0.00885
Sangat Tidak Puas	0	110	0

4.3.5 Menghitung Posterior Dan Menentukan Hasil Prediksi

Rumus Posterior tanpa normalisasi:

$$Posterior(C_k) \propto P(C_k) \times \prod_{i=1}^n P(X_i | C_k)$$

Keterangan:

C_k = kelas ke- k

$P(C_k)$ = prior probabilitas kelas C_k

$P(X_i | C_k)$ = likelihood fitur ke- i terhadap kelas C_k

$\prod$ = dikalikan semuanya

Perhitungan manual posterior tanpa normalisasi:

1. Sangat Puas

$$\begin{aligned} \text{Score} &= 0.225 \times 0.8265 \times 0.1531 \times 0.9235 \times 0.9439 \times 0.8265 \times \dots \times \\ &0.9082 \times 0.8367 = 8.221 \times 10^{-16} \end{aligned}$$

2. Puas

$$\begin{aligned} \text{Score} &= 0.387 \times 0.4822 \times 0.1805 \times 0.1834 \times 0.2101 \times 0.4467 \times \dots \times \\ &0.1953 \times 0.4083 = 6.802 \times 10^{-15} \end{aligned}$$

3. Cukup Puas

$$\begin{aligned} \text{Score} &= 0.133 \times 0.3793 \times 0.0776 \times 0.1207 \times 0.0603 \times 0.3621 \times \dots \\ &\times 0.1121 \times 0.3276 = 8.774 \times 10^{-19} \end{aligned}$$

4. Tidak Puas

$$\text{Score} = 0.129 \times 0 \times \dots = 0$$

5. Sangat Tidak Puas

$$\text{Score} = 0.126 \times 0 \times \dots = 0$$

Tabel 4. 20 Hasil Posterior Tanpa Normalisasi

<i>Class</i>	Skor Naive Bayes
Sangat Puas	8.221×10^{-16}
Puas	6.802×10^{-15}
Cukup Puas	8.774×10^{-19}

Class	Skor Naive Bayes
Sangat Puas	8.221×10^{-16}
Tidak Puas	0.0
Sangat Tidak Puas	0.0

Rumus Posterior Naïve Bayes:

$$P(C_k | X) = \frac{\tilde{P}(C_k | X)}{\sum_j \tilde{P}(C_j | X)}$$

Keterangan:

$\tilde{P}(C_k | X)$ = posterior *tanpa normalisasi* (hasil perkalian prior dan likelihood fitur).

$\sum_j \tilde{P}(C_j | X)$ = penjumlahan dari semua posterior tanpa normalisasi untuk semua kelas.

Perhitungan manual posterior Naïve Bayes:

$$S = 8.221 \times 10^{-16} + 6.802 \times 10^{-15} + 8.774 \times 10^{-19} = 7.624 \times 10^{-15}$$

Probabilitas ternormalisasi kelas "Sangat Puas":

$$P(\text{Sangat Puas} | X) = \frac{8.221 \times 10^{-16}}{7.624 \times 10^{-15}} = 0.1078$$

Probabilitas ternormalisasi kelas "Puas":

$$P(\text{Puas} | X) = \frac{6.802 \times 10^{-15}}{7.624 \times 10^{-15}} = 0.8918$$

Probabilitas ternormalisasi kelas "Puas":

$$P(\text{cukup Puas} | X) = \frac{8.774 \times 10^{-19}}{7.624 \times 10^{-15}} = 0.00011$$

Tabel 4. 21 Hasil Posterior *Naïve Bayes*

<i>Class</i>	Probabilitas Ternormalisasi	Keterangan
Sangat Puas	0.1078 (10.87%)	
Puas	0.8918 (89.18%)	Prediksi Tertinggi
Cukup Puas	0.00011(0.01%)	
Tidak Puas	0.0000	
Sangat Tidak Puas	0.0000	

4.3.6 Perhitungan Precision, Recall, F1-Score Dan Akurasi

Perhitungan manual Precision, Recall, F1-Score Dan Akurasi:

1. Sangat Puas

TP = 40, FN = 2, FP = 0

$$Precision = \frac{40}{40 + 0} = 1.0$$

$$Recall = \frac{40}{40 + 2} = 0.9524$$

$$F1 - Score = 2 \times \frac{1.0 \times 0.9524}{1.0 + 0.9524} = 0.9756$$

2. Puas

TP = 65, FN = 6, FP = 2

$$Precision = \frac{65}{65 + 2} = 0.9701$$

$$Recall = \frac{65}{65 + 6} = 0.9155$$

$$F1 - Score = 2 \times \frac{0.9701 \times 0.9155}{0.9701 + 0.9155} = 0.9420$$

3. Cukup Puas

TP = 44, FN = 0, FP = 6

$$Precision = \frac{44}{44 + 6} = 0.8800$$

$$Recall = \frac{44}{44 + 0} = 1.0$$

$$F1 - Score = 2 \times \frac{0.880 \times 1.0}{0.880 + 1.0} = 0.9362$$

4. Tidak Puas

TP = 31, FN = 0, FP = 1

$$Precision = \frac{31}{31 + 1} = 0.9688$$

$$Recall = \frac{31}{31 + 0} = 1.00$$

$$F1 - Score = 2 \times \frac{0.9688 \times 1.0}{0.9688 + 1.0} = 0.9841$$

5. Sangat Tidak Puas

TP = 30, FN = 1, FP = 0

$$Precision = \frac{30}{30 + 0} = 1.00$$

$$Recall = \frac{30}{30 + 1} = 0.9677$$

$$F1 - Score = 2 \times \frac{1.00 \times 0.9677}{1.00 + 0.9677} = 0.9836$$

Tabel 4. 22 Perhitungan *Precision*, *Recall*, *F1-Score* Dan Akurasi

<i>Class</i>	TP	FP	FN	TN	Precision	Recall	F1 Score
Sangat Puas	40	0	2	177	1.00	0.9524	0.9756
Puas	65	2	6	146	0.9701	0.9155	0.9420
Cukup Puas	44	6	0	169	0.880	1.00	0.9362
Tidak Puas	31	1	0	187	0.9688	1.00	0.9841
Sangat Tidak Puas	30	0	1	188	1.00	0.9677	0.9836

Keterangan:

TP: True Positive (diprediksi benar)

FP: False Positive (prediksi salah masuk ke kelas ini)

FN: False Negative (kelas ini diprediksi jadi kelas lain)

TN: True Negative (bukan kelas ini dan diprediksi juga bukan kelas ini)

Precision: $TP / (TP + FP)$

Recall: $TP / (TP + FN)$

F1 Score: $2 * (Precision * Recall) / (Precision + Recall)$

Proses perhitungan akurasi pada sistem menggunakan metode Naive Bayes dengan menguji berdasarkan data training yang diambil dari kuesioner sebanyak 1092 data responden. Hasil perhitungan pada tingkat akurasi yang diperoleh dari data training yaitu sebesar 0,9589 atau 95,89%. Perhitungan tersebut diproses dan dibagi oleh sistem sebanyak 80% data training dan 20% data testing atau sebanyak 873 data training dan 219 data testing.

$$Akurasi = \frac{TP(Total)}{Total Data} = \frac{40+65+44+31+30}{219} = \frac{210}{219} = 0.9589 \rightarrow 95,89\%$$

BAB V

PENUTUP

5.1 Kesimpulan

Berdasarkan hasil penelitian dan pembahasan mengenai Analisis Kepuasan Pengguna Aplikasi UMSU *Academy* Menggunakan Metode *Naïve Bayes* pada mahasiswa Universitas Muhammadiyah Sumatera Utara tahun ajaran 2021 yang terdiri dari 1092 responden, dapat disimpulkan bahwa terdapat pengguna yang merasa “Sangat Puas” ada sebanyak 238 responden, “Puas” sebanyak 409 responden, “Cukup Puas” sebanyak 160 responden, pengguna yang menyatakan “Tidak Puas” sebanyak 144 responden, dan “Sangat Tidak Puas” terdapat 141 responden. Jumlah pengguna yang puas jauh lebih banyak, sehingga menunjukkan mayoritas pengguna memiliki pengalaman positif terhadap layanan dan pengalaman pada aplikasi UMSU *Academy*.

Pada penelitian ini berhasil menggunakan algoritma *Naive Bayes* untuk menganalisis kepuasan pengguna terhadap aplikasi UMSU *Academy* dengan hasil yang cukup akurat. Penggunaan algoritma *Naïve Bayes* terbukti efektif dan efisien karena mampu memberikan prediksi yang akurat dengan data yang terbatas. Hal ini mengindikasikan bahwa aplikasi UMSU *Academy* secara keseluruhan telah memenuhi kebutuhan dan harapan pengguna, terutama dalam aspek kemudahan penggunaan yang menjadi faktor dengan kontribusi terbesar terhadap tingkat kepuasan mahasiswa.

Pada penelitian ini memungkinkan pengguna untuk mengunggah dataset, melakukan pemrosesan awal data, melatih model klasifikasi, serta melakukan evaluasi dan prediksi terhadap data baru. Pada proses evaluasi model, dilakukan

pembagian data sebesar 80% untuk pelatihan dan 20% untuk pengujian, dengan menggunakan parameter *random state* sebesar 42 untuk memastikan konsistensi hasil. Evaluasi model dilakukan menggunakan metrik akurasi, precision, recall, dan F1-score. Berdasarkan hasil evaluasi, model *Naïve Bayes* yang dibangun menghasilkan nilai akurasi total sebesar 95,89%, yang menunjukkan bahwa model mampu melakukan klasifikasi dengan baik. Evaluasi model menghasilkan nilai precision dan recall yang tinggi pada kategori “Puas” dan “Sangat Puas”.

Berdasarkan hasil evaluasi model menggunakan confusion matrix, diketahui bahwa model memiliki kinerja yang sangat baik dalam mengklasifikasikan data pada sebagian besar kelas. Kategori “Sangat Puas” menunjukkan nilai precision sebesar 1,0000 dan recall sebesar 0,9524 dengan F1-score sebesar 0,9756. Hal ini menunjukkan bahwa model mampu mengidentifikasi kelas “Sangat Puas” dengan sangat akurat, baik dalam hal ketepatan prediksi maupun kelengkapan pengenalan data. Pada kelas “Puas” juga menunjukkan performa yang tinggi dengan precision sebesar 0,9701, recall sebesar 0,9155, dan F1-score sebesar 0,9420. Nilai-nilai ini menunjukkan bahwa model cukup andal dalam mengenali kelas ini yang merupakan salah satu kelas mayoritas dalam distribusi data.

Pada kelas “Cukup Puas” memperoleh recall sempurna sebesar 1,0000 dan precision sebesar 0,8800, dengan F1-score sebesar 0,9362. Meskipun precision-nya lebih rendah dibanding kelas lain, model tetap mampu mengenali seluruh data aktual pada kelas ini tanpa kesalahan. Untuk kelas “Tidak Puas” dan “Sangat Tidak Puas”, masing-masing memperoleh F1-score sebesar 0,9841 dan 0,9836, dengan nilai recall sebesar 1,0000 dan 0,9677. Hal ini menunjukkan bahwa model

tetap mampu mengklasifikasikan data dari kelas minoritas secara akurat, meskipun jumlah datanya relatif lebih sedikit dibandingkan kelas mayoritas.

Secara keseluruhan, model Naïve Bayes yang digunakan dalam penelitian ini menunjukkan performa klasifikasi yang sangat baik, baik pada kelas mayoritas maupun minoritas. Hal ini menunjukkan bahwa model tidak hanya akurat dalam mengenali kelas yang dominan, tetapi juga mampu mempertahankan performa yang baik terhadap kelas dengan distribusi data yang lebih kecil.

5.2 Saran

Hasil penelitian pada kepuasan pengguna aplikasi UMSU *Academy* menggunakan metode *Naïve Bayes*, terdapat beberapa saran yang dapat disampaikan untuk pengembangan dan penyempurnaan sistem di masa yang akan datang. Berdasarkan kesimpulan yang diperoleh, penulis menyarankan agar pengembang aplikasi UMSU *Academy* lebih memperhatikan peningkatan keakuratan informasi pada fitur-fitur yang berkaitan langsung dengan data akademik mahasiswa, sehingga tingkat kepuasan pengguna dapat terus ditingkatkan. Pihak universitas juga diharapkan dapat secara berkelanjutan melakukan sosialisasi dan pelatihan mengenai penggunaan aplikasi UMSU *Academy*, agar seluruh mahasiswa dapat memanfaatkan seluruh fitur yang tersedia secara optimal. Disarankan untuk penelitian selanjutnya, menambah jumlah responden dari berbagai program studi dan angkatan, serta melakukan perbandingan dengan metode klasifikasi lain, seperti *Decision Tree* atau *Support Vector Machine*, sehingga hasil analisis kepuasan dapat diperoleh dengan lebih komprehensif dan akurat.

Dengan adanya upaya perbaikan dan pengembangan yang berkelanjutan, diharapkan aplikasi *UMSU Academy* dapat menjadi salah satu media pembelajaran digital yang handal dan mampu mendukung seluruh aktivitas akademik mahasiswa Universitas Muhammadiyah Sumatera Utara dengan lebih baik, sehingga kepuasan pengguna dapat terus ditingkatkan sejalan dengan perkembangan kebutuhan dan teknologi informasi di lingkungan kampus.

DAFTAR PUSTAKA

- Afrianto, I., Heryandi, A., Finadhita, A., & Atin, S. (2021). Work From Home Program. *International Journal of Information System & Technology Akreditasi*, 5(3), 270–280. <https://tt-el.my.id/>.
- Ahmad, & Muslimah. (2021). Memahami Teknik Pengolahan dan Analisis Data Kualitatif. *Proceedings*, 1(1), 173–186.
- Aulia, S. (2021). Klasterisasi Pola Penjualan Pestisida Menggunakan Metode K-Means Clustering (Studi Kasus Di Toko Juanda Tani Kecamatan Hutabayu Raja). *Djtechno: Jurnal Teknologi Informasi*, 1(1), 1–5. <https://doi.org/10.46576/djtechno.v1i1.964>
- Azis, H., Tangguh Admojo, F., & Susanti, E. (2020). Analisis Perbandingan Performa Metode Klasifikasi pada Dataset Multiclass Citra Busur Panah Performance Comparison Analysis of Classification Methods on the Multiclass Dataset of Bows. In *Agustus* (Vol. 19, Issue 3).
- Chan, V. H. Y., Chiu, D. K. W., & Ho, K. K. W. (2022). Mediating effects on the relationship between perceived service quality and public library app loyalty during the COVID-19 era. *Journal of Retailing and Consumer Services*, 67. <https://doi.org/10.1016/j.jretconser.2022.102960>
- Chen, H., Hu, S., Hua, R., & Zhao, X. (2021). Improved naive Bayes classification algorithm for traffic risk management. *Eurasip Journal on Advances in Signal Processing*, 2021(1). <https://doi.org/10.1186/s13634-021-00742-6>
- Criollo-C, S., Guerrero-Arias, A., Jaramillo-Alcázar, Á., & Luján-Mora, S.

- (2021). Mobile learning technologies for education: Benefits and pending issues. *Applied Sciences* (Switzerland), 11(9). <https://doi.org/10.3390/app11094111>
- Dari, W., & Elen Tania Hanayah. (2023). Analisis Tingkat Kepuasan Pengguna Aplikasi Ojek Online Dengan Metode Naive Bayes. *INSOLOGI: Jurnal Sains Dan Teknologi*, 2(1), 221–232. <https://doi.org/10.55123/insologi.v2i1.1693>
- Erlich, Z., & Zviran, M. (2003). Measuring IS User Satisfaction: Review and Implications. *Communications of the Association for Information Systems*, 12(1), 81–103.
- Fakhruddin, A. M., Putri, L. O., Rizqi, P., Sudirman, A. T., Annisa, R. N., Khalda, R., As, B., Studi, P., Guru, P., & Dasar, S. (2022). *Efektivitas LMS (Learning Management System) untuk Mengelola Pembelajaran Jarak Jauh pada Satuan Pendidikan*.
- Febriani, S. (2022). *Analisis Data Hasil Diagnosa Untuk Klasifikasi Gangguan Kepribadian Menggunakan Algoritma C4.5 Siska Febriani Sistem Informasi* *) Rohmansyah@gmail.com. 2(9), 1–9.
- Firdaus, D. (2017). Penggunaan Data mining dalam kegiatan pembelajaran. *Jurnal Format Volume 6 Nomor 2 Tahun 2017*, 6(2), 91–97.
- Hendrian, S. (2018). Algoritma Klasifikasi Data Mining Untuk Memprediksi Siswa Dalam Memperoleh Bantuan Dana Pendidikan. *Faktor Exacta*, 11(3), 266–274. <https://doi.org/10.30998/faktorexacta.v11i3.2777>
- Kamil, M., & Cholil, W. (2020). Analisis Perbandingan Algoritma C4.5 dan Naive Bayes pada Lulusan Tepat Waktu Mahasiswa di Universitas Islam

- Negeri Raden Fatah Palembang. *Jurnal Informatika*, 7(2), 97–106.
<https://doi.org/10.31294/ji.v7i2.7723>
- Kinanti, N., Putri1, A., & Dwi, A. (2021). Penerapan PIECES Framework sebagai Evaluasi Tingkat Kepuasan Mahasiswa terhadap Penggunaan Sistem Informasi Akademik Terpadu (SIKADU) pada Universitas Negeri Surabaya. *JEISBI*, 02. <https://siakadu.unesa.ac.id>
- Larasati, I., Yusril, A. N., & Zukri, P. Al. (2021). Systematic Literature Review Analisis Metode Agile Dalam Pengembangan Aplikasi Mobile. *Sistemasi*, 10(2), 369. <https://doi.org/10.32520/stmsi.v10i2.1237>
- Lattu, A., Sihabuddin, & Jatmika, W. (2022). 115-Article Text-391-1-10-20220210. *JURNAL RISET SISTEM INFORMASI DAN TEKNOLOGI INFORMASI (JURSISTEKNI)*, Vol 4 No 1, 39 – 50.
- Li, Q., Li, Z., & Han, J. (2021). A hybrid learning pedagogy for surmounting the challenges of the COVID-19 pandemic in the performing arts education. *Education and Information Technologies*, 26(6), 7635–7655.
<https://doi.org/10.1007/s10639-021-10612-1>
- Lukman Santoso, & Juni Amanullah. (2022). Pengembangan Sistem Informasi Akademik Berbasis Website Menggunakan Metode Rapid Application Development (Rad). *Elkom : Jurnal Elektronika Dan Komputer*, 15(2), 250–259. <https://doi.org/10.51903/elkom.v15i2.943>
- Normah, Rifai, B., Vambudi, S., & Maulana, R. (2022). Analisa Sentimen Perkembangan Vtuber Dengan Metode Support Vector Machine Berbasis SMOTE. *Jurnal Teknik Komputer AMIK BSI*, 8(2), 174–180.
<https://doi.org/10.31294/jtk.v4i2>

- Nurohman, & Nurhayati. (2021). *NUROHMAN & NURHAYATI ACADEMIC INFORMATION SYSTEM USER SATISFACTION MODEL USING TECHNOLOGY ACCEPTANCE MODEL (TAM) WITH END USER COMPUTING SATISFACTION (EUCS) MODIFICATION*.
- Oman Sumantri, S. R. G. W. S. (2015). *102820*. 67, 1–7.
- Panggabean, D. S. O., Buulolo, E., & Silalahi, N. (2020). Penerapan Data Mining Untuk Memprediksi Pemesanan Bibit Pohon Dengan Regresi Linear Berganda. *JURIKOM (Jurnal Riset Komputer)*, 7(1), 56. <https://doi.org/10.30865/jurikom.v7i1.1947>
- Parras-Burgos, D., Fernández-Pacheco, D. G., Barbosa, T. P., Soler-Méndez, M., & Molina-Martínez, J. M. (2020). An augmented reality tool for teaching application in the agronomy domain. *Applied Sciences (Switzerland)*, 10(10). <https://doi.org/10.3390/app10103632>
- Perez, J. G., & Perez, E. S. (2021). Predicting Student Program Completion Using Naïve Bayes Classification Algorithm. *International Journal of Modern Education and Computer Science*, 13(3), 57–67. <https://doi.org/10.5815/IJMECS.2021.03.05>
- Pidie, B. K. (2023). *Jurnal administrasi dan sosial sains*. 2(September), 44–51.
- Prasetyo Utomo, A., & Mariana, N. (2023). *Evaluasi Keberhasilan Sistem Informasi Universitas*. 10(1), 565–579. <http://jurnal.mdp.ac.id>
- Putro, H. F., Vlandari, R. T., & Saptomo, W. L. Y. (2020). Penerapan Metode Naive Bayes Untuk Klasifikasi Pelanggan. *Jurnal Teknologi Informasi Dan Komunikasi (TIKomSiN)*, 8(2). <https://doi.org/10.30646/tikomsin.v8i2.500>
- Putry, N. M. (2022). Komparasi Algoritma Knn Dan Naïve Bayes Untuk

- Klasifikasi Diagnosis Penyakit Diabetes Mellitus. *EVOLUSI: Jurnal Sains Dan Manajemen*, 10(1). <https://doi.org/10.31294/evolusi.v10i1.12514>
- Rahmi, A. N., & Mikola, Y. A. (2021). Implementasi Algoritma Apriori Untuk Menentukan Pola Pembelian Pada Customer (Studi Kasus : Toko Bakoel Sembako). *Information System Journal*, 4(1), 14–19. <https://jurnal.amikom.ac.id/index.php/infos/article/view/561>
- Rais, Z., Hakiki, F. T. T., & Aprianti, R. (2022). Sentiment Analysis of Peduli Lindungi Application Using the Naive Bayes Method. *SAINSMAT: Journal of Applied Sciences, Mathematics, and Its Education*, 11(1), 23–29. <https://doi.org/10.35877/sainsmat794>
- Ramadhani, A. A., Saputra, R. A., & Ningrum, I. P. (2024). Klasifikasi Tingkat Kepuasan Pengguna Google Classroom dalam Pembelajaran Online Menggunakan Algoritma Naïve Bayes. *JIKO (Jurnal Informatika Dan Komputer)*, 8(2), 310. <https://doi.org/10.26798/jiko.v8i2.1221>
- Randi Rian Putra1, C. W. (2018). *Implementasi Data Mining Pemilihan Pelanggan Potensial Menggunakan Algoritma K Means*. 1(1).
- Razak, N., & Nasution, J. (2022). Analisis Efektivitas Penatausahaan Barang Milik Negara Melalui Aplikasi SIMAK-BMN. *ALEXANDRIA (Journal of Economics, Business, & Entrepreneurship)*, 3(2), 39–41. <https://doi.org/10.29303/alexandria.v3i2.177>
- Rinanda, P. D., Delvika, B., Nurhidayarnis, S., Abror, N., & Hidayat, A. (2022). Perbandingan Klasifikasi Antara Naive Bayes dan K-Nearest Neighbor Terhadap Resiko Diabetes pada Ibu Hamil. *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 2(2), 68–75.

<https://doi.org/10.57152/malcom.v2i2.432>

Rizky Fadilla, A., & Ayu Wulandari, P. (2023). Literature Review Analisis Data Kualitatif: Tahap Pengumpulan Data. *Mitita Jurnal Penelitian*, 1(No 3), 34–46.

S Willermark, N Pantic, H. P. (2021). *Subjectively Experienced Time and User Satisfaction: An Experimental Study of Progress Indicator Design in Mobile Application*. <https://hdl.handle.net/10125/71160>

Sajiatmojo, A., Negeri, S., & Selor, T. (2021). *PENGUNAAN E-LEARNING PADA PROSES PEMBELAJARAN DARING*. 1(3), 229.

Setiawan, H., & Novita, D. (2021). Analisis Kepuasan Pengguna Aplikasi KAI Access Sebagai Media Pemesanan Tiket Kereta Api Menggunakan Metode EUCS. *Jurnal Teknologi Sistem Informasi*, 2(2), 162–175. <https://doi.org/10.35957/jtsi.v2i2.1375>

Sholeh, M., Nurnawati, E. K., & Lestari, U. (2023). Penerapan Data Mining dengan Metode Regresi Linear untuk Memprediksi Data Nilai Hasil Ujian Menggunakan RapidMiner. *JISKA (Jurnal Informatika Sunan Kalijaga)*, 8(1), 10–21. <https://doi.org/10.14421/jiska.2023.8.1.10-21>

Sinaga, S., Sembiring, R. W., & Sumarno, S. (2022). Penerapan Algoritma Naive Bayes untuk Klasifikasi Prediksi Penerimaan Siswa Baru. *Journal of Machine ...*, 1(1), 55–64. <https://journal.fkpt.org/index.php/malda/article/view/162%0Ahttps://journal.fkpt.org/index.php/malda/article/download/162/115>

Sobral, S. R. (2020). Mobile learning in higher education: A bibliometric review. *International Journal of Interactive Mobile Technologies*, 14(11), 153–170.

<https://doi.org/10.3991/ijim.v14i11.13973>

- Sudipa, I. G. I., Asana, I. M. D. P., Atmaja, K. J., Santika, P. P., & Setiawan, D. (2023). Analisis Data Kepuasan Pengguna Layanan E-Wallet Gopay Menggunakan Metode Naïve Bayes Classifier Algorithm. *Kesatria : Jurnal Penerapan Sistem Informasi (Komputer Dan Manajemen)*, 4(3), 726–735. <https://tunasbangsa.ac.id/pkm/index.php/kesatria/article/view/219%0Ahttps://tunasbangsa.ac.id/pkm/index.php/kesatria/article/download/219/218>
- Suwandi, S. (2022). Analisis Data Research dan Development Pendidikan Islam. *Journal of Islamic Education El Madani*, 1(1), 1–13. <https://doi.org/10.55438/jiee.v1i1.11>
- Wahyu, A., & Rushenda. (2022). Klasterisasi Dampak Bencana Gempa Bumi. *JEPIN (Jurnal Edukasi Dan Penelitian Informatika)*, 8(1), 175–179. <https://doi.org/10.26418/jp.v8i1>
- Widyadara, M. A. D., & Irawan, R. H. (2019). Implementasi Metode Naïve Bayes dalam Penentuan Tingkat Kesejahteraan Keluarga. *RESEARCH : Computer, Information System & Technology Management*, 2(1), 19. <https://doi.org/10.25273/research.v2i1.4259>
- Wijayanto, C., & Susetyo, Y. A. (2022). *IMPLEMENTASI FLASK FRAMEWORK PADA PEMBANGUNAN APLIKASI SISTEM INFORMASI HELPDESK (SIH)*.
- Wu, W. T., Li, Y. J., Feng, A. Z., Li, L., Huang, T., Xu, A. D., & Lyu, J. (2021). Data mining in clinical big data: the frequently used databases, steps, and methodological models. *Military Medical Research*, 8(1), 1–12. <https://doi.org/10.1186/s40779-021-00338-z>

Zai, C. (2022). Implementasi Data Mining Sebagai Pengolahan Data. *Jurnal Portal Data*, 2(3), 1–12.
<http://portaldata.org/index.php/portaldata/article/view/107>