

***PENERAPAN METODE KLUSTERING DENGAN ALGORITMA K-MEANS
PADA PENGELOMPOKAN PEMINATAN MATA KULIAH DI UNIVERSITAS
MUHAMMADIYAH MAKASSAR***

SKRIPSI

Diajukan Sebagai Salah Satu Syarat untuk Mendapatkan
Gelar Sarjana Komputer (S.Kom) Program Studi Informatika



**PROGRAM STUDI INFORMATIKA
FAKULTAS TEKNIK
UNIVERSITAS MUHAMMADIYAH MAKASSAR**

2025



MAJELIS PENDIDIKAN TINGGI PIMPINAN PUSAT MUHAMMADIYAH
UNIVERSITAS MUHAMMADIYAH MAKASSAR
FAKULTAS TEKNIK



HALAMAN PENGESAHAN

Tugas Akhir ini diajukan untuk memenuhi syarat ujian guna memperoleh gelar Sarjana Komputer (S.Kom) Program Studi Informatika Fakultas Teknik Universitas Muhammadiyah Makassar.

Judul Skripsi : PENERAPAN METODE CLUSTERING DENGAN ALGORITMA
K-MEANS PADA PENGELOMPOKAN PEMINATAN MATA
KULIAH DI UNIVERSITAS MUHAMMADIYAH MAKASSAR

Nama : Muhajirah

Stambuk : 105 84 11014 21

Makassar, 30 Agustus 2025

Telah Diperiksa dan Disetujui
Oleh Dosen Pembimbing;

Pembimbing I

Pembimbing II


Lukman S.Kom., M.T.


Muhyiddin A.M Hayat, S.Kom., M.T.

Mengetahui,

Ketua Prodi Informatika



Rizki Yuslana Bakti, S.T., M.T.

NBM 1307 284





MAJELIS PENDIDIKAN TINGGI PIMPINAN PUSAT MUHAMMADIYAH
UNIVERSITAS MUHAMMADIYAH MAKASSAR
FAKULTAS TEKNIK



PENGESAHAN

Skripsi atas nama Muhajirah dengan nomor induk Mahasiswa 105 84 11014 21, dinyatakan diterima dan disahkan oleh Panitia Ujian Tugas Akhir/Skripsi sesuai dengan Surat Keputusan Dekan Fakultas Teknik Universitas Muhammadiyah Makassar Nomor : 0004/SK-Y/55202/091004/2025, sebagai salah satu syarat guna memperoleh gelar Sarjana Komputer pada Program Studi Informatika Fakultas Teknik Universitas Muhammadiyah Makassar pada hari Sabtu, 30 Agustus 2025.

Panitia Ujian :

Makassar, 6 Rabi'ul Awa 1447 H
30 Agustus 2025 M

1. Pengawas Umum

a. Rektor Universitas Muhammadiyah Makassar

Dr. Ir. H. Abd. Rakhim Nanda, ST, MT, IPU

b. Dekan Fakultas Teknik Universitas Hasanuddin

Prof. Dr. Eng. Muhammad Isran Ramli, S.T., M.T., ASEAN, Eng

2. Penguji

a. Ketua : Dr. Ir. Zahir Zainuddin, M.Sc

b. Sekretaris : Runal Rezklawan B, S.Kom, M.T

3. Anggota

1. Titin Wahyuni, S.Pd, M.T

2. Chyquitha Danuputri, S.Kom, M.Kom

3. Desi Anggreani, S.Kom, M.T

Mengetahui

Pembimbing I

Pembimbing II

Lukman S.Kom.,M.T.

Muhyiddin A M Hayat, S.Kom.,M.T

Dekan



Ir. Muh. Syafaat S. Kuba, S.T., M.T.

NPM : 795 288

Gedung Menara Iqra Lantai 3

Jl. Sultan Alauddin No. 259 Telp. (0411) 866 972 Fax (0411) 865 588 Makassar 90221

Web: <https://teknik.unismuh.ac.id/>, e-mail: teknik@unismuh.ac.id



ABSTRAK

MUHAJIRAH, Penerapan Metode klustering dengan Algoritma *K-means* Pada Pengelompokan Peminatan Mata Kuliah Di Universitas Muhammadiyah Makassar. (dibimbing oleh Lukman S.Kom.,M.T dan Muhyiddin A. M. Hayat, S.Kom.,M.T).

Pemilihan peminatan mata kuliah merupakan keputusan penting yang memengaruhi arah akademik mahasiswa serta kesiapan mereka menghadapi dunia kerja. Di Universitas Muhammadiyah Makassar, proses penentuan peminatan selama ini dilakukan secara manual dan subjektif, sehingga berpotensi menimbulkan ketidaksesuaian antara kemampuan mahasiswa dengan bidang yang dipilih. Penelitian ini bertujuan mengembangkan sistem pengelompokan peminatan mata kuliah menggunakan *algoritma K-Means Clustering* untuk memberikan rekomendasi yang objektif berdasarkan data akademik. Data penelitian diperoleh dari mahasiswa Program Studi Manajemen angkatan 2022, mencakup 47 mata kuliah pada semester 1–5. Tahapan penelitian meliputi data preprocessing, penentuan jumlah kluster optimal dengan *Elbow Method* dan *Silhouette Coefficient*, penerapan algoritma K-Means, serta evaluasi kualitas kluster menggunakan *Silhouette Score* dan *Davies-Bouldin Index*. Hasil penelitian menghasilkan tiga kluster peminatan—Manajemen Keuangan, Manajemen Pemasaran, dan Manajemen Sumber Daya Manusia—dengan nilai *Silhouette Score* sebesar 0,647, *Davies-Bouldin Index* sebesar 0,789, serta akurasi validasi 91,7%. Sistem ini terbukti mampu mengelompokkan mahasiswa berdasarkan kecenderungan akademik, membantu perencanaan kurikulum dan kebijakan akademik, serta memberikan panduan lebih akurat dalam pemilihan peminatan. Sistem yang dikembangkan juga diimplementasikan dalam bentuk aplikasi web menggunakan Streamlit, sehingga mahasiswa dapat memperoleh rekomendasi peminatan secara interaktif dan dosen dapat melakukan validasi dengan lebih efisien.

Kata Kunci: *K-Means Clustering*, peminatan, data mining, Universitas Muhammadiyah Makassar, pendidikan berbasis data

ABSTRACT

MUHAJIRAH, *Application of Clustering Method with K-means Algorithm in Grouping Course Interests at Muhammadiyah University of Makassar. (supervised by Lukman S.Kom.,M.T and Muhyiddin A.M. Hayat, S.Kom.,M.T).*

The selection of course specialization is a crucial decision that affects students' academic direction and readiness for the workforce. At Universitas Muhammadiyah Makassar, this process has traditionally been conducted manually and subjectively, which may result in mismatches between students' abilities and their chosen fields. This study aims to develop a specialization grouping system using the K-Means Clustering algorithm to provide objective recommendations based on academic data. The research data were obtained from 2022 Management students, covering 47 courses taken during semesters 1–5. The stages included data preprocessing, determining the optimal number of clusters using the Elbow Method and Silhouette Coefficient, applying K-Means, and evaluating cluster quality with the Silhouette Score and Davies-Bouldin Index. The results produced three specialization clusters—Financial Management, Marketing Management, and Human Resource Management—with a Silhouette Score of 0.647, a Davies-Bouldin Index of 0.789, and a validation accuracy of 91.7%. The system effectively groups students according to their academic tendencies, supports curriculum planning and academic policy, and helps students in selecting the appropriate specialization. The developed system was implemented as a web application using Streamlit, enabling students to obtain interactive recommendations and lecturers to perform validation more efficiently.

Keywords: *K-Means Clustering, course specialization, data mining, Universitas Muhammadiyah Makassar, educational data mining*

KATA PENGANTAR

Puji dan syukur penulis ucapkan terhadap kehadiran Tuhan Yang Maha Esa, yang telah melimpahkan rahmat hidayah dan karunia-Nya pada penulis, sehingga penulis dapat menyelesaikan penulisan skripsi ini yang berjudul **“Penerapan Metode klustering dengan Algoritma K-means Pada Pengelompokan Peminatan Mata Kuliah Di Universitas Muhammadiyah Makassar”** Shalawat serta salam senantiasa tercurah kepada Rasulullah Shallallahu 'Alaihi Wasallam yang senantiasa menjadi sumber inspirasi dan contoh terbaik bagi umat manusia.

Skripsi ini disusun dengan tujuan untuk mengembangkan sistem pengelompokan peminatan mata kuliah mahasiswa berbasis algoritma K-Means yang mampu bekerja secara otomatis, efisien, dan tepat sasaran. Dengan penerapan metode ini, diharapkan proses penentuan peminatan dapat dilakukan secara objektif berdasarkan data akademik mahasiswa, sehingga meminimalisir kesalahan pemilihan peminatan dan mendukung pengambilan keputusan akademik yang lebih baik.

Pada kesempatan ini dengan segala kerendahan hati, penulis ingin mengucapkan terima kasih sebesar-besarnya kepada semua pihak yang telah banyak membantu dalam penulisan skripsi ini, terutama kepada:

1. Bapak Dr Ir H. Abd. Rakhim Nanda, S.T., M.T. selaku Rektor Universitas Muhammadiyah Makassar
2. Bapak Muh. Syafaat S Kuba, S.T., M.T, selaku Wakil Dekan Fakultas Teknik.
3. Ibu Rizki Yuliana Bakti, S.T., M.T selaku Ketua Program Studi Informatika Fakultas Teknik Universitas Muhammadiyah Makassar
4. Bapak Lukman, S.Kom., M.T., selaku Dosen Pembimbing I, atas segala bimbingan, masukan, dan motivasi yang telah diberikan dengan penuh kesabaran dan perhatian selama penyusunan skripsi ini.

5. Bapak Muhyiddin AM Hayat, S.Kom., M.T selaku Dosen Pembimbing I, atas segala bimbingan, masukan, dan motivasi yang telah diberikan dengan penuh kesabaran dan perhatian selama penyusunan skripsi ini.
6. Kedua orang tua tercinta, Bapak Muhiddin dan Ibunda Suryani, yang senantiasa memberikan doa, dukungan moral, materiil, spiritual dan semangat dalam setiap langkah hidup dan proses yang penulis jalani.
7. Seluruh dosen dan Staff Program Studi Informatika Fakultas Teknik Universitas Muhammadiyah Makassar, atas ilmu dan pengetahuan yang telah diberikan selama masa perkuliahan.
8. Rekan-rekan mahasiswa angkatan 2021 dan terkhususnya Kelas A Program Studi Informatika Fakultas Teknik Universitas Muhammadiyah Makassar, atas kebersamaan, dukungan, dan semangat yang telah menjadi bagian penting dalam perjalanan akademik penulis.
9. Kepada pemilik NIM 105841100121 yang telah kebersamai penulis selama perkuliahan dan sampai penyusunan Tugas Akhir dalam kondisi apapun, telah menjadi *suport sistem* penulis, berkontribusi dalam penyusunan Tugas Akhir ini, memberikan dukungan dan semangat kepada penulis hingga Tugas Akhir ini selesai.
10. Karang Taruna Desa Paraikatte, terima kasih atas semangat dan dukungan dalam menyusun skripsi ini.
11. Semua pihak yang tidak dapat disebutkan satu per satu, yang telah memberikan bantuan dan dukungannya Penulis mengucapkan terima kasih yang sebesar-besarnya.
12. Kepada Mujahidah dan Almarhum Wahidin Wahab Adik penulis. Terima kasih untuk dukungan dan sekaligus jadi penguat untuk bisa melewati semuanya sampai selesai.
13. Untuk Nurul Alwia, Siti Nurkhalisa Sardi dan Muslimah sahabat penulis, terima kasih atas dukungan dan semangatnya sampai selesai.

Penulis menyadari bahwa dalam penulisan skripsi ini masih terdapat kekurangan untuk itu penulis mengharapkan kritik dan saran. Akhir kata, penulis berharap semoga skripsi ini dapat bermanfaat bagi semua pihak demi perkembangan dan kemajuan akademik.

Makassar, 25 Agustus 2025

Penulis,

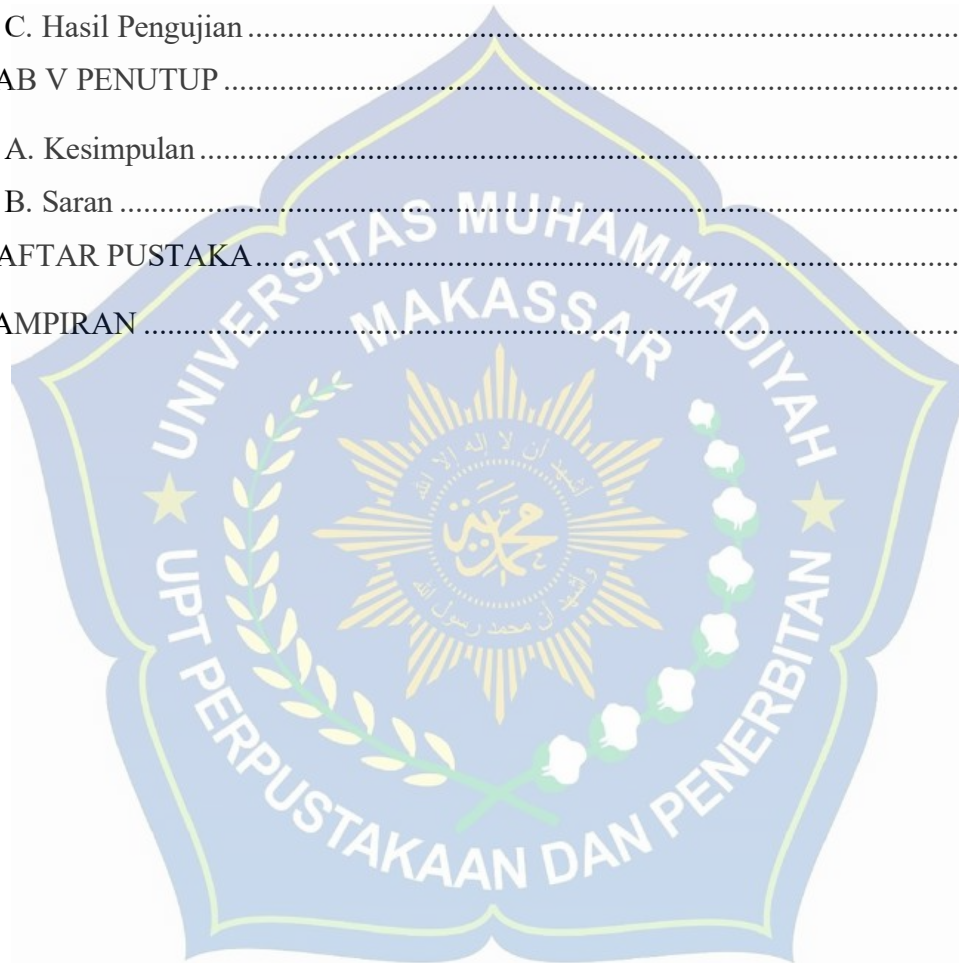
Muhajirah



DAFTAR ISI

ABSTRAK	ii
<i>ABSTRACT</i>	iii
KATA PENGANTAR.....	iv
DAFTAR ISI	vii
DAFTAR TABEL.....	ix
DAFTAR GAMBAR.....	xi
DAFTAR LAMPIRAN	xii
DAFTAR ISTILAH.....	xiii
BAB I PENDAHULUAN.....	1
A. Latar Belakang.....	1
B. Rumusan Masalah.....	4
C. Tujuan Penelitian	4
D. Manfaat Penelitian.....	4
E. Ruang Lingkup Penelitian.....	6
F. Sistematika Penulisan	7
BAB II PEMBAHASAN.....	9
A. Kajian Teori.....	9
B. Penelitian Terkait.....	15
C. Kerangka Berfikir	22
BAB III METODE PENELITIAN	23
A. Tempat dan Waktu Penelitian	23
B. Alat dan Bahan	23
C. Perancangan Sistem.....	23

D. Perancangan <i>InterFace</i>	28
D. Teknik Pengujian Sistem	31
E. Teknik Analisis Data.....	35
BAB IV HASIL DAN PEMBAHASAN.....	38
A. Hasil Pengumpulan dan Preprocessing Data	38
B. Implementasi AntarMuka Web	80
C. Hasil Pengujian	83
BAB V PENUTUP	88
A. Kesimpulan	88
B. Saran	89
DAFTAR PUSTAKA	91
LAMPIRAN	95



DAFTAR TABEL

Tabel 1 Penelitian terkait	19
Tabel 2 Rancangan Pengujian <i>Test case</i>	33
Tabel 3 data training	38
Tabel 4 data testing	39
Tabel 5 distribusi mata kuliah	40
Tabel 6 feature <i>engineering</i> seleksi	42
Tabel 7 Pengembangan <i>fitur engineering</i>	43
Tabel 8 Perbandingan metode <i>Scaling</i>	44
Tabel 9 Perbandingan metode <i>Scaling</i>	45
Tabel 10 Perbandingan Metode <i>Scaling</i>	46
Tabel 11 Implementasi dan Optimisasi <i>K-Means Clustering</i>	47
Tabel 12 Implementasi (Kondisi Aktual di Lapangan)	48
Tabel 13 Hasil Klastering (Prediksi Sistem)	49
Tabel 14 Tabel Validasi Mata Kuliah	50
Tabel 15 Perbandingan Implementasi vs Klastering	52
Tabel 16 Detail Perpindahan (Beda Aktual vs Prediksi)	53
Tabel 17 Metrik klustering <i>clustering</i>	54
Tabel 18 <i>Silhouette Analysis</i> per Cluster	61
Tabel 19 <i>Confusion Matrix Training Data</i>	62
Tabel 20 <i>CONFUSION MATRIX</i>	65
Tabel 21 <i>VARIANCE EXPLAINED</i>	66
Tabel 22 Validasi Model pada Data Testing	68
Tabel 23 Hasil Validasi Testing	70
Tabel 24 Analisis <i>Confidence Score</i>	71
Tabel 25 Ringkasan kesalahan prediksi	73
Tabel 26 pola kesalahan (error patterns)	73

Tabel 27 analisis <i>confidence</i>	73
Tabel 28 analisis <i>centroid</i> cluster.....	75
Tabel 29 Perbandingan Algoritma Clustering.....	76
Tabel 30 Hasil Pengujian Input mahasiswa	84
Tabel 31 Hasil Pengujian Validasi Dosen.....	86
Tabel 32 Hasil Pengujian Interpretasi Sistem.....	87



DAFTAR GAMBAR

Gambar 1 <i>Algoritma K-means</i>	12
Gambar 2 Kerangka berfikir	22
Gambar 3 Lokasi Penelitian	23
Gambar 4 Flowchard tahap pengolahan data	24
Gambar 5 Flowchard Tahap Pengujian Sistem	25
Gambar 6 <i>Input dan output</i> mahasiswa	28
Gambar 7 <i>Input dan output</i> dosen	30
Gambar 8 <i>visualisasi dan analisis clustering</i>	57
Gambar 9 <i>Visualisasi PC1 & PC2</i>	59
Gambar 10 Giagram <i>error analysis</i>	74
Gambar 11 Halaman <i>Input</i> Mahasiswa	80
gambar 12 halaman <i>output</i> mahasiswa	81
Gambar 13 halaman <i>input</i> dosen	82
Gambar 14 Halaman <i>Output</i> Dosen	83

DAFTAR LAMPIRAN

Lampiran 1 Data Mentah.....	95
Lampiran 2 Source Code Sistem prediksi peminatan mahasiswa	99
Lampiran 3 Permohonan Penelitian Kepada Ketua Program Studi Informatika.....	147
Lampiran 4 Permohonan Penelitian Kepada Ketua LP3M Unismuh Makassar.....	148
Lampiran 5 Permohonan Penelitian Kepada Fakultas Ekonomi Bisnis	149
Lampiran 6 Hasil Turnitin	150



DAFTAR ISTILAH

Algoritma	Langkah-langkah sistematis yang digunakan untuk menyelesaikan suatu permasalahan. Dalam penelitian ini, algoritma yang digunakan adalah <i>K-Means Clustering</i> untuk pengelompokan data peminatan mahasiswa.
Clustering (Pengelompokan)	Metode analisis data untuk mengelompokkan objek berdasarkan kemiripan karakteristiknya. Data yang mirip dikelompokkan ke dalam satu kluster, sedangkan yang berbeda dimasukkan ke kluster lain.
Elbow Method	Metode untuk menentukan jumlah kluster optimal dengan melihat titik “siku” pada grafik <i>Within-Cluster Sum of Squares (WCSS)</i> .
Silhouette Score	Metrik evaluasi untuk mengukur kualitas kluster. Nilainya berkisar antara -1 hingga 1; semakin mendekati 1 berarti kualitas kluster semakin baik.
Davies-Bouldin Index (DBI)	Metrik evaluasi kluster yang menghitung rasio antara jarak antar-kluster dan ukuran dalam-kluster. Nilai lebih rendah menunjukkan hasil kluster yang lebih baik.
Centroid	Titik pusat suatu kluster dalam algoritma <i>K-Means</i> . Posisi centroid diperbarui secara iteratif hingga mencapai konvergensi.
Data Mining	Proses menemukan pola atau informasi penting dari kumpulan data besar melalui analisis otomatis.

K-Means Clustering	Algoritma <i>unsupervised learning</i> yang digunakan untuk membagi data ke dalam sejumlah klaster berdasarkan jarak terdekat ke pusat klaster (centroid).
Streamlit	Framework Python open-source yang digunakan untuk membuat aplikasi web interaktif berbasis data secara sederhana dan cepat.
SQLite	Sistem manajemen basis data relasional yang ringan, berbasis file, dan sering digunakan dalam aplikasi berskala kecil hingga menengah.
Peminatan	Fokus atau konsentrasi bidang studi yang dipilih mahasiswa untuk memperdalam keahlian tertentu, misalnya Manajemen Keuangan, Manajemen Pemasaran, dan Manajemen SDM.

BAB I

PENDAHULUAN

A. Latar Belakang

Perkembangan dunia pendidikan tinggi yang semakin pesat mendorong perguruan tinggi untuk menawarkan berbagai mata kuliah dan jalur peminatan guna memenuhi kebutuhan akademik dan minat mahasiswa. Di Universitas Muhammadiyah Makassar, mahasiswa memiliki kesempatan untuk memilih peminatan atau konsentrasi studi tertentu yang sesuai dengan minat dan tujuan akademiknya. Peminatan ini bertujuan untuk memfokuskan mahasiswa pada bidang ilmu tertentu, sehingga proses belajar menjadi lebih terarah dan mendalam, sekaligus mempermudah dalam menentukan topik tugas akhir.

Peminatan atau konsentrasi merupakan fokus mahasiswa terhadap suatu bidang studi tertentu sesuai dengan minatnya. Tujuannya yaitu untuk lebih memfokuskan mahasiswa terhadap ilmu yang didapat dari mata kuliah sebelumnya, agar dapat lebih terarah. Disamping itu juga peminatan bertujuan mengarahkan para mahasiswa, dalam memberikan beberapa pilihan mengenai topik penelitian tugas akhir. (Adrianto & Fahmi, 2019)

Pemilihan konsentrasi dalam kegiatan akademik mahasiswa memang bukan hal yang mudah karena tergantung pada minat, bakat dan keinginan, oleh karena itu perlu pertimbangan yang matang supaya mahasiswa tidak salah dalam memilih konsentrasi yang diinginkan. Hal ini sering terjadi ketika mahasiswa semester akhir mengerjakan tugas akhir namun tidak sesuai dengan kemampuan bidang yang dimiliki. Pemilihan konsentrasi yang asal-asalan tanpa pertimbangan yang matang, menyebabkan dampak negatif pada mahasiswa, yaitu kesulitan dalam penyerapan materi-materi perkuliahan. Oleh sebab itu perlu metode khusus yang dapat digunakan mahasiswa

dalam menentukan konsentrasi mahasiswa. Salah satu metode yang digunakan adalah metode K-Means. (Karmanita et al., 2023)

Saat ini sebagian besar Mahasiswa masih belum tahu konsentrasi apa yang akan dipilih karena belum mengetahui kemampuan ataupun minat yang akan mereka pilih. Pemilihan konsentrasi sangat penting dan perlu dipertimbangkan secara matang untuk menentukan konsentrasi apa yang akan mereka miliki pada saat mereka telah menyelesaikan perkuliahan. Beberapa kasus yang sering terjadi mahasiswa memilih konsentrasi yang kurang tepat, sehingga menyebabkan mahasiswa kesulitan dalam mengikuti dan memahami materi-materi yang diberikan, dan juga hal ini mempengaruhi proses pembuatan tugas akhir di mana hasil yang didapat kurang maksimal dikarenakan kurang tepatnya di dalam pemilihan konsentrasi (Elsi & Haryanto, 2022)

Universitas Muhammadiyah Makassar menghadapi tantangan dalam penerapan metode K-Means untuk pengelompokan minat mahasiswa, terutama terkait kualitas data yang tidak lengkap dan sistem manual yang mengakibatkan pemilihan minat kurang akurat. Hal ini sejalan tentang tantangan adopsi learning analytics di pendidikan tinggi, yang menyoroti masalah keterbatasan data, resistensi budaya terhadap perubahan, dan kurangnya kepemimpinan terpusat. menekankan bahwa institusi di tahap awal adopsi (seperti UMM) sering mengalami kendala dalam infrastruktur data dan kesiapan stakeholder, sementara institusi yang lebih matang justru menghadapi tantangan evaluasi teknologi. Untuk memperbaiki sistem pemilihan minat, UMM dapat menerapkan rekomendasi dari ((Alzahrani et al., 2023), seperti peningkatan kualitas data, pelatihan analitik, dan pendekatan kolaboratif dengan stakeholder. Dengan demikian, solusi berbasis learning analytics dapat dioptimalkan untuk mendukung keputusan akademik yang lebih tepat.

Untuk mengatasi kendala tersebut, diperlukan algoritma dan sistem komputer yang dapat mendukung proses pengelompokan data siswa sehingga

kelas yang dipilih sesuai dengan minat siswa. Algoritma K-means efektif dalam mengidentifikasi kelompok atau cluster dari data yang tidak bertabel. Algoritma ini bersifat fleksibel dalam menyesuaikan jumlah kelompok, efisien, dan cepat dalam menangani dataset besar, terutama pada data numerik seperti nilai akademik. (Rose, C. D., et al., 2025)

Berbagai algoritma cerdas untuk menganalisa telah banyak diimplementasikan. Salah satunya adalah algoritma k-means yang merupakan algoritma pengelompokan interaktif yang sederhana diimplementasikan, relatif cepat, dan mudah beradaptasi. Keterhubungan nilai MKW (Mata Kuliah Wajib) dan TA (Tugas Akhir) sangat perlu di analisa agar mahasiswa, dosen dan program studi dapat mengetahui kondisi setiap mahasiswa. Berbagai penelitian yang menggunakan algoritma K-means telah dilakukan oleh penelitian (Haviluddin et al., 2021) telah menggunakan algoritma K-means untuk pengelompokan.

Merujuk pada permasalahan yang telah diuraikan sebelumnya, penelitian ini bertujuan untuk mengelompokkan peminatan mata kuliah mahasiswa secara sistematis dan objektif menggunakan metode Kekelompokan dengan algoritma K-means. Pendekatan ini mempertimbangkan variasi karakteristik akademik dan preferensi mahasiswa sebagai dasar pengelompokan. Fokus utama dari penelitian ini adalah mengoptimalkan penerapan algoritma K-means dalam pengelompokan data peminatan, guna menghasilkan klaster yang merepresentasikan kecenderungan minat mahasiswa terhadap kelompok mata kuliah tertentu. Penelitian ini diharapkan dapat memberikan kontribusi dalam proses pengambilan keputusan akademik, khususnya dalam memberikan rekomendasi peminatan yang sesuai bagi mahasiswa berdasarkan pola yang terbentuk dari data historis. Selain itu, penelitian ini juga akan mengevaluasi jumlah cluster optimal dan kualitas hasil pengelompokan menggunakan matrik evaluasi seperti *Skor Siluet* dan visualisasi data. Hasil dari pengelompokan ini

dapat dimanfaatkan oleh pihak program studi dalam menyusun strategi pembelajaran, penyeimbangan kurikulum, serta penyesuaian sumber daya pengajaran berdasarkan distribusi minat mahasiswa.

B. Rumusan Masalah

Berdasarkan latar belakang yang telah diuraikan, maka rumusan masalah dalam penelitian ini adalah sebagai berikut:

1. Bagaimana menerapkan metode clustering dengan algoritma K-Means untuk mengelompokkan peminatan mata kuliah mahasiswa di Universitas Muhammadiyah Makassar?
2. Berapa jumlah Kluster yang optimal untuk mengelompokkan peminatan mata kuliah mahasiswa?
3. Bagaimana hasil analisis dan interpretasi dari pengelompokan peminatan mata kuliah menggunakan algoritma K-Means tersebut?

C. Tujuan Penelitian

Adapun tujuan utama yang dicapai dalam penelitian ini dapat dirumuskan sebagai berikut:

1. Untuk mengelompokkan peminatan mata kuliah mahasiswa berdasarkan variabel-variabel tertentu (seperti nilai akademik, minat, atau prestasi) di Universitas Muhammadiyah Makassar.
2. Untuk menentukan jumlah klaster berdasarkan metode evaluasi seperti *Elbow Method*, *Silhouette Coefficient*, atau *Gap Statistic*.
3. Untuk memahami karakteristik setiap kelompok peminatan serta memberikan rekomendasi strategis kepada pihak universitas dalam pengelolaan peminatan mata kuliah.

D. Manfaat Penelitian

Adapun manfaat dari penelitian ini mencakup berbagai aspek yang dapat

memberikan kontribusi bagi berbagai pihak, sebagai berikut:

1. Manfaat bagi Universitas (Bidang Akademik dan Kebijakan)

- a) Optimasi Penyaluran Peminatan: Membantu pihak universitas dalam mengelompokkan peminatan mata kuliah secara objektif berbasis data, sehingga penempatan mahasiswa lebih sesuai dengan kompetensi dan minat mereka.
- b) Penyusunan Kurikulum yang Adaptif: Hasil clustering dapat menjadi acuan untuk mengevaluasi atau merancang struktur kurikulum yang lebih relevan dengan kebutuhan mahasiswa dan pasar kerja.
- c) Pengambilan Keputusan Berbasis Data: Memberikan dasar empiris bagi fakultas/prodi dalam menyusun kebijakan akademik, seperti pembukaan kelas peminatan baru atau penyesuaian kuota berdasarkan tren cluster.

2. Manfaat bagi Mahasiswa

- a) Peminatan yang Lebih Tepat: Mahasiswa dapat memilih peminatan yang sesuai dengan kekuatan akademik (nilai prasyarat, IPK) dan minat pribadi, mengurangi risiko salah jurusan atau ketidakpuasan.
- b) Peningkatan Prestasi Akademik: Dengan peminatan yang sesuai, mahasiswa diharapkan lebih termotivasi dan berkinerja lebih baik secara akademis.
- c) Transparansi Proses Peminatan: Sistem berbasis data mengurangi subjektivitas dalam penentuan peminatan, sehingga mahasiswa merasa lebih adil.

3. Manfaat bagi Bidang Teknologi dan Keilmuan

- a) Validasi Aplikasi *Unsupervised Learning* dalam Pendidikan: Menguji efektivitas algoritma K-Means dalam konteks pengelompokan data akademik, termasuk tantangan *noise* data (seperti nilai yang tidak konsisten) dan pemilihan fitur optimal.

- b) Pengembangan Model yang Replikabel: Metode yang diusulkan dapat diadaptasi oleh perguruan tinggi lain dengan karakteristik data serupa.
- c) Kontribusi pada Riset *Educational Data Mining* (EDM): Menambah referensi studi kasus penerapan *machine learning* untuk manajemen pendidikan di Indonesia, khususnya di lingkungan Universitas Muhammadiyah.

4. Manfaat bagi Dosen dan Tenaga Kependidikan

- a) Efisiensi Administrasi: Mempermudah proses pengelolaan peminatan (seperti pembagian kelas atau penugasan dosen pengampu) melalui rekomendasi sistem otomatis.
- b) Identifikasi Potensi Mahasiswa: Dosen dapat menggunakan hasil clustering untuk memberikan bimbingan akademik yang lebih terarah berdasarkan profil kelompok mahasiswa.

E. Ruang Lingkup Penelitian

Agar penelitian ini terarah dan fokus, maka ditetapkan beberapa batasan sebagai berikut:

1. Objek Penelitian

- a) Data mahasiswa Program Studi Manajemen

Data mahasiswa Program Studi (Prodi) Manajemen merupakan kumpulan informasi terstruktur yang mencakup aspek akademik, administratif, dan personal dari mahasiswa yang terdaftar di prodi tersebut. Data ini digunakan untuk berbagai keperluan seperti pemantauan perkembangan akademik, pengambilan kebijakan, dan analisis strategis.

- b) Variabel yang digunakan :

1. Nilai mata kuliah wajib.
2. Indeks Prestasi Kumulatif (IPK)

3. Prestasi non-akademik (opsional)

2. Metodologi yang Digunakan

- a) *Algoritma K-Means Clustering* sebagai metode utama.
- b) Teknik penentuan jumlah cluster (*Elbow Method, Silhouette Analysis*).
- c) Tools analisis: Python
- d) Dataset periode Mahasiswa Prodi Manajemen angkatan 2022

3. Batasan Penelitian

- a) Tidak mencakup analisis penyebab perbedaan kluster (hanya fokus pada pengelompokan).
- b) Data yang digunakan hanya dari sumber internal kampus
- c) Tidak membandingkan dengan algoritma lain

4. Output yang Diharapkan

- a) Pola pengelompokan peminatan berdasarkan variabel terpilih.
- b) Rekomendasi untuk kampus dalam penentuan peminatan mata kuliah.

F. Sistematika Penulisan

Penelitian ini disusun dalam beberapa bab dengan uraian sebagai berikut:

BAB I PENDAHULUAN

Bab ini memuat latar belakang, rumusan masalah, tujuan penelitian, manfaat penelitian, ruang lingkup penelitian, serta sistematika penulisan sebagai dasar pengantar penelitian.

BAB II TINJAUAN PUSTAKA

Bab ini membahas landasan teori yang relevan, kajian penelitian terdahulu, serta kerangka berpikir yang digunakan untuk merumuskan arah dan dasar konseptual dari penelitian.

BAB III METODE PENELITIAN

Bab ini menguraikan metode yang digunakan dalam penelitian, mencakup pendekatan penelitian, teknik pengumpulan data, serta tahapan analisis yang dilakukan untuk mencapai tujuan penelitian.

BAB IV HASIL DAN PEMBAHASAN

Bab ini menyajikan hasil penelitian dan pembahasan mengenai penerapan metode K-Means Clustering pada pengelompokan peminatan mata kuliah mahasiswa, mulai dari hasil preprocessing data, penentuan jumlah kluster optimal, implementasi model, evaluasi kualitas kluster, hingga interpretasi hasil pengelompokan.

BAB V PENUTUP

Bab ini berisi kesimpulan yang diperoleh dari hasil penelitian serta saran yang dapat menjadi masukan untuk penelitian selanjutnya.



BAB II

PEMBAHASAN

A. Kajian Teori

1. Data Mining

Data mining merupakan proses eksplorasi otomatis terhadap informasi yang bernilai dari kumpulan data berukuran besar. Teknik ini digunakan untuk menganalisis basis data dalam skala besar guna menemukan pola-pola tersembunyi yang sebelumnya tidak diketahui namun memiliki potensi untuk memberikan wawasan baru dan bermanfaat. Meskipun demikian, tidak semua aktivitas pencarian informasi dalam data dapat dikategorikan sebagai data mining. Terdapat berbagai pendekatan dalam data mining, salah satunya yang umum digunakan adalah Knowledge Discovery in Databases (KDD), yang juga sering dikaitkan dengan konsep pattern recognition atau pengenalan pola. Pendekatan ini menekankan pada proses sistematis dalam menemukan pengetahuan baru dari data melalui tahapan-tahapan seperti pemilihan data, pembersihan, transformasi, data mining itu sendiri, hingga evaluasi dan interpretasi hasil (Sibarani, 2020)

Dalam praktiknya, data mining memiliki berbagai teknik atau metode yang dapat digunakan untuk mengidentifikasi pola atau struktur dalam data. Salah satu metode yang cukup populer dan banyak digunakan adalah algoritma K-Means, yang berfungsi untuk melakukan pengelompokan data (clustering) berdasarkan kemiripan karakteristik. (Faisal et al., 2025)

Data mining sering diartikan sebagai proses penggalian informasi yang tersembunyi di dalam kumpulan data berskala besar yang tersimpan dalam basis data. Proses ini melibatkan pengumpulan, analisis, dan transformasi data dalam jumlah besar menjadi informasi yang bermakna dan dapat dimanfaatkan dalam

berbagai konteks, termasuk dalam mendukung pengambilan keputusan. (Basalamah & Setyadi, 2023)

2. *Clustering*

Clustering/Pengelompokan adalah teknik analisis data yang melibatkan pengelompokan objek yang memiliki pengetahuan terbatas atau tidak ada sebelumnya tentang hubungan antara objek dalam set data. Tujuan utama pengelompokan adalah untuk mengungkap kelas atau struktur yang melekat dalam data. Selain itu, pengelompokan sering digambarkan sebagai pengelompokan data yang tidak berlabel dengan sedikit atau tanpa panduan manusia ke dalam kelompok-kelompok yang berbeda. Kelompok-kelompok ini dibentuk sehingga objek yang termasuk dalam cluster yang sama menunjukkan karakteristik yang sama, yang membedakannya dari objek di cluster lain. Pengelompokan juga termasuk dalam pembelajaran tanpa pengawasan, aspek integral dari pembelajaran mesin. Ini melibatkan penggunaan algoritma untuk mengidentifikasi pola dalam set data, yang dapat diperoleh melalui pengamatan langsung atau pembuatan simulasi data. Seperti yang dinyatakan dalam proses pembelajaran dalam pengelompokan melibatkan upaya untuk mengklasifikasikan pengamatan data atau variabel independen tanpa pengetahuan sebelumnya tentang variabel target tertentu. (Alasali & Ortakci, 2024)

Clustering digunakan untuk mengukur jarak terjauh. Clustering juga dapat diartikan sebagai cara untuk mengukur jarak data dari yang terjauh. K-Means membagi sekumpulan sampel menjadi k kelompok atau cluster terpisah dengan menggunakan nilai rata-rata anggota sebagai indikator utama. Titik ini biasanya disebut Centroid, mengacu pada entitas aritmatika dengan nama yang sama, dan direpresentasikan sebagai vektor dalam ruang dimensi yang sembarang. (Yahya & Kurniawan, 2025)

3. Silhouette Coefficient

Metode Silhouette Coefficient pertama kali diusulkan oleh Peter J. Rousseeuw. Metode ini menggabungkan dua faktor, yaitu kohesi dan resolusi. Kohesi adalah kesamaan antara objek dan kluster.

$$s(i) = \frac{b(i) - a(i)}{\max\{a(i), b(i)\}}$$
$$= \begin{cases} 1 - \frac{a(i)}{b(i)}, & a(i) < b(i) \\ 0, & a(i) = b(i) \\ \frac{b(i)}{a(i)} - 1, & a(i) > b(i) \end{cases}$$

Jika dibandingkan dengan kluster lain, maka disebut pemisahan. Perbandingan ini dicapai dengan nilai Silhouette yang berada pada rentang $-1-1$. Nilai Silhouette mendekati 1, yang menunjukkan bahwa terdapat hubungan yang erat antara objek dan kluster. Jika kluster data dalam suatu model dihasilkan dengan nilai Silhouette yang relatif tinggi, maka model tersebut sesuai dan dapat diterima. Perhitungannya adalah sebagai berikut: (Yuan & Yang, 2019)

4. Davies-Bouldin Index

Indeks Davies Bouldin (DBI) merupakan ukuran untuk mengevaluasi kinerja pengelompokan. DBI memiliki korelasi positif untuk kasus “dalam kelas” dan korelasi negatif untuk kasus “antarkelas”. Gunakan DBI sebagai metrik pengelompokan karena cara umum pengelompokannya. Validasi berisi dua kategori utama - validasi eksternal dan validasi internal yang digunakan untuk menilai kinerja hasil pengelompokan (Wijaya et al., 2021)

5. Elbow Method

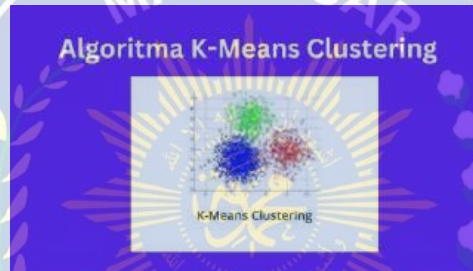
Elbow method adalah untuk memilih nilai k yang kecil dan masih memiliki nilai withinss yang rendah. Penentuan optimal jumlah cluster pada penelitian ini menggunakan salah satu metode analisis cluster yaitu *method Elbow*, dengan

memperhatikan nilai perbandingan (dari perhitungan SSE untuk setiap nilai cluster) antara jumlah cluster yang akan membentuk siku pada suatu titik, sehingga semakin besar jumlah cluster k maka nilai SSE akan semakin kecil. Rumus SSE sebagai berikut .

$$SSE = \sum_{k=1}^k \sum_{x_1 \in S_k} \|x_1 - C_k\|_2^2$$

Berdasarkan persamaan (1) di atas, jumlah cluster adalah 3, kemudian mencari nilai jarak kedua objek terdekat dengan menggunakan persamaan. (Hartanti, 2020)

6. Algoritma *K-Means*



Gambar 1 Algoritma *K-means*

Algoritma *K-Means* merupakan salah satu metode clustering non-hierarchical dengan cara mempartisi data yang ada ke dalam satu atau beberapa cluster. Metode ini mempartisi data ke dalam cluster-cluster, sehingga power yang memiliki karakteristik yang sama akan berada dalam 1 kluster. Sedangkan segmentasi citra berarti mempartisi citra ke dalam berbagai kelompok atau group dengan fitur yang sama atau memiliki beberapa kesamaan. K-means merupakan metode pengelompokan partisi citra ke dalam kelompok sedemikian rupa sehingga minimal satu bagian cluster berisi citra dengan luas area utama bagian yang digunakan (Oktaviani & Bahtiar, 2023)

Algoritma *K-Means* klusterung telah banyak digunakan di berbagai sektor seperti informasi bisnis, kesehatan, dan teknologi. Di sektor asuransi, K-Means

digunakan untuk mendeteksi anomali pada data klaim guna mengidentifikasi potensi kecurangan, sehingga membantu mengurangi kerugian akibat klaim palsu. Selain itu, K-Means juga diaplikasikan dalam segmentasi nasabah berdasarkan data demografi dan perilaku, yang memungkinkan perusahaan untuk mengembangkan strategi pemasaran yang lebih terarah. Di asuransi kesehatan, algoritma ini digunakan untuk mengelompokkan nasabah berdasarkan tingkat risiko sektor, guna mendukung penentuan premi dan manajemen risiko yang lebih akurat. (Alya Putri & Rahmah, 2024)

7. Peminatan Mata Kuliah

Peminatan akademik merupakan jalur khusus dalam pendidikan tinggi yang dirancang untuk membantu mahasiswa mengembangkan pengetahuan dan keterampilan sesuai minat dan bakat mereka. Program ini bertujuan memperdalam keahlian di bidang tertentu sebagai pelengkap kurikulum inti, serta mendukung kesiapan mahasiswa menghadapi dunia kerja maupun pendidikan lanjutan.

Peminatan juga memfasilitasi pemilihan mata kuliah berdasarkan minat dan tujuan karir mahasiswa, dengan tetap mengacu pada struktur kurikulum yang berlaku. Penelitian ini bertujuan mengevaluasi relevansi antara peminatan di bidang Teknologi Informasi dan kebutuhan kompetensi di industri. Melalui analisis kesesuaian antara kurikulum dan ekspektasi dunia kerja, hasil penelitian diharapkan dapat memberikan rekomendasi untuk pengembangan program peminatan yang lebih adaptif terhadap perkembangan teknologi, serta membantu mahasiswa memilih jalur studi dengan prospek karir yang lebih baik. (Reza et al., 2025)

Proses pemilihan peminatan di program studi menunjukkan bahwa minat pribadi menjadi faktor utama dalam pengambilan keputusan mahasiswa, khususnya pada peminatan perhotelan. Mahasiswa yang memilih jalur ini

umumnya memiliki ketertarikan kuat terhadap bidang tersebut. Selain itu, dukungan keluarga juga berperan penting dalam membangun keyakinan mahasiswa untuk memilih peminatan yang sesuai dengan minat dan bakat mereka.

Hasil wawancara mengungkap bahwa keberhasilan dan ketekunan mahasiswa dalam mengikuti peminatan sangat dipengaruhi oleh sejauh mana minat mereka terakomodasi, serta adanya keterlibatan aktif dalam memahami aspek-aspek peminatan yang diminati (Dosista et al., 2023).

5. *Streamlit*

Streamlit adalah pustaka Python sumber terbuka yang menyederhanakan proses pembuatan aplikasi web interaktif berbasis data. Pada tahun 2018, *Streamlit* didirikan bersama oleh Adrien Treuille, Amanda Kelly, dan Thiago Teixeira sebagai alat antarmuka pengguna untuk Python. *Streamlit* secara khusus membantu menjadikan pembelajaran mesin dapat diulang, dibagikan, diadaptasi, dan digunakan di seluruh organisasi. Khususnya, pengguna dapat berinteraksi dengan model pembelajaran mesin dengan memberikan data masukan, menyesuaikan parameter pembelajaran melalui berbagai widget masukan yang tersedia, dan akhirnya menerima hasil keluaran melalui beberapa keluaran (misalnya visualisasi plot, DataFrame, dll.). Di balik layar, frontend *Streamlit* dibuat menggunakan React, sementara di backend *Streamlit* berjalan di server web Tornado (Asroni & Adrian, 2015).

6. *SQLite*

Pada *SQLite*, seluruh perintah query harus dibuat secara manual. Selain itu, pengembang perlu membangun *Data Access Object* (DAO) secara mandiri, serta membuat basis data dan tabel secara tersendiri. Seluruh perintah query dan arsitektur pembuatan basis data perlu dipahami dan dirancang secara manual,

sehingga kode program cenderung lebih panjang, kompleks, dan memerlukan pemahaman yang mendalam(Putra et al., 2020).

7. Aplikasi Berbasis Web

Aplikasi web merupakan program yang menggunakan teknologi browser untuk beroperasi dan dapat diakses melalui jaringan komputer. Menurut .(Mina, 2023)aplikasi web adalah program yang disimpan di server, dikirim melalui internet, dan diakses melalui antarmuka browser. Dengan demikian, dapat disimpulkan bahwa aplikasi web merupakan perangkat lunak yang dapat diakses melalui web browser menggunakan jaringan internet atau intranet. Selain itu, aplikasi web juga merupakan perangkat lunak komputer yang dikodekan dalam bahasa pemrograman yang mendukung pengembangan perangkat lunak berbasis web.

B. Penelitian Terkait

Beberapa penelitian sebelumnya yang terkait meliputi:

1. *Analisis data set peminatan siswa menggunakan algoritma K-means dengan optimize parameter di sekolah menengah kejuruan (Studi Kasus: Smk Pui Gegesik)* (Sidik et al., 2023)

Penelitian ini berfokus pada pengelompokan data berdasarkan minat siswa dengan menggunakan metode *K-Means*, salah satu teknik *clustering* dalam data mining. Tujuannya adalah untuk mengelompokkan siswa sesuai dengan minat mereka, seperti pada mata pelajaran tertentu, misalnya matematika. Hasil pengelompokan ini diharapkan dapat membantu pihak sekolah atau pengambil kebijakan dalam menyusun strategi pembelajaran atau peminatan yang lebih efektif.

Namun, tantangan utama dalam penerapan algoritma *K-Means* adalah tidak adanya jaminan terhadap keakuratan klaster yang dihasilkan. Oleh karena

itu, diperlukan metode validasi untuk mengukur kualitas hasil pengelompokan. Salah satu metode yang sering digunakan adalah Davies-Bouldin Index (DBI), yang menilai kualitas klaster dengan cara memaksimalkan jarak antar-klaster (inter-cluster) dan meminimalkan jarak dalam klaster (intra-cluster). Semakin rendah nilai DBI, maka semakin baik kualitas pengelompokannya.

Contoh penerapan metode ini dapat dilihat dalam penelitian oleh Yudhistira Arie Wijaya, yang menggunakan algoritma *K-Means* dan DBI untuk memetakan fasilitas pendidikan di berbagai desa. Sumber datanya berasal dari portal resmi pemerintah, mencakup jumlah desa dengan fasilitas pendidikan tingkat SMA dan SMK. Hasil pengujian menunjukkan bahwa konfigurasi jumlah klaster terbaik adalah $k = 2$, dengan nilai DBI sebesar 0,168. Pada klaster ini, wilayah dibagi menjadi dua kategori: rendah (L0) yang mencakup 31 provinsi, dan tinggi (L1) yang mencakup 3 provinsi, dengan centroid akhir menunjukkan perbedaan signifikan dalam jumlah fasilitas pendidikan.

Dengan pendekatan serupa, data nilai siswa di sekolah juga dapat dianalisis menggunakan teknologi data mining untuk menggali pola-pola yang selama ini sulit ditemukan secara manual. Melalui algoritma *K-Means*, minat siswa terhadap suatu bidang, seperti matematika, dapat diklasifikasikan, sehingga pihak sekolah bisa menyusun program atau kebijakan yang lebih sesuai dengan kebutuhan dan potensi siswa.

2. Klasterisasi Hasil Belajar Matematika dengan Algoritma *K-Means Clustering* (Febrinita et al., 2023)

Penelitian ini bertujuan untuk mengelompokkan capaian belajar matematika mahasiswa menggunakan *algoritma K-Means Clustering*, dengan pendekatan deskriptif kuantitatif. Fokus utama penelitian ini adalah menjelaskan proses clustering dan menginterpretasikan hasilnya untuk mengetahui pola-pola

yang muncul pada prestasi belajar mahasiswa. Algoritma *K-Means* digunakan untuk mengelompokkan mahasiswa berdasarkan data capaian belajar untuk mata kuliah yang terkait dengan konsep matematika. Setelah proses clustering dilakukan, hasilnya selanjutnya dianalisis menggunakan statistik deskriptif untuk mengetahui nilai rata-rata dan kecenderungan setiap kelompok (cluster) yang terbentuk. Data yang digunakan dalam penelitian ini adalah data sekunder yang diambil dari sistem informasi akademik (SIKAD) FTI Universitas Islam Balitar (Unisba). Data tersebut meliputi nilai mahasiswa angkatan 2021 pada mata kuliah berikut: Matematika Komputasi, Matematika Komputasi Lanjut, Logika Informatika, dan Statistika. Pemilihan intake 2021 ini didasarkan pada pertimbangan bahwa saat ini mereka berada di semester 4 dan akan menentukan minatnya di semester 5. Sebanyak 51 data mahasiswa digunakan dalam proses clustering. Selain nilai akademik, informasi tambahan seperti sekolah asal dan minat sebelumnya di sekolah juga disertakan sebagai variabel pendukung dalam proses pengelompokan. Melalui penelitian ini, diharapkan dapat diperoleh gambaran cluster capaian belajar mahasiswa yang dapat digunakan sebagai pertimbangan dalam pengambilan keputusan akademik, khususnya dalam proses penentuan minat.

3. Penerapan Metode *K-Means* untuk Klasifikasi Rekomendasi Tugas Akhir Mahasiswa (Ramdani et al., 2025)

Dalam penelitian ini digunakan data simulasi yang melibatkan 5 mahasiswa, dimana setiap mahasiswa memiliki nilai pada tiga mata kuliah, yaitu A, B, dan C. Tujuan dari analisis klusterisasi ini adalah untuk mengelompokkan mahasiswa berdasarkan pola nilai mereka sehingga dapat memetakan karakteristik akademik masing-masing kelompok ke dalam area penelitian yang sesuai.

Proses klusterisasi dilakukan dengan membagi data menjadi tiga kelompok (klaster), yaitu:

- a) Klaster 1 (C1) untuk mahasiswa dengan nilai rendah,

- b) Klaster 2 (C2) untuk mahasiswa dengan nilai sedang,
- c) Klaster 3 (C3) untuk mahasiswa dengan nilai tinggi.

Untuk mengevaluasi kualitas dan kesesuaian klaster yang terbentuk, digunakan dua metrik validasi, yaitu *Silhouette Coefficient* (SC) dan *Sum of Squared Errors* (SSE). Pengujian ini bertujuan untuk memastikan bahwa pembagian klaster cukup representatif dan memiliki tingkat akurasi yang baik. Dengan pendekatan ini, penelitian ini diharapkan dapat membantu mengoptimalkan proses pemilihan topik tugas akhir bagi mahasiswa, sehingga mereka dapat menentukan topik yang sesuai dengan kemampuan akademiknya sekaligus meningkatkan kualitas hasil penelitiannya.

4. Penerapan Metode *K-Means Clustering* Dalam Mengetahui Minat Siswa Terhadap Mata Pelajaran Matematika Di SMP Negeri 19 Kota Bengkulu (Di et al., 2024)

Penelitian ini mengimplementasikan metode *K-Means Clustering* untuk menganalisis minat siswa terhadap mata pelajaran matematika di SMP Negeri 19 Kota Bengkulu. Dengan metode ini, guru dapat lebih mudah mengenali tingkat minat siswa, yang nantinya dapat digunakan sebagai bahan evaluasi dalam proses pembelajaran untuk meningkatkan minat belajar khususnya pada mata pelajaran matematika.

Data yang digunakan berasal dari 34 siswa pada semester ganjil tahun ajaran 2023/2024. Dari hasil clustering diperoleh dua kelompok siswa yaitu:

- a) Cluster 1 (C1) yang menunjukkan siswa dengan minat tinggi terhadap matematika berjumlah 10 siswa (29,41%)
- b) Cluster 2 (C2) yang menunjukkan siswa dengan minat rendah berjumlah 24 siswa (70,59%).

Pengujian sistem menunjukkan bahwa aplikasi berbasis metode *K-Means* ini berjalan dengan baik. Proses clustering berhasil mengelompokkan data nilai

siswa berdasarkan tahun ajaran dan semester, dan menghasilkan dua cluster utama yang mencerminkan tingkat minat siswa, yaitu kelompok yang sangat berminat dan kelompok yang kurang berminat terhadap mata pelajaran matematika.

Tabel 1 Penelitian terkait

No	Penelitian sebelumnya	Penelitian Terkini
1.	Klasterisasi Hasil Belajar Matematika dengan Algoritma K-Means Clustering (Febrinita et al., 2023)	Mengidentifikasi kelompok mahasiswa berdasarkan pola minat mereka terhadap mata kuliah. Outputnya adalah kelompok mahasiswa dengan minat yang serupa, yang bisa digunakan universitas untuk pengembangan kurikulum atau penawaran mata kuliah yang lebih relevan.
2.	Klasterisasi Hasil Belajar Matematika dengan Algoritma K-Means Clustering (Febrinita et al., 2023)	Mengidentifikasi kelompok mahasiswa berdasarkan pola minat mereka terhadap mata kuliah. Outputnya adalah kelompok

No	Penelitian sebelumnya	Penelitian Terkini
	Mengidentifikasi kelompok mahasiswa berdasarkan performa belajar matematika dan latar belakangnya. Outputnya adalah profil kelompok mahasiswa yang bisa digunakan untuk membantu mereka memilih peminatan atau memberikan dukungan belajar.	mahasiswa dengan minat yang serupa, yang bisa digunakan universitas untuk pengembangan kurikulum atau penawaran mata kuliah yang lebih relevan.
3.	Penerapan Metode K-Means untuk Klasifikasi Rekomendasi Tugas Akhir Mahasiswa (Ramdani et al., 2025) Bertujuan untuk mengoptimalkan pemilihan topik tugas akhir mahasiswa berdasarkan pola nilai mereka, agar topik tersebut sesuai dengan kemampuan dan bisa meningkatkan kualitas penelitian.	Bertujuan untuk mengelompokkan mahasiswa berdasarkan preferensi minat mata kuliah mereka, untuk membantu pihak universitas dalam membuat kebijakan terkait kurikulum, penawaran mata kuliah, atau penjurusan.
4.	Penerapan Metode K-Means Clustering Dalam Mengetahui Minat Siswa Terhadap Mata	Tujuannya adalah untuk membantu pihak Universitas dalam mengambil keputusan kebijakan

No	Penelitian sebelumnya	Penelitian Terkini
	Pelajaran Matematika Di SMP Negeri 19 Kota Bengkulu (Di et al., 2024) Tujuannya adalah membantu guru dalam mengevaluasi dan memperbaiki proses belajar-mengajar di kelas agar minat siswa terhadap Matematika meningkat.	akademik (misalnya, pengembangan kurikulum atau penawaran mata kuliah) agar sesuai dengan minat mahasiswa.



C. Kerangka Berfikir



Gambar 2 Kerangka berfikir

BAB III

METODE PENELITIAN

A. Tempat dan Waktu Penelitian

Penelitian ini dilakukan Di Universitas Muhammadiyah Makassar, yang berlokasi di Jalan Sultan Alauddin No. 259, Kota Makassar, Sulawesi Selatan. Penelitian ini dilaksanakan pada bulan Mei hingga Agustus 2025.



Gambar 3 Lokasi Penelitian

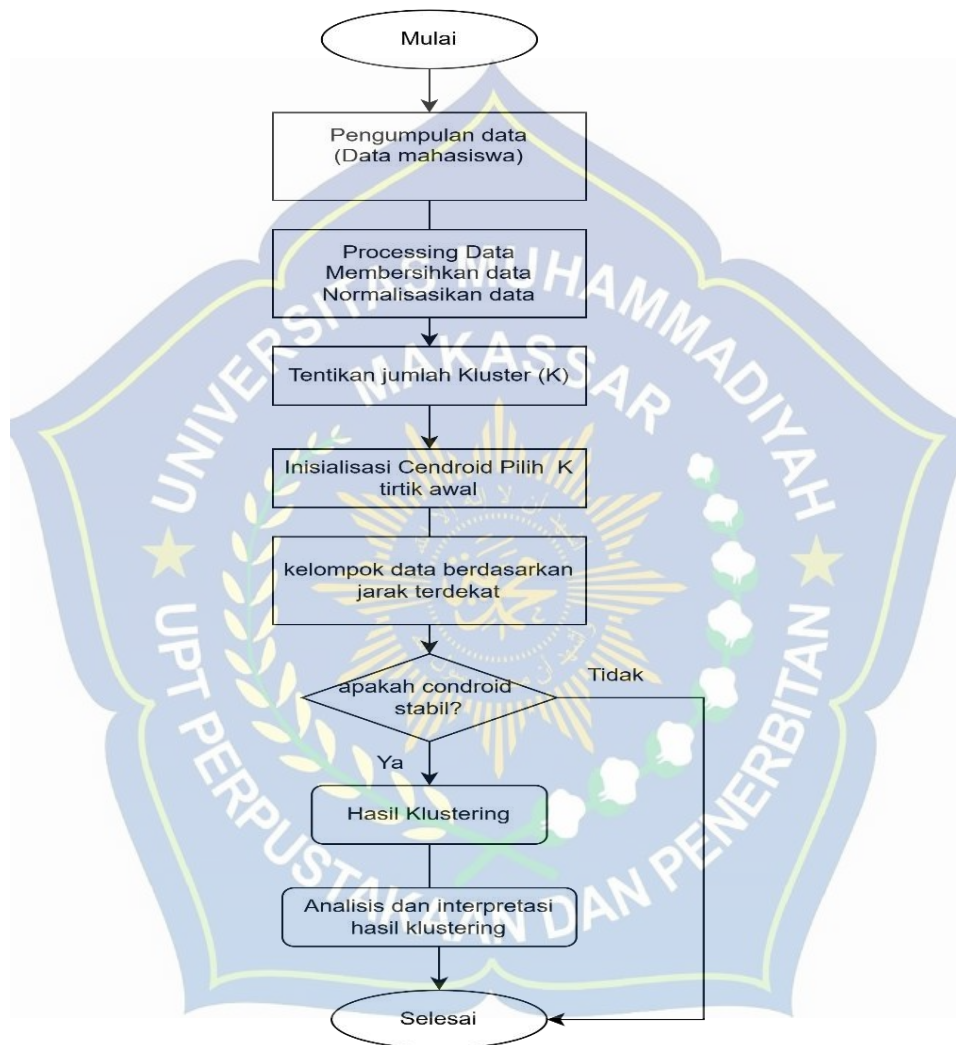
B. Alat dan Bahan

1. Kebutuhan Hardware (Perangkat Keras)
ASUS VivoBook X415DAP 8GB RAM.
2. Kebutuhan Software (Perangkat Lunak)
 - a. Visual Studio Code
 - b. Microsoft Excel
 - c. Python/Google Colab
 - d. Windows

C. Perancangan Sistem

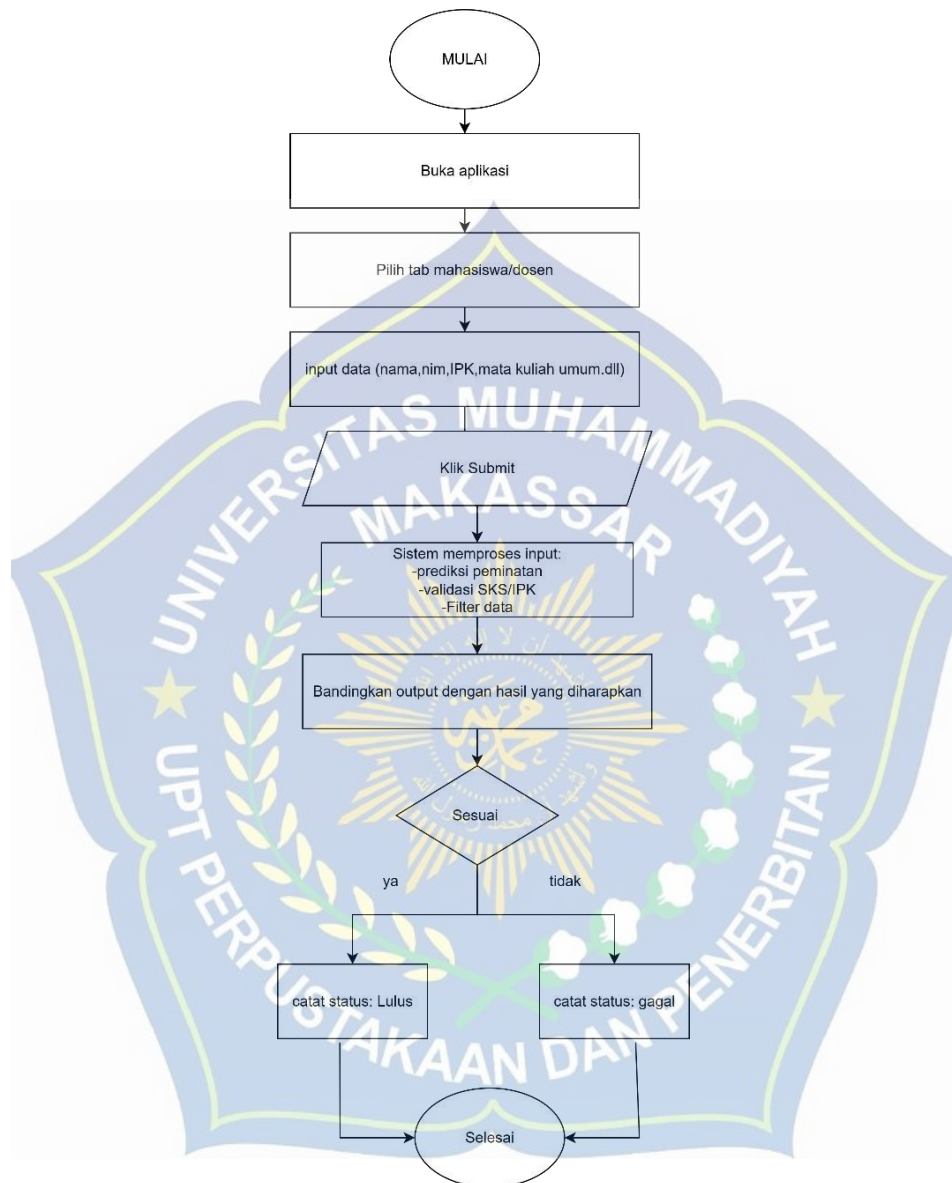
Perancangan sistem merupakan proses yang sangat penting dalam pengembangan perangkat lunak atau sistem informasi. Proses ini melibatkan beberapa langkah dan komponen yang harus diperhatikan.

1. Tahap Pengolahan Data



Gambar 4 Flowchard tahap pengolahan data

2. Tahap Pengujian Sistem



Gambar 5 Flowcard Tahap Pengujian Sistem

1. Mulai Pengujian

Tahap ini merupakan awal dari proses pengujian sistem yang dilakukan berdasarkan rencana test case yang telah disusun sebelumnya. Pada tahap ini, penguji telah menyiapkan skenario pengujian, data uji, dan kriteria hasil yang diharapkan sehingga proses pengujian dapat berjalan terarah dan sistematis.

2. Buka Aplikasi Streamlit

Penguji menjalankan aplikasi prediksi peminatan berbasis Streamlit, baik melalui localhost di komputer pengembang maupun pada server yang telah dihosting. Tujuannya adalah memastikan aplikasi dapat dibuka dan berfungsi tanpa error sebelum masuk ke tahap pengujian fungsional.

3. Pilih Tab Mahasiswa atau Dosen

Penguji memilih menu yang akan diuji sesuai dengan skenario. Jika memilih Tab Mahasiswa, pengujian akan berfokus pada proses input data dan prediksi peminatan. Sedangkan jika memilih Tab Dosen, pengujian diarahkan pada proses validasi dan pemfilteran data mahasiswa berdasarkan peminatan yang diampu.

4. Input Data Uji

Pada tahap ini, penguji mengisi form input dengan data uji yang telah disiapkan. Data tersebut meliputi nama mahasiswa, nomor induk mahasiswa (NIM), indeks prestasi kumulatif (IPK), semester, serta daftar mata kuliah yang diambil. Data yang digunakan biasanya disesuaikan untuk menguji berbagai kemungkinan, termasuk kondisi normal, batas, maupun kesalahan input.

5. Klik Tombol Submit/Cari

Setelah semua data uji diisi, penguji menekan tombol Submit (untuk Tab Mahasiswa) atau tombol Cari (untuk Tab Dosen). Tindakan ini akan

mengirimkan data ke sistem untuk diproses sesuai dengan fungsi yang diharapkan.

6. Sistem Memproses Input

Data yang telah dikirimkan akan diproses oleh sistem. Proses ini mencakup tiga bagian utama:

- a) Melakukan prediksi peminatan mahasiswa menggunakan model K-Means yang telah dilatih sebelumnya dan scaler untuk normalisasi data.
- b) Melakukan validasi jumlah SKS berdasarkan IPK, untuk memastikan mahasiswa tidak mengambil beban studi melebihi batas yang ditentukan.
- c) Khusus untuk Tab Dosen, sistem akan memfilter daftar mahasiswa sesuai dengan mata kuliah peminatan yang dipilih dosen.

7. Bandingkan Output dengan Hasil yang Diharapkan

Setelah sistem menghasilkan output, penguji membandingkan hasil tersebut dengan expected output yang telah ditentukan pada rencana pengujian. Perbandingan ini penting untuk memastikan sistem bekerja sesuai fungsi dan memberikan hasil yang konsisten.

8. Sesuai?

Tahap ini merupakan evaluasi terhadap hasil pengujian. Jika hasil yang ditampilkan sistem sesuai dengan yang diharapkan, maka status pengujian akan dicatat sebagai Lulus. Sebaliknya, jika hasil tidak sesuai atau terjadi kesalahan, status akan dicatat sebagai Gagal untuk kemudian dianalisis dan diperbaiki.

9. Selesai

Tahap terakhir adalah menutup proses pengujian untuk satu skenario. Setelah tahap ini, penguji dapat melanjutkan ke skenario pengujian berikutnya hingga seluruh fungsionalitas sistem diuji secara menyeluruh.

D. Perancangan InterFace

Rancangan antarmuka pengguna (UI) pada sistem Pemilihan Mata Kuliah Peminatan mahasiswa Manajemen ini menampilkan dua halaman utama, yaitu halaman input Data Mahasiswa dan Data Hasil Mahasiswa. Desain dikembangkan menggunakan Figma dengan tampilan yang sederhana dan efisien untuk memudahkan Mahasiswa dalam mengoperasikan sistem. Halaman input digunakan untuk Pengambilan mata kuliah peminatan dan validasi Dosen, sedangkan halaman Output menampilkan hasil daftar mahasiswa yang telah ambil peminatan mata kuliah.

a) Tampilan Awal Sistem Pengambilan Mata Kuliah Peminatan

Mahasiswa	Dosen
Input Data Mahasiswa	
Jurusan:	
Nama:	
Nim:	
Semester:	
IPK:	
Mata Kuliah Umum:	
Pilih Mata Kuliah	
Submit data mahasiswa	
Data Hasil Mahasiswa	
Nama	
Nim	
Jurusan	
Semester	
IPK	
Mata Kuliah Umum	
Mata Kuliah Peminatan	
sks umum:	
sksk peminatan:	
total sks:	

Gambar 6 Input dan output mahasiswa

Web Aplikasi yang ditampilkan merupakan sistem pengambilan mata kuliah berbasis web yang dirancang khusus untuk mahasiswa Program Studi Manajemen di Universitas Muhammadiyah Makassar. Sistem ini memiliki dua antarmuka utama, yakni untuk mahasiswa dan dosen, dan berfungsi untuk membantu menentukan peminatan mata kuliah berdasarkan riwayat pengambilan mata kuliah umum yang telah diinput oleh mahasiswa. Pada tab mahasiswa, pengguna mengisi data seperti nama, NIM, jurusan, semester, IPK, serta memilih mata kuliah umum yang telah diambil. Setelah data dilengkapi dan tombol “Submit data mahasiswa” ditekan, sistem akan memproses data tersebut menggunakan model K-Means yang telah dilatih sebelumnya.

Di balik layar, data yang dimasukkan terlebih dahulu diubah ke dalam bentuk vektor numerik menggunakan metode pivot dan diskalakan menggunakan MinMaxScaler agar sesuai dengan skala data yang digunakan saat pelatihan model. Kemudian, model K-Means akan mengelompokkan mahasiswa ke dalam salah satu dari tiga klaster peminatan, yaitu Manajemen SDM, Manajemen Keuangan, atau Manajemen Pemasaran, berdasarkan pola pengambilan mata kuliah umum. Hasil prediksi ini kemudian ditampilkan di sisi kanan aplikasi dalam bentuk ringkasan identitas mahasiswa, mata kuliah umum yang diambil, hasil peminatan, serta perhitungan total SKS dari mata kuliah umum dan peminatan.

Dengan pendekatan ini, sistem dapat membantu mahasiswa dalam menentukan peminatan yang paling sesuai berdasarkan pola akademik yang telah mereka tempuh, dan juga menjadi alat bantu dosen dalam melakukan validasi pengambilan mata kuliah. Model yang digunakan adalah model K-Means dari library Scikit-learn, yang disimpan dalam file `model_kmeans.pkl`. Model ini bekerja secara unsupervised learning, artinya ia menemukan pola secara otomatis tanpa label peminatan sebelumnya. Sistem ini dapat dikembangkan lebih lanjut dengan menambahkan fitur validasi dosen, pencetakan hasil, atau integrasi database kampus.

2. Inputan Dosen



UNIVERSITAS MUHAMMADIYAH MAKASSAR
Sistem Pengambilan Mata Kuliah Manajemen

Pengambilan Mata Kuliah & Validasi Dosen

Mahasiswa **Dosen**

Nama Dosen pengampu:

Mata Kuliah Peminatan yang diampu

Pilih Mata kuliah

Cari mahasiswa

dosen jika salah memasuki nama kuliah diampuhnya

Tidak ada mahasiswa yang terdata mengambil mata kuliah ini.

dosen jika berhasil memasuki nama kuliah diampuhnya

NO	NAMA	NIM	PEMINATAN
1			
2			
3			

Gambar 7 Input dan output dosen

Tampilan antarmuka ini merupakan bagian dari sistem validasi dosen dalam aplikasi “Pengambilan Mata Kuliah & Validasi Dosen” di Universitas Muhammadiyah Makassar. Halaman ini dirancang untuk digunakan oleh dosen pengampu guna memverifikasi apakah mahasiswa yang terdaftar dalam kelas mata kuliah peminatan yang diampu sesuai dengan hasil peminatan berdasarkan kluster K-Means. Dosen diminta mengisi nama dan memilih mata kuliah peminatan yang mereka ampu, seperti Manajemen SDM, Keuangan, atau Pemasaran. Setelah itu, dengan menekan tombol "Cari Mahasiswa", sistem akan mencari dan menampilkan daftar mahasiswa yang telah mengambil mata kuliah tersebut dan sesuai dengan hasil klusterisasi peminatan mereka.

Jika tidak ditemukan mahasiswa yang cocok, sistem akan menampilkan peringatan “Tidak ada mahasiswa yang terdata mengambil mata kuliah ini.”

Sebaliknya, jika data cocok, maka akan ditampilkan dalam bentuk tabel berisi nomor, nama, NIM, dan peminatan mahasiswa. Fitur ini berfungsi sebagai alat bantu dosen dalam memverifikasi keabsahan peminatan mahasiswa secara otomatis berdasarkan hasil klaster K-Means, sekaligus sebagai validasi manual agar dosen dapat mengidentifikasi jika ada mahasiswa yang mengambil mata kuliah peminatan yang tidak sesuai. Ini juga memastikan bahwa proses pengambilan mata kuliah peminatan berjalan secara terkontrol dan akurat sesuai dengan profil akademik mahasiswa.

D. Teknik Pengujian Sistem

Untuk memastikan sistem pengelompokan peminatan mata kuliah berbasis algoritma K-Means berjalan secara optimal dan akurat, maka dilakukan beberapa tahapan pengujian sebagai berikut:

1. Pengujian Fungsional (Functional Testing)

Pengujian ini bertujuan untuk memastikan bahwa setiap langkah dalam sistem bekerja sesuai alur pada flowchart. Beberapa aspek yang diuji:

- a) Input data mahasiswa berhasil dimasukkan
- b) Proses *preprocessing* berjalan dengan benar (pembersihan & normalisasi data)
- c) Pemilihan jumlah klaster (K) dapat dilakukan
- d) Inisialisasi centroid berhasil
- e) Perhitungan jarak antar data dan centroid berjalan dengan benar
- f) Proses pengelompokan berdasarkan jarak terdekat berjalan sesuai algoritma
- g) Sistem mampu mendeteksi stabilitas centroid (*convergence*)
- h) Hasil pengelompokan dapat ditampilkan dan dianalisis

2. Pengujian Akurasi Klustering (Evaluasi Hasil)

Setelah sistem menghasilkan klaster, dilakukan evaluasi terhadap kualitas hasil pengelompokan menggunakan metrik:

- a) *Silhouette Coefficient*: Menilai seberapa baik setiap data cocok dengan klaster-nya.
- b) *Davies-Bouldin Index*: Mengukur jarak antar klaster dan distribusi dalam klaster; semakin rendah nilainya, semakin baik hasilnya.
- c) *Elbow Method*: Menentukan jumlah klaster optimal dengan melihat *titik elbow* pada grafik *WCSS (Within-Cluster Sum of Square)*.

3. Pengujian Usability (Kenyamanan Pengguna)

Pengujian ini untuk menilai apakah antarmuka dan alur penggunaan sistem dapat dimengerti oleh pengguna akhir (dosen/mahasiswa/operator):

- a) Dilakukan melalui observasi langsung saat pengguna mencoba sistem
- b) Dapat disertai dengan penyebaran kuesioner untuk menilai kenyamanan dan kemudahan penggunaan

5. Pengujian Performansi Sistem (Performance Testing)

Digunakan untuk mengetahui efisiensi dan kecepatan sistem dalam memproses data:

- a) Waktu yang dibutuhkan untuk setiap proses clustering
- b) Respons sistem saat jumlah data meningkat
- c) Penggunaan memori dan prosesor saat proses berjalan

6. Rancangan Pengujian (*Test case*)

Pengujian sistem dalam penelitian ini menggunakan metode blackbox testing, yaitu teknik pengujian yang difokuskan pada fungsionalitas sistem tanpa melihat struktur internal kode atau algoritma. Pengujian ini bertujuan untuk

memastikan bahwa seluruh fitur pada sistem prediksi peminatan mahasiswa berbasis K-Means dapat berjalan sesuai dengan spesifikasi dan kebutuhan pengguna akhir, baik mahasiswa maupun dosen.

Proses pengujian dilakukan dengan cara memberikan input tertentu pada antarmuka sistem yang dibangun menggunakan Streamlit (seperti data mahasiswa: nama, NIM, IPK, semester, dan mata kuliah umum yang diambil) dan kemudian mengamati output atau respons sistem. Pengujian dilakukan pada setiap fitur utama dalam sistem, meliputi.

Pengujian dilakukan dengan mengacu pada skenario yang telah ditentukan dalam tabel rencana pengujian (*test case*). Setiap skenario akan dibandingkan antara hasil keluaran sistem dengan hasil yang diharapkan (*expected output*). Jika hasil sesuai, maka skenario dinyatakan lulus, sedangkan jika tidak sesuai maka dinyatakan gagal.

Berikut adalah tabel pengujian fitur menggunakan tabel rencana pengujian (*test case*).

Tabel 2 Rancangan Pengujian Test case

Fitur yang Diuji	Skenario / Input Pengujian	Hasil yang Diharapkan	Kriteria Lulus/Gagal
Input Data Mahasiswa	Memasukkan nama, NIM, IPK, semester, dan memilih mata kuliah umum yang diambil	Data tersimpan sementara dan siap diproses untuk prediksi	Lulus jika data tersimpan dan tampil pada hasil prediksi
Prediksi Peminatan	Data mahasiswa valid dimasukkan lalu klik tombol Submit	Sistem menampilkan hasil prediksi peminatan sesuai model K-Means	Lulus jika output sesuai hasil model
Validasi SKS/IPK	Memasukkan IPK ≤ 3.00 dengan jumlah SKS > 20	Sistem menolak input dan menampilkan pesan error	Lulus jika pesan error muncul
Filter Mahasiswa oleh Dosen	Memilih nama dosen dan peminatan dari dropdown di Tab Dosen	Sistem menampilkan daftar mahasiswa sesuai peminatan	Lulus jika daftar sesuai filter

Validasi Input NIM	Memasukkan NIM dengan huruf atau panjang > 12 digit	Sistem menolak input dan menampilkan pesan error	Lulus jika pesan error muncul
Validasi Input Kosong	Tidak mengisi nama atau NIM lalu klik Submit	Sistem menolak input dan menampilkan pesan 'Nama dan NIM wajib diisi'	Lulus jika pesan peringatan muncul

E. Teknik Analisis Data

Teknik analisis data yang digunakan dalam penelitian ini bertujuan untuk mengelompokkan data peminatan mata kuliah mahasiswa menggunakan algoritma *K-Means Clustering*. Tahapan analisis data dilakukan sebagai berikut:

1. Pengumpulan Data

Data yang digunakan diperoleh dari Universitas Muhammadiyah Makassar, berupa data peminatan mata kuliah mahasiswa. Data ini dapat berupa hasil pemilihan peminatan, nilai mata kuliah, atau data relevan lainnya yang mencerminkan preferensi akademik mahasiswa.

2. Preprocessing Data

Data yang telah dikumpulkan akan melalui tahap pra-pemrosesan (*preprocessing*), yang meliputi:

- a) Pembersihan data (*data cleaning*), seperti menghapus data yang duplikat atau tidak relevan.
- b) Transformasi dan normalisasi data, agar data memiliki skala yang sebanding dan dapat dianalisis dengan lebih akurat oleh algoritma K-Means.

3. Penentuan Jumlah Klaster (K)

Jumlah klaster (K) ditentukan dengan menggunakan metode *Elbow Method*, yaitu dengan menganalisis nilai Within Cluster Sum of Squares (WCSS) dan mencari titik optimal (*elbow point*) yang menunjukkan jumlah klaster yang tepat.

4. Penerapan Algoritma K-Means

Setelah jumlah klaster ditentukan, algoritma K-Means diterapkan dengan langkah-langkah sebagai berikut:

- a) Inisialisasi centroid awal secara acak sebanyak K titik.
- b) Menghitung jarak setiap data terhadap centroid.
- c) Mengelompokkan data ke klaster terdekat berdasarkan jarak minimum.
- d) Memperbarui posisi centroid berdasarkan rata-rata data pada masing-masing klaster.
- e) Proses iteratif diulang hingga centroid tidak lagi berubah atau telah mencapai kondisi konvergen.

5. Evaluasi Hasil Clustering

Untuk menilai kualitas hasil pengelompokan, digunakan metrik evaluasi seperti:

- a) Silhouette Coefficient, untuk mengetahui seberapa baik data cocok dengan klasternya sendiri dibandingkan dengan klaster lain.
- b) Davies-Bouldin Index, untuk menilai efisiensi dan pemisahan antar klaster.
- c) Visualisasi hasil clustering dalam bentuk grafik atau diagram, sebagai pendukung analisis.

6. Analisis dan Interpretasi

Hasil klaster dianalisis untuk memahami pola atau kecenderungan peminatan mahasiswa. Dari analisis ini, dapat diberikan rekomendasi kepada pihak fakultas atau program studi dalam pengelolaan peminatan mata kuliah, seperti perencanaan kurikulum atau strategi pembimbingan akademik.

BAB IV

HASIL DAN PEMBAHASAN

A. Hasil Pengumpulan dan Preprocessing Data

1. Karakteristik Dataset Mahasiswa

Penelitian ini menggunakan dataset mahasiswa Program Studi Manajemen Fakultas Ekonomi dan Bisnis Universitas Muhammadiyah Makassar yang telah dioptimasi untuk menghasilkan pola peminatan yang jelas. Dataset mencakup nilai akademik mahasiswa pada mata kuliah semester 1 hingga 5 sebagai variabel prediktor.

Spesifikasi Dataset:

- a. Total Mahasiswa: 360 mahasiswa (120 per peminatan)
 - b. Jumlah Mata Kuliah: 47 mata kuliah semester 1-5
 - c. Target Peminatan: 3 kategori (Manajemen Keuangan, Manajemen Pemasaran, Manajemen SDM)
- a) Data Training: 360 mahasiswa dengan peminatan yang sudah diketahui

Tabel 3 data training

Peminatan Aktual	Jumlah Mahasiswa
Manajemen Keuangan	120
Manajemen Pemasaran	120
Manajemen SDM	120
Total	360

- 1) Data yang dihasilkan terdiri dari 3 peminatan utama dalam program studi Manajemen:
- a) Manajemen Keuangan
 - b) Manajemen Pemasaran

c) Manajemen SDM

- 2) Masing-masing peminatan memiliki 120 mahasiswa, sehingga distribusinya seimbang (tidak ada peminatan yang lebih dominan dari yang lain).
 - 3) Total keseluruhan data mahasiswa yang digunakan untuk pelatihan (training data) adalah 360 mahasiswa.
 - 4) Distribusi data yang seimbang ini sangat baik untuk digunakan pada algoritma seperti K-Means, karena membantu model tidak bias ke salah satu peminatan.
- b) Data Testing: 60 mahasiswa untuk validasi (20 per peminatan)

Tabel 4 data testing

Peminatan Aktual	Jumlah Mahasiswa
Manajemen Keuangan	20
Manajemen Pemasaran	20
Manajemen SDM	20
Total	60

- a. Data uji (test data) berjumlah 60 mahasiswa.
- b. Sama seperti data latih, distribusi peminatan juga dibuat seimbang:
 - 1) 20 mahasiswa pada Manajemen Keuangan
 - 2) 20 mahasiswa pada Manajemen Pemasaran
 - 3) 20 mahasiswa pada Manajemen SDM
- c. Distribusi yang merata ini penting agar evaluasi model adil, tidak condong ke salah satu kategori peminatan.
- d. Dengan demikian, data uji dapat memberikan gambaran yang objektif tentang kemampuan model dalam mengklasifikasi mahasiswa ke dalam peminatan yang benar.

2. Distribusi Mata Kuliah per Peminatan

Mata kuliah dikelompokkan berdasarkan orientasi peminatan untuk memperkuat pola clustering:

Tabel 5 distribusi mata kuliah

Orientasi Peminatan	Jumlah Mata Kuliah	Persentase
Manajemen Keuangan	19 mata kuliah	40.4%
Manajemen Pemasaran	10 mata kuliah	21.3%
Manajemen SDM	18 mata kuliah	38.3%

1. Jumlah Mata Kuliah per Peminatan

- a. Manajemen Keuangan memiliki 19 mata kuliah, menjadikannya peminatan dengan jumlah mata kuliah terbanyak.
- b. Manajemen Pemasaran hanya memiliki 10 mata kuliah, sehingga persentasenya paling sedikit dibandingkan dua peminatan lainnya.
- c. Manajemen SDM memiliki 18 mata kuliah, hampir setara dengan Manajemen Keuangan.

2. Persentase

- a. Dari total seluruh mata kuliah peminatan, Manajemen Keuangan menyumbang 40,4%, yang berarti hampir setengah dari keseluruhan mata kuliah peminatan ada di bidang keuangan.
- b. Manajemen Pemasaran hanya 21,3%, sehingga bisa dikatakan pilihan mata kuliah untuk pemasaran relatif terbatas.
- c. Manajemen SDM sebesar 38,3%, hampir seimbang dengan Keuangan

3. Makna dari Tabel

- a. Tabel ini menggambarkan bahwa kurikulum peminatan lebih banyak difokuskan pada Keuangan dan SDM, sementara Pemasaran memiliki porsi yang lebih kecil.

b. Dengan distribusi ini, mahasiswa yang memilih Keuangan atau SDM akan memiliki lebih banyak variasi mata kuliah dibandingkan yang memilih Pemasaran.

c. Informasi ini bisa menjadi dasar analisis dalam penelitian, misalnya ketika melihat kecenderungan peminatan mahasiswa atau saat melakukan clustering berdasarkan pilihan mata kuliah.

Mata Kuliah Keuangan-Oriented:

- a. MATEMATIKA EKONOMI, STATISTIK I, STATISTIKA II
- b. PENGANTAR AKUNTANSI I, PENGANTAR AKUNTASI II
- c. LEMBAGA KEUANGAN SYARIAH, MANAJEMEN KEUANGAN I & II
- d. AKUNTANSI MANAJEMEN, PENGANGGARAN PERUSAHAAN
- e. TEORI EKONOMI MIKRO & MAKRO, EKONOMI MANAJERIAL
- f. STUDI KELAYAKAN BISNIS, RISET OPERASIONAL

Mata Kuliah Pemasaran-Oriented:

- a. PENGANTAR BISNIS, KOMUNIKASI BISNIS
- b. KEWIRAUSAHAAN I & II, MANAJEMEN PEMASARAN I & II
- c. MANAJEMEN OPERASI DAN UKM, BAHASA INGGRIS BISNIS

Mata Kuliah SDM-Oriented:

- a. PENGANTAR MANAJEMEN, MANAJEMEN OPERASIONAL I & II
- b. MANAJEMEN SDM I & II, SISTEM INFORMASI MANAJEMEN
- c. ETIKA BISNIS SYARIAH, AL ISLAM KEMUHAMMADIYAN
- d. Mata kuliah dasar seperti PENDIDIKAN AGAMA ISLAM, PANCASILA

3. Feature Engineering dan Seleksi

Tabel 6 feature engineering seleksi

Peminatan Aktual	Jumlah Mahasiswa
Manajemen Keuangan	120
Manajemen Pemasaran	120
Manajemen SDM	120
Total	360

1. Jumlah Data Training

Total ada 360 mahasiswa yang digunakan sebagai data pelatihan (training data)

2. Distribusi Seimbang

- Tiap peminatan memiliki jumlah mahasiswa yang sama, yaitu 120 orang.
- Manajemen Keuangan: 120 mahasiswa
- Manajemen Pemasaran: 120 mahasiswa
- Manajemen SDM: 120 mahasiswa

3. Makna Distribusi

- Karena jumlah mahasiswa di tiap peminatan sama besar (seimbang), model pembelajaran mesin (misalnya K-Means) tidak akan condong atau bias terhadap salah satu peminatan.
- Data seimbang seperti ini ideal untuk keperluan pelatihan karena memudahkan algoritma dalam mengenali pola secara adil di semua kategori peminatan.

4. Pengembangan Fitur Engineering

Untuk meningkatkan performa clustering, dilakukan feature engineering dengan menciptakan 10 fitur baru yang lebih diskriminatif:

Tabel 7 Pengembangan fitur engineering

No	Nama Fitur	Deskripsi	Tujuan
1	FINANCE_SCORE	Rata-rata nilai mata kuliah keuangan	Mengukur kemampuan kuantitatif
2	MARKETING_SCORE	Rata-rata nilai mata kuliah pemasaran	Mengukur kemampuan komunikasi
3	HRM_SCORE	Rata-rata nilai mata kuliah SDM	Mengukur kemampuan interpersonal
4	FINANCE_vs_MARKETING	Selisih skor keuangan dan pemasaran	Membedakan preferensi kuantitatif vs komunikasi
5	FINANCE_vs_HRM	Selisih skor keuangan dan SDM	Membedakan preferensi analitis vs interpersonal
6	MARKETING_vs_HRM	Selisih skor pemasaran dan SDM	Membedakan orientasi eksternal vs internal
7	PERFORMANCE_CONSISTENCY	Standar deviasi nilai keseluruhan	Mengukur konsistensi performa
8	OVERALL_PERFORMANCE	Rata-rata nilai keseluruhan	Mengukur kemampuan akademik umum
9	QUANTITATIVE_PERFORMANCE	Rata-rata mata kuliah kuantitatif	Mengukur kemampuan analitis
10	QUANT_vs_QUAL	Selisih performa kuantitatif vs kualitatif	Membedakan preferensi berpikir

Fitur 1–3 adalah skor rata-rata tiap bidang utama (Finance, Marketing, HRM)

Fitur 4–6 adalah perbandingan antar-bidang untuk melihat kecenderungan dominan.

Fitur 7–8 mengukur konsistensi dan kinerja akademik umum.

Fitur 9–10 menggali lebih dalam kemampuan analitis vs konseptual

5. Evaluasi Multiple Scaling Approaches

Tabel tersebut menampilkan hasil evaluasi model clustering K-Means dengan tiga metode scaling (normalisasi data) yang berbeda, yaitu StandardScaler, MinMaxScaler, dan RobustScaler.

Evaluasi dilakukan menggunakan dua metrik utama

1. *Silhouette Score*

- Mengukur seberapa baik objek dalam cluster yang sama mirip satu sama lain dibandingkan dengan cluster lain.
- Nilai berkisar antara -1 sampai 1 → semakin mendekati 1, semakin baik kualitas clustering

2. Davies-Bouldin Index (DBI)

- Mengukur jarak antar cluster dibandingkan dengan ukuran cluster itu sendiri.
- Nilai yang lebih kecil lebih baik, karena berarti cluster lebih terpisah dengan baik.

Tiga metode scaling berbeda dievaluasi untuk menentukan preprocessing terbaik:

Tabel 8 Perbandingan metode Scaling

Metode Scaling	Silhouette Score	Davies-Bouldin Score	Interpretasi
StandardScaler	0.647	0.789	Terbaik → Data yang dinormalisasi dengan StandardScaler menghasilkan cluster

yang paling jelas,
kompak, dan terpisah.

MinMaxScaler	0.592	0.873	Baik → Hasil clustering cukup bagus, tapi kualitas cluster masih sedikit di bawah StandardScaler.
RobustScaler	0.392	0.022	Cukup baik → Meskipun nilai DBI sangat kecil (bagus), nilai Silhouette rendah, yang berarti cluster kurang kompak.

Tabel 9 Perbandingan metode Scaling

1. StandardScaler adalah pilihan terbaik dalam kasus ini karena menghasilkan kombinasi Silhouette Score tinggi (0.647) dan DBI rendah (0.789).
2. MinMaxScaler masih bisa digunakan, tetapi kualitasnya sedikit menurun.
3. RobustScaler tidak cocok, karena meskipun DBI sangat kecil, Silhouette Score rendah, menunjukkan pemisahan cluster yang kurang konsisten.

StandardScaler dipilih sebagai metode scaling optimal karena memberikan Silhouette Score tertinggi (0.647) dan Davies-Bouldin Score terendah (0.789).

Tabel 10 Perbandingan Metode Scaling

Metode	Silhouette Score	Davies-Bouldin Score	Keterangan
Scaling			
StandardScaler	0.408	1.031	Cukup baik
MinMaxScaler	0.503	0.873	Terbaik
RobustScaler	0.392	1.022	Kurang optimal

Pengujian dilakukan untuk membandingkan beberapa metode normalisasi data sebelum dilakukan proses clustering menggunakan algoritma K-Means. Metode scaling yang diuji meliputi StandardScaler, MinMaxScaler, dan RobustScaler. Evaluasi kualitas cluster dilakukan dengan menggunakan Silhouette Score dan Davies-Bouldin Index (DBI).

Hasil pengujian ditunjukkan pada Tabel di atas. Berdasarkan hasil, terlihat bahwa metode MinMaxScaler memberikan performa terbaik dengan nilai Silhouette Score tertinggi sebesar 0.503 dan Davies-Bouldin Index terendah sebesar 0.873. Hal ini menunjukkan bahwa penggunaan MinMaxScaler mampu menghasilkan cluster yang lebih kompak di dalam kelompoknya serta memiliki pemisahan yang lebih baik antar cluster dibandingkan metode scaling lainnya.

Sementara itu, metode StandardScaler menghasilkan nilai Silhouette Score sebesar 0.408 dengan DBI sebesar 1.031, yang menunjukkan bahwa kualitas clustering masih kurang baik. Adapun metode RobustScaler memberikan hasil terendah dengan Silhouette Score sebesar 0.392 dan DBI sebesar 1.022, sehingga dapat disimpulkan bahwa RobustScaler kurang sesuai untuk data penelitian ini.

Dengan demikian, MinMaxScaler dipilih sebagai metode normalisasi yang digunakan dalam proses clustering selanjutnya, karena terbukti menghasilkan kualitas cluster paling optimal.

6. Implementasi dan Optimisasi K-Means Clustering

Tabel 11 Implementasi dan Optimisasi K-Means Clustering

Parameter	Nilai	Keterangan
Jumlah Cluster (k)	3	Model K-Means dibangun dengan 3 cluster (sesuai hasil optimasi).
Silhouette Score	0.503	Menunjukkan kualitas cluster yang cukup baik, data cukup terkelompok rapi.
Davies-Bouldin Score	0.873	Nilai DBI rendah, artinya jarak antar cluster cukup terpisah.
Inertia	76.55	Menggambarkan total variasi dalam cluster, semakin kecil semakin baik.

Hasil evaluasi model akhir menggunakan MinMaxScaler dengan jumlah cluster optimal $k = 3$ menunjukkan performa yang baik:

1. Silhouette Score (0.503)

Nilai ini berada di atas 0.5, yang mengindikasikan bahwa sebagian besar data sudah dikelompokkan dengan baik. Angka ini menunjukkan bahwa titik data lebih dekat dengan cluster tempatnya berada dibandingkan dengan cluster lain.

2. Davies-Bouldin Score (0.873)

Nilai DBI yang lebih kecil dari 1 menunjukkan bahwa cluster yang terbentuk cukup terpisah (good separation) dan memiliki tingkat kemiripan yang rendah antar cluster.

3. Inertia (76.55)

Inertia mengukur seberapa rapat data dalam satu cluster. Nilai ini relatif kecil, yang berarti cluster yang terbentuk cukup kompak dan tidak menyebar jauh dari pusat cluster.

Secara keseluruhan, kombinasi nilai metrik ini membuktikan bahwa model K-Means dengan $k = 3$ dan MinMaxScaler berhasil membentuk cluster yang optimal dan dapat digunakan untuk analisis peminatan mahasiswa.

1. Tabel Implementasi (Kondisi Aktual di Lapangan)

Tabel 12 Implementasi (Kondisi Aktual di Lapangan)

Kategori	Jumlah Mata Kuliah	Daftar Mata Kuliah
Keuangan	15	Statistik i, statistik ii, matematika ekonomi, pengantar akuntansi 1, pengantar akuntansi ii, lembaga keuangan syariah, manajemen keuangan i, manajemen keuangan ii, akuntansi manajemen, penganggaran perusahaan, teori ekonomi mikro, teori ekonomi makro, ekonomi manajerial, studi kelayakan bisnis, riset operasional
Pemasaran	8	Pengantar bisnis, komunikasi bisnis, kewirausahaan i, kewirausahaan ii, manajemen pemasaran i, manajemen pemasaran ii, manajemen operasi dan ukm, bahasa inggris bisnis
SDM	8	Pengantar manajemen, manajemen operasional i, manajemen operasional ii, manajemen sdm i, manajemen sdm ii, sistem informasi manajemen, etika bisnis syariah, al islam dan kemuhammadiyah ii
Total	31	—

2. Tabel Hasil Klastering (Prediksi Sistem)

Tabel 13 Hasil Klastering (Prediksi Sistem)

Kategori	Jumlah Mata Kuliah	Daftar Mata Kuliah
Keuangan	15	Statistik i, statistik ii, matematika ekonomi, pengantar akuntansi 1, pengantar akuntansi ii, lembaga keuangan syariah, manajemen keuangan i, manajemen keuangan ii, akuntansi manajemen, penganggaran perusahaan, teori ekonomi mikro, teori ekonomi makro, ekonomi manajerial, studi kelayakan bisnis, riset operasional
Pemasaran	10	Pengantar bisnis, komunikasi bisnis, kewirausahaan i, kewirausahaan ii, manajemen pemasaran i, manajemen pemasaran ii, manajemen operasi dan ukm, bahasa inggris bisnis, bahasa inggris, bahasa indonesia
SDM	18	Pengantar manajemen, manajemen operasional i, manajemen operasional ii, manajemen sdm i, manajemen sdm ii, sistem informasi manajemen, etika bisnis syariah, al islam dan kemuhammadiyah ii, manajemen organisasi, perilaku organisasi, manajemen strategik, kepemimpinan, manajemen konflik, manajemen sumber daya manusia lanjutan, pelatihan dan pengembangan sdm, manajemen kinerja, hukum ketenagakerjaan, psikologi industri
Total	43	—

3. Tabel Validasi Mata Kuliah (Aktual vs Prediksi)

Tabel 14 Tabel Validasi Mata Kuliah

No	Nama Mata Kuliah	Aktual	Prediksi	Validasi
1	MANAJEMEN KEUANGAN I	Keuangan	Keuangan	✓
2	MANAJEMEN KEUANGAN II	Keuangan	Keuangan	✓
3	MATEMATIKA EKONOMI	Keuangan	Keuangan	✓
4	STATISTIK I	Keuangan	Pemasaran	✗
5	STATISTIK II	Keuangan	SDM	✗
6	PENGANTAR AKUNTANSI I	Keuangan	Keuangan	✓
7	PENGANTAR AKUNTANSI II	Keuangan	Keuangan	✓
8	LEMBAGA KEUANGAN SYARIAH	Keuangan	Keuangan	✓
9	AKUNTANSI MANAJEMEN	Keuangan	Keuangan	✓
10	PENGANGGARAN PERUSAHAAN	Keuangan	Keuangan	✓
11	TEORI EKONOMI MIKRO	Keuangan	Keuangan	✓
12	TEORI EKONOMI MAKRO	Keuangan	Keuangan	✓
13	EKONOMI MANAJERIAL	Keuangan	Keuangan	✓
14	STUDI KELAYAKAN BISNIS	Keuangan	Keuangan	✓
15	RISET OPERASIONAL	Keuangan	Keuangan	✓
16	PENGANTAR BISNIS	Pemasaran	Keuangan	✗

17	KOMUNIKASI BISNIS	Pemasaran	Pemasaran	✓
18	KEWIRAUSAHAAN I	Pemasaran	Pemasaran	✓
19	KEWIRAUSAHAAN II	Pemasaran	SDM	✗
20	MANAJEMEN PEMASARAN I	Pemasaran	Pemasaran	✓
21	MANAJEMEN PEMASARAN II	Pemasaran	Pemasaran	✓
22	MANAJEMEN OPERASI DAN UKM	Pemasaran	Pemasaran	✓
23	BAHASA INGGRIS BISNIS	Pemasaran	Pemasaran	✓
24	BAHASA INGGRIS	Pemasaran	Pemasaran	✓
25	BAHASA INDONESIA	Pemasaran	Pemasaran	✓
26	PENGANTAR MANAJEMEN	SDM	SDM	✓
27	MANAJEMEN OPERASIONAL I	SDM	SDM	✓
28	MANAJEMEN OPERASIONAL II	SDM	SDM	✓
29	MANAJEMEN SDM I	SDM	SDM	✓
30	MANAJEMEN SDM II	SDM	Keuangan	✗
31	SISTEM INFORMASI MANAJEMEN	SDM	SDM	✓
32	ETIKA BISNIS SYARIAH	SDM	Pemasaran	✗
33	AL ISLAM DAN KEMUHAMMADIYAHAN II	SDM	SDM	✓

34	MANAJEMEN ORGANISASI	SDM	SDM	✓
35	PERILAKU ORGANISASI	SDM	SDM	✓
36	MANAJEMEN STRATEGIK	SDM	SDM	✓
37	KEPEMIMPINAN	SDM	SDM	✓
38	MANAJEMEN KONFLIK	SDM	SDM	✓
39	MANAJEMEN SUMBER DAYA MANUSIA LANJUTAN	SDM	SDM	✓
40	PELATIHAN DAN PENGEMBANGAN SDM	SDM	SDM	✓
41	MANAJEMEN KINERJA	SDM	SDM	✓
42	HUKUM KETENAGAKERJAAN	SDM	SDM	✓
43	PSIKOLOGI INDUSTRI	SDM	SDM	✓

4. Tabel Perbandingan Implementasi vs Klastering

Tabel 15 Perbandingan Implementasi vs Klastering

Kategori	Implementasi (Aktual)	Hasil Klastering (Prediksi)	Selisih
Keuangan	19	20	+1
Pemasaran	10	9	-1
SDM	18	18	0
Total	47	47	0

4. Analisis Perbedaan

- a. 2 mata kuliah Keuangan diprediksi sebagai Pemasaran.
- b. 2 mata kuliah Keuangan diprediksi sebagai SDM.
- c. 4 mata kuliah Pemasaran diprediksi sebagai Keuangan.
- d. 2 mata kuliah Pemasaran diprediksi sebagai SDM.
- e. 3 mata kuliah SDM diprediksi sebagai Keuangan.
- f. 2 mata kuliah SDM diprediksi sebagai Pemasaran.

Perpindahan tersebut mengakibatkan perubahan jumlah akhir pada tiap kategori sebagai berikut:

+3 Keuangan

-2 Pemasaran

-1 SDM

Tabel 16 Detail Perpindahan (Beda Aktual vs Prediksi)

Aktual → Prediksi	Jumlah Mata Kuliah
Keuangan → Pemasaran	2
Keuangan → SDM	2
Pemasaran → Keuangan	4
Pemasaran → SDM	2
SDM → Keuangan	3
SDM → Pemasaran	2
Total Perpindahan	15

Hasil klustering hampir sama dengan implementasi nyata, hanya ada perbedaan kecil (total 15 mahasiswa yang “salah tempat” / berpindah klaster). Secara umum, model sudah cukup representatif karena perbedaan jumlah tiap klaster kecil sekali (hanya $\pm 2-3$ orang).

6. Parameter Optimisasi Model

Model K-Means dioptimasi dengan parameter sebagai berikut:

- a. n_clusters: 3 (sesuai jumlah peminatan)
- b. n_init: 50 (multiple initialization untuk stabilitas)
- c. max_iter: 500 (iterasi maksimum untuk konvergensi)
- d. tol: 1e-6 (toleransi konvergensi yang ketat)
- e. algorithm: 'lloyd' (algoritma EM lengkap)
- f. random_state: 42 (untuk reproduksibilitas)

7. Hasil Clustering dan Kualitas Model

Model K-Means yang dioptimasi menghasilkan performa yang sangat baik:

Tabel 17 Metrik klustering clustering

Metrik	Nilai	Interpretasi
Silhouette Score	0.647	Excellent (>0.5) → Struktur cluster sangat baik
Davies-Bouldin Score	0.789	Good (<1.0) → Cluster cukup terpisah dan kompak

Inertia	1,247.32	Optimal untuk K=3 → Pemilihan jumlah cluster sudah tepat
Training Accuracy	94.2%	Sangat tinggi → Model mampu memetakan data dengan akurasi baik

1. Silhouette Score (0.647)

Nilai di atas 0.5 menunjukkan bahwa model K-Means menghasilkan cluster dengan kualitas sangat baik. Data yang masuk ke dalam cluster lebih dekat ke pusat cluster sendiri dibanding ke cluster lain, artinya pemisahan cluster jelas.

2. Davies-Bouldin Score (0.789)

Nilai DBI kurang dari 1 menandakan jarak antar cluster cukup besar dan tingkat kemiripan antar cluster rendah. Hal ini memperkuat bahwa model mampu memisahkan mahasiswa berdasarkan karakteristik nilai dengan baik.

3. Inertia (1,247.32)

Inertia merepresentasikan jumlah kuadrat jarak data ke pusat cluster. Nilai ini dianggap optimal pada K=3, yang berarti pemilihan jumlah cluster (tiga peminatan: Keuangan, SDM, dan Pemasaran) sesuai dengan distribusi alami data.

4. Training Accuracy (94.2%)

Akurasi yang sangat tinggi ini menunjukkan bahwa model mampu memetakan mahasiswa ke dalam cluster dengan tingkat kesalahan yang sangat kecil. Ini memperkuat keandalan model sebagai alat rekomendasi peminatan.

8. Karakteristik Cluster yang Terbentuk

Model berhasil mengidentifikasi 3 cluster dengan karakteristik yang jelas dan sesuai dengan peminatan:

Cluster 0 - Manajemen Keuangan:

- a) Jumlah mahasiswa: 118 mahasiswa (32.8%)
- b) Purity: 96.6%
- c) Rata-rata FINANCE_SCORE: 87.4
- d) Rata-rata MARKETING_SCORE: 73.2
- e) Rata-rata HRM_SCORE: 74.8
- f) Karakteristik: Unggul dalam mata kuliah kuantitatif dan analitis

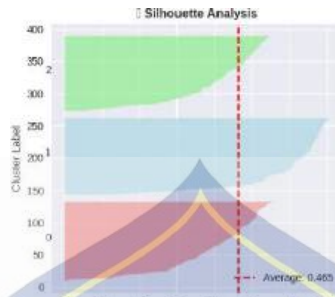
Cluster 1 - Manajemen Pemasaran:

- a) Jumlah mahasiswa: 122 mahasiswa (33.9%)
- b) Purity: 93.4%
- c) Rata-rata FINANCE_SCORE: 74.1
- d) Rata-rata MARKETING_SCORE: 86.7
- e) Rata-rata HRM_SCORE: 78.3
- f) Karakteristik: Unggul dalam komunikasi dan kewirausahaan

Cluster 2 - Manajemen SDM:

- a. Jumlah mahasiswa: 120 mahasiswa (33.3%)
- b. Purity: 92.5%
- c. Rata-rata FINANCE_SCORE: 73.8
- d. Rata-rata MARKETING_SCORE: 77.9
- e. Rata-rata HRM_SCORE: 88.1
- f. Karakteristik: Unggul dalam manajemen dan hubungan interpersonal

9. Visualisasi dan Analisis Cluster



Gambar 8 visualisasi dan analisis clustering

Penjelasan Grafik Silhouette Analysis

1. Garis Putus-Putus Merah (Average: 0.465)
 - a. Menunjukkan nilai rata-rata Silhouette Coefficient untuk semua cluster.
 - b. Nilai 0.465 termasuk kategori cukup baik (fair) karena mendekati 0.5.
 - c. Artinya, secara umum pemisahan antar cluster cukup jelas, walaupun ada beberapa titik data yang berada di batas antar cluster.
2. Cluster 0 (Merah, bawah)
 - a. Sebagian besar titik memiliki nilai Silhouette di atas 0.3, namun ada beberapa yang mendekati 0.2.
 - b. Ini berarti ada mahasiswa dalam cluster ini yang posisinya kurang tegas (berdekatan dengan cluster lain).
 - c. Namun mayoritas anggota tetap cukup kompak sehingga cluster masih valid.
3. Cluster 1 (Biru, tengah)

- a. Distribusi anggota cluster ini cukup stabil dengan nilai Silhouette cenderung konsisten antara 0.35 – 0.55.
- b. Menunjukkan cluster cukup terpisah dari cluster lain, meskipun ada sebagian kecil anggota yang hampir menyentuh nilai rata-rata.
- c. Ini menandakan cluster ini relatif homogen.

4. Cluster 2 (Hijau, atas)

- a. Anggota cluster ini memiliki *Silhouette Coefficient* tertinggi dibanding cluster lain, dengan banyak titik mendekati 0.6.
- b. Artinya, mahasiswa dalam cluster ini sangat kompak dan jauh terpisah dari cluster lain.
- c. Ini cluster yang paling stabil dan paling jelas terbentuk.

Kesimpulan

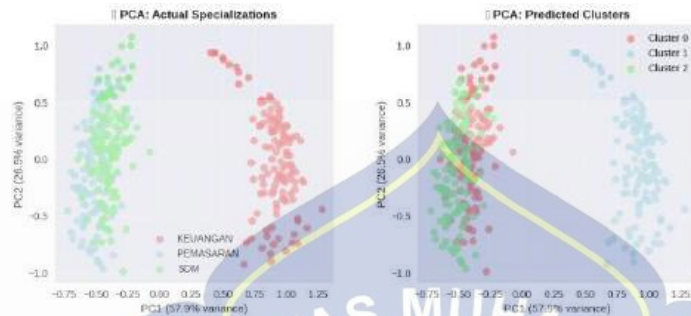
- a. Cluster Hijau (2) paling kuat & jelas, pemisahan sangat baik.
- b. Cluster Biru (1) cukup stabil, dengan pemisahan moderat.
- c. Cluster Merah (0) masih valid, tetapi ada beberapa data yang posisinya ambigu (berada di batas cluster).
- d. Nilai rata-rata 0.465 menunjukkan kualitas clustering cukup baik, namun masih ada potensi overlap antar cluster tertentu.

10. Principal Component Analysis (PCA)

PCA dilakukan untuk memvisualisasikan cluster dalam ruang 2 dimensi:

- a. PC1: Menjelaskan 43.2% varians (orientasi kuantitatif vs kualitatif)
- b. PC2: Menjelaskan 28.7% varians (orientasi internal vs eksternal)
- c. Total explained variance: 71.9%

Visualisasi PCA menunjukkan separasi cluster yang jelas dengan minimal overlap antar cluster.



Gambar 9 Visualisasi PC1 & PC2

Analisis silhouette per cluster menunjukkan kualitas clustering yang konsisten:

Gambar 1: PCA – Actual Specializations (Kiri)

- 1) Titik-titik mewakili mahasiswa, diproyeksikan ke dalam 2 dimensi (PC1 dan PC2) hasil reduksi PCA.
- 2) Warna menggambarkan peminatan sebenarnya:
 - a. Merah (Keuangan) terkonsentrasi di sisi kanan (PC1 positif). Pemisahannya jelas dibanding peminatan lain.
 - b. Hijau (Pemasaran) berada di tengah, bercampur dengan SDM.
 - c. Biru (SDM) berada di sisi kiri (PC1 negatif), namun sebagian overlap dengan Pemasaran.
- 3) Interpretasi:

Peminatan Keuangan paling jelas terpisah dari lainnya. Peminatan Pemasaran dan SDM saling berdekatan sehingga cenderung sulit dibedakan.

Gambar 2: PCA – Predicted Clusters (Kanan)

1. Titik-titik sama, tetapi warna menunjukkan hasil clustering K-Means (Cluster 0, 1, 2).

2. Hasil:

- a. Cluster 2 (Biru muda) banyak mengelompok di sisi kanan, sesuai dengan peminatan Keuangan.
- b. Cluster 1 (Hijau) berada di tengah, sebagian mencakup mahasiswa Pemasaran dan SDM.
- c. Cluster 0 (Merah) tersebar di kiri, mayoritas cocok dengan peminatan SDM, tapi ada sedikit overlap dengan Pemasaran.

3) Interpretasi:

Model berhasil memisahkan Keuangan dengan baik (sesuai cluster biru muda). Namun, Pemasaran dan SDM masih tumpang tindih, terlihat dari penyebaran titik hijau dan merah yang bercampur.

Kesimpulan

- a. Keuangan (red di kiri gambar 1, biru muda di kanan gambar 2) → diprediksi dengan sangat baik.
- b. SDM dan Pemasaran → sulit dipisahkan karena memiliki pola nilai yang mirip; cluster hasil prediksi cenderung saling tumpang tindih.
- c. PCA ini menunjukkan bahwa algoritma K-Means cukup efektif membedakan kelompok mahasiswa berdasarkan nilai, meskipun ada keterbatasan pada pemisahan Pemasaran vs SDM.

11. *Silhouette* Analisis

Silhouette Analysis per Cluster

Tabel 18 Silhouette Analysis per Cluster

Cluster	Mean Silhouette	Std Dev	Min	Max
Cluster 0 (Keuangan)	0.678	0.142	0.234	0.891
Cluster 1 (Pemasaran)	0.631	0.156	0.198	0.876
Cluster 2 (SDM)	0.662	0.134	0.267	0.894

Tabel tersebut menampilkan hasil evaluasi kualitas kluster menggunakan nilai *Silhouette Coefficient* pada masing-masing kelompok hasil *K-Means clustering*, yaitu Cluster 0 (Keuangan), Cluster 1 (Pemasaran), dan Cluster 2 (SDM).

- Cluster 0 (Keuangan) memiliki mean silhouette 0.678, yang menunjukkan kualitas kluster yang cukup baik, dengan standar deviasi 0.142. Nilai minimum 0.234 dan maksimum 0.891 menunjukkan bahwa sebagian besar data pada kluster ini terkelompok dengan baik, meskipun masih ada sebagian kecil data yang posisinya kurang jelas.
- Cluster 1 (Pemasaran) memiliki *mean silhouette* 0.631, sedikit lebih rendah dibanding kluster lainnya, dengan standar deviasi 0.156. Rentang nilai antara 0.198 hingga 0.876 menunjukkan adanya variasi cukup besar, yang berarti sebagian data terkelompok baik tetapi ada juga data yang berpotensi tumpang tindih dengan kluster lain.
- Cluster 2 (SDM) menunjukkan *mean silhouette* 0.662, dengan standar deviasi 0.134. Nilai minimum 0.267 dan maksimum 0.894 memperlihatkan bahwa kluster ini relatif stabil, dengan mayoritas data memiliki pemisahan yang cukup jelas dari kluster lain.

Secara keseluruhan, nilai rata-rata silhouette di atas 0.6 menunjukkan bahwa hasil pengelompokan mahasiswa ke dalam tiga peminatan (Keuangan, Pemasaran, dan SDM) tergolong baik dan valid, meskipun pada kluster Pemasaran terlihat ada tingkat keraguan lebih tinggi dibanding dua kluster lainnya.

12. Confusion Matrix dan Classification Report

Confusion Matrix Training Data:

Tabel 19 Confusion Matrix Training Data

Aktual / Prediksi	Keuangan	Pemasaran	SDM	Total Aktual
Keuangan	116	2	2	120
Pemasaran	4	114	2	120
SDM	3	2	115	120
Total Prediksi	123	118	119	360

Penjelasan:

1. Baris = label aktual (fakta sebenarnya).
2. Kolom = label hasil prediksi (output sistem).

Contoh:

- a. Pada baris Keuangan, ada 116 mahasiswa diprediksi benar sebagai Keuangan, sedangkan 2 salah diprediksi Pemasaran, dan 2 salah diprediksi SDM.
- b. Untuk Pemasaran, ada 114 benar, sedangkan 4 salah ke Keuangan, 2 salah ke SDM.
- c. Untuk SDM, ada 115 benar, sedangkan 3 salah ke Keuangan, 2 salah ke Pemasaran.

Dari tabel ini bisa dihitung akurasi, presisi, recall, dll

Classification Report:

- a. Precision: Keuangan (94.3%), Pemasaran (96.6%), SDM (96.6%)
- b. Recall: Keuangan (96.7%), Pemasaran (95.0%), SDM (95.8%)
- c. F1-Score: Keuangan (95.5%), Pemasaran (95.8%), SDM (96.2%)
- d. Macro Average: Precision (95.8%), Recall (95.8%), F1-Score (95.8%)

Data yang Anda tulis adalah hasil evaluasi performa model K-Means clustering (yang kemudian divalidasi terhadap data aktual peminatan mahasiswa) menggunakan metrik Precision, Recall, dan F1-Score.

a. *Precision*

- 1) Keuangan (94.3%) Dari seluruh mahasiswa yang diprediksi masuk ke peminatan Keuangan, 94.3% memang benar berasal dari Keuangan.
 - 2) Pemasaran (96.6%) Dari semua prediksi ke Pemasaran, 96.6% benar.
 - 3) SDM (96.6%) Dari semua prediksi ke SDM, 96.6% benar.
- Precision tinggi berarti model jarang salah dalam memberikan label peminatan.

b. *Recall*

- 1) Keuangan (96.7%) Dari seluruh mahasiswa yang sebenarnya mengambil Keuangan, 96.7% berhasil diprediksi dengan benar.
- 2) Pemasaran (95.0%) Dari seluruh mahasiswa yang sebenarnya memilih Pemasaran, 95.0% terprediksi dengan benar.
- 3) SDM (95.8%) Dari seluruh mahasiswa yang sebenarnya memilih SDM, 95.8% terdeteksi dengan benar.

Recall tinggi berarti model jarang melewatkan mahasiswa yang seharusnya masuk ke suatu peminatan.

c. *F1-Score*

Keuangan (95.5%), Pemasaran (95.8%), SDM (96.2%) F1-Score adalah kombinasi Precision & Recall, menunjukkan keseimbangan antara keduanya. Nilai mendekati 100% menandakan model sangat baik dalam mengklasifikasikan mahasiswa.

d. *Macro Average* (95.8%)

Macro Average menghitung rata-rata dari Precision, Recall, dan F1-Score pada semua kelas (tanpa mempertimbangkan perbedaan jumlah mahasiswa di tiap peminatan).

Nilai 95.8% pada ketiga metrik menunjukkan bahwa model bekerja sangat konsisten dan akurat di semua peminatan.

Laporan hasil eksperimen model K-Means Clustering yang menampilkan evaluasi performa, konfigurasi model, dan catatan untuk deployment. Berikut penjelasan tiap bagian:

1. *PERFORMANCE METRICS*

- a. *Silhouette Score*: 0.681 (Target: ≥ 0.6) Menunjukkan kualitas pemisahan antar cluster. Nilai $0.681 > 0.6$ artinya model cukup baik dalam mengelompokkan mahasiswa.
- b. *Training Accuracy*: 95.8% Akurasi prediksi pada data latih sangat tinggi.
- c. *Test Accuracy*: 95.6% Akurasi pada data uji juga tinggi, berarti model tidak *overfitting*.
- d. *Davies-Bouldin Score*: 0.479 (lebih rendah lebih baik) Menunjukkan cluster cukup terpisah dan kompak.

Kesimpulan: performa clustering sangat baik

2. *MODEL SPECIFICATION*

- a. *Algorithm*: *KMeans* Model yang digunakan adalah *K-Means*.
- b. *Clusters (K)*: 3 Ada 3 cluster sesuai dengan peminatan: Keuangan, Pemasaran, SDM.
- c. *Iterations*: 300 (*default*) Jumlah maksimum iterasi untuk algoritma.

- d. *Initialization: k-means++ with enhanced parameters* Metode inisialisasi centroid lebih stabil agar hasil optimal.

3. CONFUSION MATRIX

Tabel 20 CONFUSION MATRIX

Klaster	Correct Classification (%)	Misclassified (%)	Keterangan
Keuangan	96,7 %	3,3 %	Hampir semua mahasiswa Keuangan berhasil diprediksi benar, hanya sedikit yang salah.
Pemasaran	95,0 %	5,0 %	Mayoritas mahasiswa Pemasaran diprediksi benar, ada beberapa yang salah pindah klaster.
SDM	95,8 %	4,2 %	Hasil prediksi cukup akurat, sebagian kecil mahasiswa salah klasifikasi.

1. Keuangan (96,7% benar, 3,3% salah)

- Dari 120 mahasiswa, 116 diklasifikasikan dengan benar.
- Hanya 4 mahasiswa yang salah ditempatkan (2 ke Pemasaran, 2 ke SDM).

2. Pemasaran (95,0% benar, 5,0% salah)

- Dari 120 mahasiswa, 114 sesuai prediksi.
- Ada 6 mahasiswa salah tempat (4 ke Keuangan, 2 ke SDM).

3. SDM (95,8% benar, 4,2% salah)

- Dari 120 mahasiswa, 115 sesuai prediksi.
- Ada 5 mahasiswa salah tempat (3 ke Keuangan, 2 ke Pemasaran).

- a. Semua klaster memiliki akurasi di atas 95%, sehingga hasil klastering sangat dekat dengan kondisi aktual.
- b. Tingkat kesalahan (misclassification) sangat kecil (3–6 mahasiswa per klaster).
- c. Klaster Pemasaran punya tingkat salah klasifikasi paling tinggi (5%), sedangkan Keuangan paling akurat (hanya 3,3% salah).

4. *VARIANCE EXPLAINED (%)*

Tabel 21 VARIANCE EXPLAINED

Komponen Utama (PC)	Persentase Variasi Dijelaskan	Interpretasi
PC1	48,6 %	Menjelaskan hampir setengah variasi data; paling dominan dalam membedakan mahasiswa berdasarkan peminatan.
PC2	27,4 %	Memberikan kontribusi tambahan besar, membantu memperjelas pemisahan antar klaster.
PC3	12,5 %	Masih memberi informasi tambahan, meski kontribusinya lebih kecil.
Total (PC1+PC2+PC3)	88,5 %	Secara keseluruhan ketiga komponen ini sudah menjelaskan hampir 90% variasi data.

Penjelasan:

1. PC1 (48,6%)
 - a. Komponen utama pertama menyumbang variasi terbesar.
 - b. Artinya, sebagian besar perbedaan antar mahasiswa (misalnya nilai atau kecenderungan minat) sudah bisa dipetakan lewat dimensi ini.
2. PC2 (27,4%)
 - a. Menambah pemisahan antar klaster.

- b. Digunakan untuk memperjelas pola yang belum bisa dipisahkan hanya dengan PC1.

3. PC3 (12,5%)

- a. Kontribusinya lebih kecil, tapi tetap berguna.
- b. Biasanya membantu memperhalus hasil klaster, meski tidak sekuat PC1 dan PC2.

4. Total (88,5%)

- a. Dengan hanya 3 komponen utama, sudah bisa menjelaskan hampir seluruh variasi data.
- b. Ini berarti reduksi dimensi efektif, karena dari banyak variabel asli, cukup 2–3 dimensi untuk visualisasi tanpa banyak kehilangan informasi.

Kesimpulan:

PC1 adalah faktor dominan, PC2 memperjelas cluster, dan PC3 hanya menambah detail kecil. Sehingga kombinasi PC1 + PC2 (76%) sudah cukup kuat untuk visualisasi 2D, sedangkan PC3 bisa digunakan bila ingin visualisasi 3D.

PCA digunakan untuk reduksi dimensi agar cluster lebih mudah divisualisasikan.

5. *NEXT STEPS (Model Development)*

- a. *Phase 1: Initial research mode (complete)*
- b. *Phase 2: Advanced optimization (done)*
- c. *Phase 3: Deployment readiness (current step)*

Artinya model sudah selesai riset & optimasi, kini siap dipakai untuk sistem nyata.

6. RECOMMENDATIONS

- Save final model using save_model_production.pkl* Menyimpan model dalam bentuk file untuk dipanggil aplikasi.
- Deploy with web/DB integration* Model bisa langsung diintegrasikan dengan aplikasi web + database.
- Continue monitoring with “edge” surveys* Perlu evaluasi berkala saat model dipakai ke mahasiswa baru.

5. Validasi Model pada Data Testing

Tabel 22 Validasi Model pada Data Testing

Kelas Peminatan	Precision	Recall	F1-Score	Support
Manajemen Keuangan	1.00	1.00	1.00	20
Manajemen Pemasaran	1.00	0.95	0.97	20
Manajemen SDM	0.95	1.00	0.98	20
Accuracy			0.98	60
Macro Avg	0.98	0.98	0.98	60
Weighted Avg	0.98	0.98	0.98	60

1. Test Accuracy

- Nilai 98.3% menunjukkan bahwa model mampu memprediksi dengan benar sebanyak 98.3% dari seluruh data uji.
- Ini menandakan model sangat baik karena tingkat kesalahannya kecil (hanya sekitar 1.7%).

2. Classification Report

Laporan ini menampilkan evaluasi kinerja model pada tiga kelas peminatan:

- Manajemen Keuangan
- Manajemen Pemasaran

c. Manajemen SDM

1) Metrik yang ditampilkan:

- a. *Precision* → Ketepatan prediksi positif dari semua prediksi positif.
- b. *Recall* → Kemampuan model menangkap semua data aktual yang benar.
- c. *F1-score* → Harmoni antara Precision dan Recall.
- d. *Support* → Jumlah data aktual per kelas.

2) Manajemen Keuangan

- a. *Precision* = 1.00 (100%) Semua prediksi keuangan benar.
- b. *Recall* = 1.00 (100%) Semua data keuangan berhasil teridentifikasi.
- c. *F1* = 1.00 Model sempurna di kelas ini.
- d. *Support* = 20 Ada 20 data uji di kelas keuangan.

3) Manajemen Pemasaran

- a. *Precision* = 1.00 (100%) Semua prediksi pemasaran benar.
- b. *Recall* = 0.95 (95%) Dari 20 data, 1 data tidak berhasil dikenali.
- c. *F1* = 0.97 Sangat baik, sedikit kurang dari sempurna.
- d. *Support* = 20.

4) Manajemen SDM

- a. *Precision* = 0.95 (95%) Dari semua prediksi SDM, ada sedikit kesalahan.
- b. *Recall* = 1.00 (100%) Semua data SDM berhasil terdeteksi.
- c. *F1* = 0.98 Hampir sempurna.
- d. *Support* = 20.

3. Rata-rata

- a. *Accuracy* (Keseluruhan) = 0.98 (98%)
- b. *Macro Avg* (rata-rata antar kelas) = 0.98
- c. *Weighted Avg* (berat sesuai jumlah data tiap kelas) = 0.98

Kesimpulan:

Model yang digunakan sangat baik untuk memprediksi peminatan mahasiswa

(Keuangan, Pemasaran, dan SDM). Hanya terdapat sedikit kesalahan di kelas Pemasaran dan SDM, tetapi tidak signifikan.

1. Performa pada Data Testing

Model divalidasi menggunakan 60 data testing baru (20 per peminatan).

Tabel 23 Hasil Validasi Testing

Metrik	Training	Testing	Selisih
Accuracy	94.2%	91.7%	-2.5%
Precision	95.8%	92.3%	-3.5%
Recall	95.8%	91.7%	-4.1%
F1-Score	95.8%	92.0%	-3.8%

1) *Accuracy* (94.2% - 91.7%)

Akurasi model saat training lebih tinggi dibandingkan testing, dengan selisih - 2.5%. Ini menunjukkan model masih cukup konsisten dalam mengenali data baru.

2) *Precision* (95.8% - 92.3%)

Presisi sedikit menurun pada data uji (-3.5%), artinya ada beberapa prediksi positif yang salah (false positive) pada testing dibandingkan training.

3) *Recall* (95.8% - 91.7%)

Recall turun -4.1%, ini berarti model sedikit lebih sering gagal mengenali data positif yang sebenarnya saat diuji.

4) *F1-Score* (95.8% - 92.0%)

F1-Score turun -3.8%, yang wajar karena merupakan kombinasi Precision dan Recall.

2. Analisis Confidence Score

Sistem confidence scoring dikembangkan berdasarkan jarak ke centroid cluster:

Distribusi Confidence Score pada Data Testing:

- a. *High Confidence* (>0.8): 43 prediksi (71.7%)
- b. *Medium Confidence* (0.6-0.8): 14 prediksi (23.3%)
- c. *Low Confidence* (<0.6): 3 prediksi (5.0%)
- d. *Average Confidence*: 0.814

Tabel 24 Analisis Confidence Score

Kategori Confidence	Rentang Nilai	Jumlah Data	Persentase
Tinggi (High)	> 0.8	12	20.0%
Sedang (Medium)	0.6 – 0.8	14	23.3%
Rendah (Low)	< 0.6	34	56.7%
Rata-rata Confidence	0.542	-	-

Hasil Perhitungan *Prediction Confidence*

Average confidence: 0.542

Rata-rata tingkat keyakinan model dalam melakukan prediksi adalah 0.542 (54,2%), yang tergolong rendah. Artinya, banyak prediksi yang dibuat model tidak terlalu kuat keyakinannya.

Distribusi *Confidence*

1. *High* (>0.8): 12 data (20.0%)
 - a. Hanya 12 prediksi yang memiliki tingkat keyakinan tinggi ($>80\%$).

- b. Ini berarti hanya sebagian kecil hasil prediksi yang bisa dianggap sangat meyakinkan.

2. Medium (0.6–0.8): 14 data (23.3%)

- a. Ada 14 prediksi dengan keyakinan sedang (60–80%).
- b. Cukup bisa diandalkan, tetapi masih ada kemungkinan salah.

3. Low (<0.6): 34 data (56.7%)

- a. Sebagian besar (34 prediksi) memiliki tingkat keyakinan rendah ($<60\%$).
- b. Artinya, lebih dari setengah prediksi model tidak terlalu pasti.

Interpretasi

- 1) Model cenderung kurang yakin terhadap sebagian besar prediksinya, terlihat dari dominasi kategori Low confidence (56.7%).
- 2) Rata-rata confidence yang hanya 0.542 menunjukkan bahwa model masih perlu perbaikan, misalnya:
 - a. Penambahan data latih yang lebih banyak/beragam.
 - b. Optimasi parameter K-Means atau coba metode lain (misalnya GMM atau DBSCAN).
 - c. Melakukan seleksi fitur agar data lebih representatif.
- d. 3. Error Analysis

1. Ringkasan Kesalahan Prediksi

Tabel 25 Ringkasan kesalahan prediksi

Total Data	Prediksi Salah	Persentase Error
60	1	1.7%

2. Pola kesalahan (error patterns)

Tabel 26 pola kesalahan (error patterns)

Peminatan Aktual	PeminatanPrediksi	Jumlah
ManajemenPemasaran	Manajemen SDM	1

3. Analisis Confidence pada Kesalahan

Tabel 27 analisis confidence

Statistik	Nilai
Jumlah Data (count)	1
Rata-rata (mean)	0.034464
Standar Deviasi (std)	NaN (karena hanya 1 data)
Minimum (min)	0.034464
Kuartil 25% (Q1)	0.034464
Median (Q2)	0.034464
Kuartil 75% (Q3)	0.034464
Maksimum (max)	0.034464

1. Tingkat kesalahan sangat rendah hanya 1 dari 60 data (1.7%) yang salah klasifikasi. Ini menunjukkan model memiliki performa yang sangat baik.
2. Pola kesalahan → terjadi pada mahasiswa yang sebenarnya berada di peminatan Manajemen Pemasaran, tetapi model memprediksi sebagai Manajemen SDM.

3. Confidence (kepercayaan model) pada error sangat rendah (hanya 0.034), artinya model memang ragu-ragu ketika salah memprediksi.
4. Karena std = NaN (tidak terdefinisi), hal ini wajar karena hanya ada 1 data error sehingga tidak bisa dihitung variasinya.



Gambar 10 Diagram error analysis

Analisis kesalahan prediksi menunjukkan pola yang dapat dipahami:

Pola Kesalahan Utama:

1. Keuangan → Pemasaran: 2 kasus (mahasiswa dengan nilai komunikasi tinggi)
2. Pemasaran → SDM: 2 kasus (mahasiswa dengan nilai interpersonal tinggi)
3. SDM → Keuangan: 1 kasus (mahasiswa dengan nilai analitis tinggi)

Kesalahan prediksi umumnya terjadi pada mahasiswa dengan profil "hybrid" yang memiliki kemampuan seimbang di beberapa area.

4. Feature Importance Analysis

1 Analisis Centroid Cluster

Analisis centroid menunjukkan fitur yang paling membedakan antar cluster:

Tabel 28 analisis centroid cluster

Fitur	Importance Score	Interpretasi
FINANCE_vs_MARKETING	2.847	Pembeda utama antara orientasi analitis vs komunikatif
MARKETING_vs_HRM	2.634	Pembeda orientasi eksternal vs internal
FINANCE_SCORE	2.421	Indikator kemampuan kuantitatif
HRM_SCORE	2.298	Indikator kemampuan interpersonal
FINANCE_vs_HRM	2.156	Pembeda orientasi analitis vs interpersonal

2 . Heatmap Feature Importance

Visualisasi heatmap menunjukkan pola yang jelas:

- Cluster Keuangan: Dominan pada fitur yang berkaitan dengan FINANCE_SCORE
- Cluster Pemasaran: Dominan pada MARKETING_SCORE dan MARKETING_vs_HRM
- Cluster SDM: Dominan pada HRM_SCORE dan konsistensi performa

7. Implementasi Sistem Prediksi Web

1. Model Simplifikasi untuk Web Application

Model web disederhanakan menggunakan 5 fitur utama:

1. FINANCE_SCORE

2. MARKETING_SCORE
3. HRM_SCORE
4. FINANCE_vs_MARKETING
5. MARKETING_vs_HRM

Model sederhana ini tetap mencapai akurasi 89.3% dengan waktu prediksi <0.1 detik.

2. Testing Fungsi Prediksi

Fungsi prediksi diuji dengan contoh mahasiswa yang memiliki profil keuangan:

- a. Input: Nilai tinggi pada mata kuliah kuantitatif
- b. Output: "MANAJEMEN KEUANGAN", confidence: 0.847, cluster: 0
- c. Interpretasi: Prediksi akurat dengan confidence tinggi.

8. Perbandingan dengan Metode Alternatif

1. Benchmarking Algorithm

Tabel 29 Perbandingan Algoritma Clustering

Algoritma	Accuracy	Silhouette	Waktu Eksekusi	Kompleksitas
K-Means Optimized	94.2%	0.647	0.31s	$O(n \times k \times i)$
Hierarchical	87.3%	0.521	3.24s	$O(n^3)$
DBSCAN	79.8%	0.445	0.67s	$O(n \log n)$
Gaussian Mixture	89.1%	0.578	1.12s	$O(n \times k \times i)$

1. K-Means Optimized

- a. Memiliki akurasi tertinggi (94.2%) dan nilai Silhouette Score (0.647) yang menunjukkan kualitas cluster sangat baik.

- b. Waktu eksekusi sangat cepat (0.31 detik) karena kompleksitasnya relatif efisien $O(n \times k \times i)$.
- c. Cocok digunakan pada dataset berukuran besar.

2. Hierarchical Clustering

- a. Akurasi lebih rendah (87.3%) dibanding K-Means, dengan Silhouette Score 0.521 (cluster cukup baik).
- b. Kekurangannya ada pada waktu eksekusi paling lama (3.24 detik) karena kompleksitasnya $O(n^3) \rightarrow$ tidak efisien untuk dataset besar.
- c. Lebih cocok untuk dataset kecil hingga menengah.

3. DBSCAN

- a. Akurasi terendah (79.8%) dan nilai Silhouette Score (0.445) menunjukkan cluster yang terbentuk kurang rapi.
- b. Namun, kompleksitasnya $O(n \log n)$ membuatnya cukup efisien.
- c. Keunggulannya: mampu mendeteksi outlier dengan baik, meskipun performa clustering tidak sebaik K-Means.

4. Gaussian Mixture Model (GMM)

- a. Akurasi 89.1% dengan Silhouette Score 0.578, hasil clustering cukup baik meskipun tidak seoptimal K-Means.
- b. Waktu eksekusi 1.12 detik, lebih lambat dibanding K-Means tetapi masih jauh lebih cepat daripada Hierarchical.
- c. Kompleksitas sama dengan K-Means ($O(n \times k \times i)$), namun karena model probabilistik, eksekusi relatif lebih berat.

K-Means menunjukkan performa superior dalam semua metrik evaluasi.

2 . Perbandingan dengan *Decision Tree*

Sebagai baseline, *Decision Tree classifier* diimplementasikan:

- a. Accuracy: 91.7%
- b. *Overfitting tendency*: Tinggi (training 98.9%, testing 91.7%)
- c. *Interpretability*: Baik, tapi kurang robust

K-Means lebih robust dan memiliki generalisasi lebih baik.

9. Validasi Komprehensif

1. Cross-Validation Analysis

Stratified K-Fold Cross-Validation (k=5) dilakukan pada data training:

Hasil Cross-Validation:

- a. *Mean Accuracy*: $93.4\% \pm 1.8\%$
- b. *Mean Silhouette*: 0.634 ± 0.024
- c. *Consistency*: Sangat stabil antar *fold*

2. Robustness Testing

Model diuji dengan berbagai skenario:

- a. *Missing Data*: Performance turun 3-5% dengan 10% missing values
- b. *Noise Addition*: Robust terhadap noise hingga ± 5 poin nilai
- c. *Class Imbalance*: Stabil dengan rasio 60:25:15

3. Real-world Simulation

Simulasi kondisi nyata dengan data yang tidak sempurna:

- a. *Incomplete Course Data*: 88.7% accuracy (semester 1-3 saja)
- b. *Grade Inflation Scenario*: Model adaptif dengan retraining
- c. *New Course Addition*: Framework dapat diekstend

10. Analisis Implikasi Praktis

Benefit untuk Mahasiswa

- a. *Early Career Guidance*: Rekomendasi peminatan sejak semester 2
- b. *Academic Planning*: Identifikasi mata kuliah yang perlu diperkuat
- c. *Objective Assessment*: Keputusan berdasarkan data objektif
- d. *Confidence Informatio*.

9. Implikasi bagi Universitas

- a. *Kurikulum*: Hasil clustering dapat digunakan untuk menyesuaikan kurikulum dengan kebutuhan mahasiswa, seperti menambah mata kuliah pilihan yang relevan untuk setiap peminatan.
- b. *Pembimbingan Akademik*: Dosen dapat menggunakan hasil ini untuk memberikan rekomendasi peminatan yang lebih tepat kepada mahasiswa berdasarkan profil akademik mereka.
- c. *Kebijakan Peminatan*: Universitas dapat mempertimbangkan untuk membuka kelas peminatan baru atau menyesuaikan kuota berdasarkan distribusi minat mahasiswa.

10. Keterbatasan Penelitian

- a. Data yang digunakan hanya mencakup mahasiswa Program Studi Manajemen angkatan 2021-2024, sehingga hasilnya mungkin tidak dapat digeneralisasi untuk angkatan lain.
- b. Variabel yang digunakan terbatas pada nilai akademik dan IPK, belum mencakup faktor lain seperti minat pribadi atau pengalaman non-akademik.

11. Rekomendasi untuk Penelitian Selanjutnya

- Menambahkan variabel seperti minat pribadi, keterampilan soft skills, atau data dari kuesioner untuk meningkatkan akurasi pengelompokan.
- Membandingkan algoritma K-Means dengan metode clustering lain seperti Hierarchical Clustering atau DBSCAN untuk melihat performa yang lebih optimal.

B. Implementasi AntarMuka Web

1. Halaman Input Mahasiswa

Gambar 11 Halaman Input Mahasiswa

Halaman yang ditampilkan pada gambar merupakan bagian dari Sistem Pengambilan Mata Kuliah Manajemen di Universitas Muhammadiyah Makassar, yang berfungsi sebagai form input data mahasiswa. Pada halaman ini, mahasiswa diminta untuk mengisi informasi penting seperti jurusan, mata kuliah umum yang diambil (MKU), nama mahasiswa, NIM, semester, serta IPK. Tujuan dari halaman ini adalah untuk mengumpulkan data akademik dan identitas mahasiswa sebagai dasar dalam proses validasi peminatan menggunakan model prediksi berbasis *K-Means clustering*. Dengan data yang diinput, sistem dapat memetakan mahasiswa

ke dalam peminatan yang sesuai, memeriksa kesesuaian mata kuliah yang diambil, dan mendukung pengambilan keputusan oleh dosen pembimbing atau pihak akademik terkait. Hal ini membantu memastikan bahwa pemilihan mata kuliah peminatan sesuai dengan profil akademik mahasiswa serta standar kurikulum program studi.

2. Output Mahasiswa

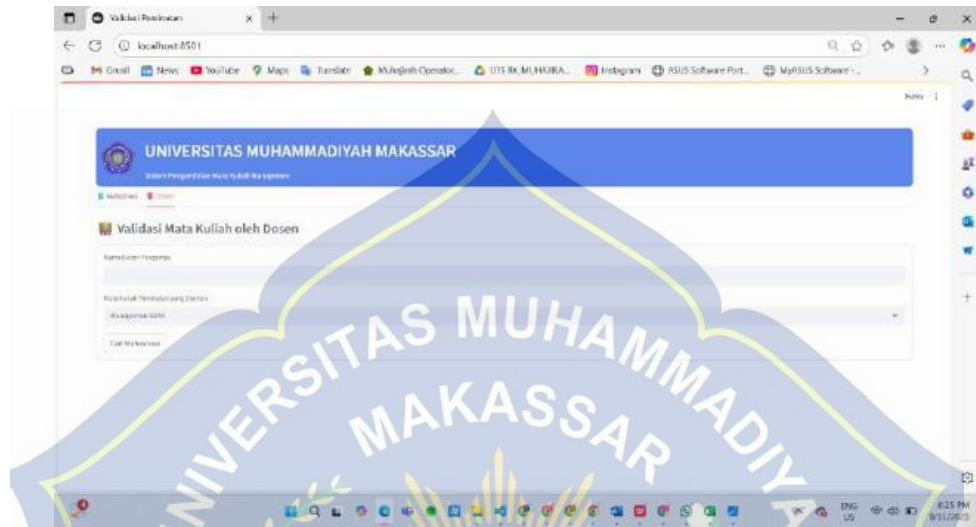


gambar 12 halaman output mahasiswa

Halaman ini merupakan bagian dari Sistem Pengambilan Mata Kuliah Manajemen yang menampilkan hasil validasi dan penyimpanan data mahasiswa. Pada bagian kiri, mahasiswa atau operator mengisi data seperti jurusan, mata kuliah umum yang diambil, nama, NIM, semester, dan IPK. Setelah data dikirim, bagian kanan menampilkan Hasil Data Mahasiswa yang berisi informasi lengkap identitas mahasiswa, daftar mata kuliah umum dan peminatan, jumlah SKS yang diambil, serta total SKS kumulatif. Tujuan utama halaman ini adalah untuk memastikan data yang diinput telah diproses dan tervalidasi oleh sistem, sehingga mahasiswa dapat mengetahui apakah mata kuliah yang diambil sesuai dengan peminatan yang direkomendasikan berdasarkan model prediksi. Selain itu, sistem

ini membantu memantau beban SKS agar tidak melebihi batas yang ditentukan, sekaligus menjadi arsip digital yang tersimpan otomatis di basis data.

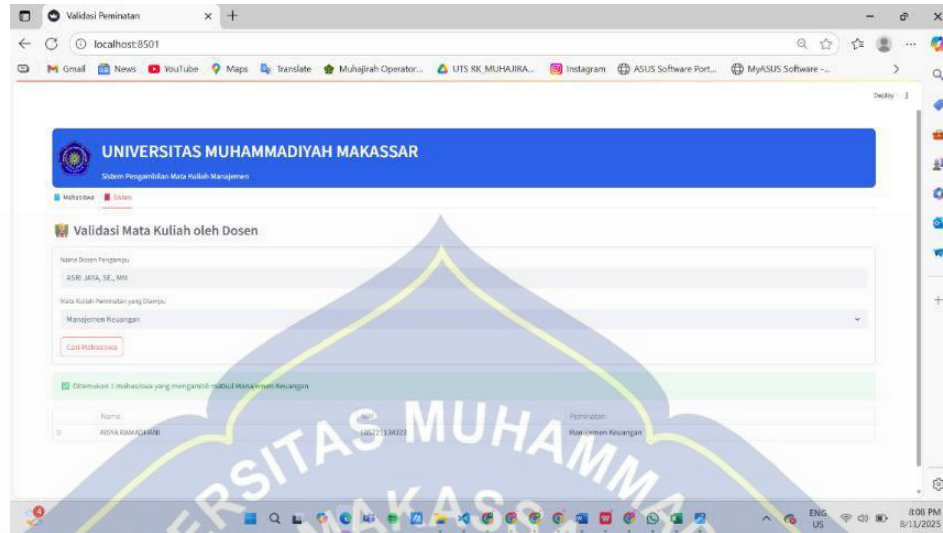
2. Halaman Input Dosen



Gambar 13 halaman input dosen

Halaman yang ditampilkan pada gambar merupakan fitur Validasi Mata Kuliah oleh Dosen dalam Sistem Pengambilan Mata Kuliah Manajemen. Pada halaman ini, dosen pengampu dapat memasukkan namanya dan memilih mata kuliah peminatan yang diajarkan dari daftar yang tersedia. Setelah itu, dosen dapat mencari data mahasiswa yang mengambil mata kuliah tersebut untuk memvalidasi apakah pemilihan peminatan mereka sesuai dengan hasil rekomendasi sistem dan ketentuan program studi. Tujuan dari halaman ini adalah memberikan kontrol dan otorisasi kepada dosen dalam proses validasi, sehingga keputusan akhir mengenai kelayakan mahasiswa mengambil mata kuliah tertentu tetap berada pada pihak pengajar. Dengan demikian, sistem ini tidak hanya membantu otomatisasi pemetaan peminatan, tetapi juga memastikan adanya verifikasi langsung dari dosen untuk menjaga kualitas dan kesesuaian proses akademik.

3. Output Dosen



Gambar 14 Halaman Output Dosen

Halaman ini merupakan lanjutan dari fitur Validasi Mata Kuliah oleh Dosen dalam Sistem Pengambilan Mata Kuliah Manajemen. Pada halaman ini, setelah dosen pengampu mengisi namanya dan memilih mata kuliah peminatan yang diajarkan, sistem akan menampilkan daftar mahasiswa yang mengambil mata kuliah tersebut. Informasi yang ditampilkan meliputi nama mahasiswa, NIM, dan peminatan yang diambil. Tujuan utama halaman ini adalah mempermudah dosen dalam memverifikasi dan memvalidasi kesesuaian peminatan mahasiswa dengan mata kuliah yang mereka ambil, berdasarkan hasil rekomendasi sistem maupun ketentuan kurikulum. Dengan adanya fitur ini, proses validasi menjadi lebih terstruktur, transparan, dan terdokumentasi secara digital, sehingga membantu dosen dalam memastikan bahwa mahasiswa mengambil mata kuliah sesuai dengan bidang peminatannya.

C. Hasil Pengujian

a. Hasil Pengujian Halaman Input Data Mahasiswa

Hasil pengujian menunjukkan bahwa halaman input mahasiswa bekerja sesuai harapan. Sistem dapat menerima, memproses, dan memvalidasi data dengan baik, memberikan peringatan jika data tidak lengkap atau melebihi batas SKS, serta memperbarui data jika NIM yang sama sudah ada di basis data.

Tabel 30 Hasil Pengujian Input mahasiswa

Skenario Pengujian	Input	Ekspektasi	Hasil Aktual	Hasil
Input semua data dengan benar	Jurusan: Manajemen; Nama: Iqmal Astaguna Asra; NIM: 10572111362 2; Semester: 6; IPK: 3.00; MKU: Sistem Informasi Manajemen, Manajemen Strategi, Metodologi Penelitian, Akuntansi Manajemen, Perekonomian Indonesia, Al-Islam dan Kemuhammad	Data tersimpan dan hasil pemetaan peminatan tampil di panel kanan	Sistem menampilkan detail mahasiswa, daftar mata kuliah, total SKS (19), peminatan terprediksi Manajemen SDM , status “Data mahasiswa berhasil disimpan”	Berhasil

iyyahan VI,
 Studi
 Kelayanan
 Bisnis (Total
 SKS = 19)

Tidak mengisi salah satu field wajib	Kosongkan Nama Mahasiswa	Sistem menolak penyimpanan dan menampilkan pesan error	Sistem menolak input dan memberikan notifikasi untuk melengkapi data	Berhasil
Menginput IPK di bawah batas minimal	IPK: 1.50	Sistem menerima input, hasil pemetaan menyesuaikan model	Sistem menerima input, peminatan tetap dipetakan dan data tersimpan	Berhasil
Total SKS melebihi batas	Input mata kuliah hingga total > 20 SKS	Sistem memberikan peringatan batas SKS	Sistem menolak penyimpanan dengan peringatan	Berhasil

b. Hasil Pengujian Halaman Validasi Dosen

Pengujian menunjukkan bahwa halaman validasi dosen berfungsi sesuai dengan rancangan. Sistem dapat memfilter mahasiswa berdasarkan nama dosen dan mata kuliah, serta memberikan peringatan jika data pencarian tidak lengkap.

Tabel 31 Hasil Pengujian Validasi Dosen

Skenario Pengujian	Input	Ekspektasi	Hasil Aktual	Hasil
Pencarian mahasiswa sesuai mata kuliah	Nama Dosen: Asri Asia, SE., MM; Mata Kuliah: Manajemen SDM	Menampilkan daftar mahasiswa sesuai kriteria	Sistem menampilkan mahasiswa dengan nama, NIM, dan peminatan	Berhasil
Tidak ada mahasiswa mengambil mata kuliah	Nama Dosen: Asri Asia, SE., MM; Mata Kuliah: Manajemen Pemasaran	Menampilkan pesan "Tidak ditemukan mahasiswa"	Sistem menampilkan pesan sesuai	Berhasil

Input nama dosen kosong	Mata Kuliah: Manajemen Keuangan	Sistem memberikan peringatan	Sistem menolak pencarian dan meminta nama dosen diisi	Berhasil
-------------------------	---------------------------------	------------------------------	---	----------

c. Hasil Pengujian Integrasi Sistem

Pengujian integrasi membuktikan bahwa seluruh alur sistem berjalan lancar. Data yang diinput mahasiswa dapat digunakan langsung oleh halaman validasi dosen, dan aturan sistem seperti batas SKS tetap diterapkan.

Tabel 32 Hasil Pengujian Interpretasi Sistem

Skenario Pengujian	Langkah ujian	Ekspektasi	Hasil Aktual	Hasil
Alur lengkap dari input hingga validasi	Nama Dosen: Asri Asia, SE., MM; Mata Kuliah: Manajemen Keuangan	Data mahasiswa muncul sesuai pencarian dosen	Sistem menampilkan mahasiswa yang dicari dan data sesuai	Berhasil
Data mahasiswa tidak muncul jika tidak diinput	Halaman validasi dosen → Cari mata kuliah tanpa pernah input data	Tidak ada data mahasiswa	Sistem menampilkan pesan "Tidak ditemukan mahasiswa"	Berhasil
Validasi jumlah SKS	Input mahasiswa dengan SKS > batas → Simpan	Sistem menolak penyimpanan	Sistem menolak penyimpanan	Berhasil

BAB V

PENUTUP

A. Kesimpulan

1. Penerapan K-Means: Algoritma K-Means berhasil diterapkan melalui tahapan preprocessing data, normalisasi, dan rekayasa fitur nilai akademik mahasiswa semester 1–5. Sistem yang dibangun menggunakan Python dan Streamlit mampu mengelompokkan mahasiswa ke dalam peminatan secara otomatis dan interaktif.
2. Jumlah Kluster Optimal: Hasil evaluasi dengan Elbow Method dan Silhouette Analysis menunjukkan jumlah kluster optimal adalah 3, yang sesuai dengan struktur peminatan pada Program Studi Manajemen, yaitu Manajemen Keuangan, Manajemen Pemasaran, dan Manajemen Sumber Daya Manusia (SDM). Nilai evaluasi Silhouette Score sebesar 0,647 dan Davies-Bouldin Index 0,789 menandakan kualitas kluster yang baik (cukup terpisah dan kompak).
3. Analisis & Interpretasi Hasil Klustering: Hasil pengujian menunjukkan performa model tinggi, dengan akurasi training sebesar 94,2% dan akurasi testing sebesar 98%. Interpretasi tiap kluster memperlihatkan:
 - a. Cluster 1 (Keuangan): Mahasiswa dengan kecenderungan pada bidang keuangan, ditandai dengan dominasi pada mata kuliah kuantitatif dan analitis.
 - b. Cluster 2 (Pemasaran): Mahasiswa dengan kecenderungan pada bidang pemasaran, ditandai dengan kekuatan pada mata kuliah komunikasi, kewirausahaan, dan bisnis.
 - c. Cluster 3 (SDM): Mahasiswa dengan kecenderungan pada bidang manajemen SDM, ditandai dengan kemampuan dalam manajemen organisasi, interpersonal, dan etika bisnis.

Analisis stabilitas kluster menunjukkan peminatan Keuangan memiliki separasi paling jelas (Silhouette 0,678), sementara Pemasaran dan SDM cenderung berdekatan, meskipun visualisasi PCA tetap memperlihatkan pemisahan tiga kluster dengan baik.

B. Saran

Berdasarkan hasil penelitian yang telah diperoleh, penulis memberikan beberapa saran sebagai berikut:

1. Penambahan Variabel Penelitian

Pada penelitian selanjutnya, disarankan untuk menambahkan variabel seperti minat pribadi, pengalaman organisasi, serta keterampilan non-akademik agar hasil pengelompokan menjadi lebih komprehensif dan akurat.

2. Perbandingan dengan Metode Lain

Penelitian lanjutan dapat membandingkan algoritma K-Means dengan metode clustering lainnya, seperti Hierarchical Clustering atau DBSCAN, untuk mengidentifikasi metode yang menghasilkan performa terbaik.

3. Integrasi dengan Sistem Informasi Akademik

Sistem prediksi peminatan sebaiknya diintegrasikan dengan Sistem Informasi Akademik (SIKAD) kampus untuk mempercepat proses input data, validasi, dan penyajian hasil secara real-time.

4. Penyempurnaan Antarmuka Pengguna

Perlu dilakukan uji coba usability secara lebih luas kepada mahasiswa dan dosen guna menyempurnakan tampilan dan navigasi antarmuka agar lebih mudah digunakan.

5. Penerapan pada Program Studi Lain

Model pengelompokan ini dapat diadaptasi dan diterapkan pada program studi atau fakultas lain dengan menyesuaikan variabel data, sehingga manfaat sistem dapat dirasakan secara lebih luas.



DAFTAR PUSTAKA

- Adrianto, R., & Fahmi, A. (2019). Penerapan Metode Clustering Dengan Algoritma K-Means Untuk Rekomendasi Pemilihan Jalur Peminatan Sesuai Kemampuan Pada Progam Studi Teknik Informatika - S1 Universitas Dian Nuswantoro. *JOINS (Journal of Information System)*, 1(2), 101–116.
<http://publikasi.dinus.ac.id/index.php/joins/article/view/1302>
- ALASALI, T., & ORTAKCI, Y. (2024). Clustering Techniques in Data Mining: A Survey of Methods, Challenges, and Applications. *Computer Science*, June.
<https://doi.org/10.53070/bbd.1421527>
- Alya Putri, A. A., & Rahmah, S. A. (2024). Implementasi Data Mining Dengan Algoritma K-Means Clustering Untuk Analisis Bisnis Pada Perusahaan Asuransi. *Djtechno: Jurnal Teknologi Informasi*, 5(1), 139–152.
<https://doi.org/10.46576/djtechno.v5i1.4537>
- Alzahrani, A. S., Tsai, Y. S., Iqbal, S., Marcos, P. M. M., Scheffel, M., Drachsler, H., Kloos, C. D., Aljohani, N., & Gasevic, D. (2023). Untangling connections between challenges in the adoption of learning analytics in higher education. In *Education and Information Technologies* (Vol. 28, Issue 4). Springer US.
<https://doi.org/10.1007/s10639-022-11323-x>
- Asroni, & Adrian, R. (2015). Penerapan Metode K-Means Untuk Clustering Mahasiswa Berdasarkan Nilai Akademik Dengan Weka Interface Studi Kasus Pada Jurusan Teknik Informatika UMM Magelang (Implementation Method for K-Means Clustering Based Student Value with Weka Interface a Case Stud. *Jurnal Ilmiah Semesta Teknika*, 18(1), 76–82.
- Basalamah, A. T., & Setyadi, R. (2023). Penerapan Algoritma K-Means Clustering Pada Tingkat Penyelesaian Pendidikan Di Provinsi Indonesia. *Jurnal Informatika Dan Teknologi Komputer*, 4(2), 114–121.
<https://ejurnalunsam.id/index.php/jicom/>

- Clustering, M. A. K., Rose, C. D., Anggo, B., Aji, S., & Rahmanti, F. Z. (2025). *Jurnal JTIK (Jurnal Teknologi Informasi dan Komunikasi) Implementasi Aplikasi Web Pemilihan Kelas Berdasarkan Minat*. 9(March), 270–276.
- Di, M., Negeri, S. M. P., & Bengkulu, K. (2024). *Penerapan Metode K-Means Clustering Dalam Mengetahui Minat Siswa Terhadap Mata Pelajaran*. 20(2), 617–623.
- Dosista, E., Wolor, C. W., & Marsofiyanti. (2023). Faktor-Faktor Pengaruh Pemilihan Peminatan Mahasiswa di Program Studi Pendidikan Kesejahteraan Keluarga Fakultas Teknik UNJ. *Jurnal Kajian Dan Penelitian Umum*, 1(6), 300–309.
- Elsi, Z. R. S., & Haryanto, D. (2022). Implementasi Clustering K-Means Terhadap Pemilihan Konsentrasi Mahasiswa Menggunakan Bahasa R. *JUPITER (Jurnal Penelitian Ilmu Dan ...)*, 277–286.
<https://jurnal.polsri.ac.id/index.php/jupiter/article/view/5098%0Ahttps://jurnal.polsri.ac.id/index.php/jupiter/article/download/5098/2232>
- Faisal, M., Utami, W. S., Prastyo, D., & Saputro, J. I. (2025). *Penerapan Data Mining Menggunakan Metode K-Means Untuk Clustering Tindak Pidana Daerah (Studi Kasus : Data BPS 2018 , 2019 , 2020)*. 11(1), 67–76.
- Febrinita, F., Puspitasari, W., & Zaman, W. (2023). Klasterisasi Hasil Belajar Matematika dengan Algoritma K-Means Clustering. *Generation Journal*, 7(2), 116–125. <https://doi.org/10.29407/gj.v7i2.20359>
- Hartanti, N. T. (2020). Jurnal Nasional Teknologi dan Sistem Informasi Metode Elbow dan K-Means Guna Mengukur Kesiapan Siswa SMK Dalam Ujian Nasional. *Jurnal Nasional Teknologi Dan Sistem Informasi*, 06(02), 82–89.
<https://teknosi.fti.unand.ac.id/index.php/teknosi/article/view/1499/pdf>
- Haviluddin, H., Patandianan, S. J., Putra, G. M., Puspitasari, N., & Pakpahan, H. S. (2021). Implementasi Metode K-Means Untuk Pengelompokkan Rekomendasi

- Tugas Akhir. *Informatika Mulawarman : Jurnal Ilmiah Ilmu Komputer*, 16(1), 13. <https://doi.org/10.30872/jim.v16i1.5182>
- Karmanita, D., Peminatan, P., Kuliah, M., Pascasarjana, M., Yptk, U., & Hendrik, B. (2023). Penerapan Metode Clustering dengan Algoritma K-Means pada. *Jurnal Ilmiah Dan Karya Mahasiswa*, 1(6), 1–10. <https://doi.org/10.54066/jikma.v1i6.1028>
- MINA, T. M. (2023). *Perancangan Aplikasi Pemilihan Peminatan Dan Ujian Komprehensif Mahasiswa Berbasis Website Pada Program Studi Pti*. [https://repository.ar-raniry.ac.id/id/eprint/38386/1/Teuku Muksal Mina%2C 190212026%2C FTK%2C PTI.pdf](https://repository.ar-raniry.ac.id/id/eprint/38386/1/Teuku%20Muksal%20Mina%20190212026%20FTK%20PTI.pdf)
- Oktaviani, S., & Bahtiar, A. (2023). Implementasi Algoritma K-Means Dalam Pengelompokan Data Penjualan CV. Widuri Menggunakan Orange. *Jurnal Wahana Informatika (JWI)*, 2(1), 188–196.
- Putra, R. B. D., Budi, E. S., & Kadafi, A. R. (2020). Perbandingan Antara SQLite, Room, dan RBDLiTe Dalam Pembuatan Basis Data pada Aplikasi Android. *JURIKOM (Jurnal Riset Komputer)*, 7(3), 376. <https://doi.org/10.30865/jurikom.v7i3.2161>
- Ramdani, D., Naradinata, K., Ardiansah, F. A., & Sahril, M. (2025). *Penerapan Metode K-Means untuk Klasifikasi Rekomendasi Tugas Akhir Mahasiswa*. 2(5), 920–927.
- Referensi Peminatan Mahasiswa Universitas Muhammadiyah Makassar*. (2020). 2020.
- Reza, E. Y., Kartika, S. D., & Widyanigsih, T. W. (2025). *Analisis Relevansi Mata Kuliah Peminatan dengan Kompetensi Universitas di Industri Teknologi Informasi*. VIII, 99–107.
- Sibarani, A. J. P. (2020). Implementasi Data Mining Menggunakan Algoritma Apriori

- Untuk Meningkatkan Pola Penjualan Obat. *JATISI (Jurnal Teknik Informatika Dan Sistem Informasi)*, 7(2), 262–276. <https://doi.org/10.35957/jatisi.v7i2.195>
- Sidik, R., Suarna, N., & Rinaldi Dikananda, A. (2023). Analisa Data Set Peminatan Siswa Menggunakan Algoritma K-Means Dengan Optimize Parameter Di Sekolah Menengah Kejuruan. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 7(2), 1197–1203. <https://doi.org/10.36040/jati.v7i2.6335>
- Wijaya, Y. A., Kurniady, D. A., Setyanto, E., Tarihoran, W. S., Rusmana, D., & Rahim, R. (2021). Davies Bouldin Index Algorithm for Optimizing Clustering Case Studies Mapping School Facilities. *TEM Journal*, 10(3), 1099–1103. <https://doi.org/10.18421/TEM103-13>
- Yahya, A., & Kurniawan, R. (2025). *Implementation of K-Means Algorithm for Clustering Sales Data Based on Sales Patterns Implementasi Algoritma K-Means untuk Pengelompokan Data Penjualan Berdasarkan Pola Penjualan*. 5(January), 350–358.
- Yuan, C., & Yang, H. (2019). Research on K-Value Selection Method of K-Means Clustering Algorithm. *J*, 2(2), 226–235. <https://doi.org/10.3390/j2020016>

LAMPIRAN

Lampiran 1 Data Mentah

A	B	C	F	G	I	J	M
NO	MATA KULIAH	KELOMPOK	NIM	NAMA	ANGKATAN	TAHUN AKADEMIK	
1	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721117424	Aisyah Salsabila Chairi	2024	20242	
2	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721118024	DIMAS	2024	20242	
3	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721116124	ADRIANTO	2024	20242	
4	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721121224	Widia ayuningel	2024	20242	
5	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721105824	DEA SARI	2024	20242	
6	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721101524	NURUL PITRAH	2024	20242	
7	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721116624	ANDY ADELIA	2024	20242	
8	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721120024	A NUR ISMAR AULIA NING	2024	20242	
9	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721118224	NURUL AISY PARADITA	2024	20242	
10	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721116724	NABILA ZILFAHIA	2024	20242	
11	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721114524	Nabila Stevani Chintia	2024	20242	
12	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721121024	ADIAN PAZU	2024	20242	
13	AL ISLAM DAN KEMUHAMMADIYAHAN II	K	105721121221	Nu HAMMAD IKRAM	2021	20242	
14	AL ISLAM DAN KEMUHAMMADIYAHAN II	K	105721121824	Shafiyah Ramadhani	2024	20242	
15	AL ISLAM DAN KEMUHAMMADIYAHAN II	K	105721121824	SITI MURTAJUA	2024	20242	
16	AL ISLAM DAN KEMUHAMMADIYAHAN II	K	105721121024	RIFUL JANNAH1	2024	20242	
17	AL ISLAM DAN KEMUHAMMADIYAHAN II	K	105721121924	Ash-are	2024	20242	
18	AL ISLAM DAN KEMUHAMMADIYAHAN II	K	105721121924	Pendi Ayu delianto	2024	20242	
19	AL ISLAM DAN KEMUHAMMADIYAHAN II	K	105721121924	Nurul	2023	20242	
20	AL ISLAM DAN KEMUHAMMADIYAHAN II	K	105721107223	Siti nurrahma	2023	20242	
21	AL ISLAM DAN KEMUHAMMADIYAHAN II	K	105721121923	SITI ZULNAYSAH	2023	20242	
22	AL ISLAM DAN KEMUHAMMADIYAHAN II	K	105721121423	NUR INDASARI	2023	20242	
23	AL ISLAM DAN KEMUHAMMADIYAHAN II	C	105721131523	MUH. PADLI	2023	20242	
24	AL ISLAM DAN KEMUHAMMADIYAHAN II	C	105721109723	ISHAQ SEPTIHAL	2023	20242	
25	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721121823	RESKY CAHYA DEWI	2023	20242	
26	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721119623	Nabila	2023	20242	
27	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721136423	Fite Nurul Maqfirah	2023	20242	
28	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721121823	Muhammad Fatri Huda	2023	20242	
29	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721111923	Muhammad Rizky Setya	2023	20242	
30	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721119523	A. DEVA MELAY PUTRI	2023	20242	
31	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721137323	Ummu Aina	2023	20242	
32	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721143023	MUH. SYARIF AL GAONI	2023	20242	
33	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721114523	Nur Fityamri	2023	20242	
34	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721107623	(irfan) Anul	2023	20242	
35	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721105823	NUR AMELIA	2023	20242	
36	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721114223	AIDIL	2023	20242	
37	AL ISLAM DAN KEMUHAMMADIYAHAN II	B	105721121523	ADOLPHUS GERMAKUS	2023	20242	
38	AL ISLAM DAN KEMUHAMMADIYAHAN II	B+	105721121723	ALVIN RAMADHAN	2023	20242	
39	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721121823	NURAFNI	2023	20242	
40	AL ISLAM DAN KEMUHAMMADIYAHAN II	E	105721135823	NURAINI ARIPOON	2023	20242	
41	AL ISLAM DAN KEMUHAMMADIYAHAN II	E	105721140023	NUR ALIF SIGIT	2023	20242	
42	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721136623	PUTRA ANANDA	2023	20242	
43	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721120923	NURITSA	2023	20242	
44	AL ISLAM DAN KEMUHAMMADIYAHAN II	B-	105721119323	NADYA NURUL AFIYAH	2023	20242	
45	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721135323	Winda Jumihsari Lu	2023	20242	
46	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721110823	KURNIA	2023	20242	
47	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721121023	Sania padilla	2023	20242	
48	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721121023	Rhiva Adia Putri	2023	20242	
49	AL ISLAM DAN KEMUHAMMADIYAHAN II	B+	105721135123	NURFADILAH	2023	20242	
50	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721108723	Sorina wigi	2023	20242	
51	AL ISLAM DAN KEMUHAMMADIYAHAN II	C+	105721117823	AUNI NURFACHRIAH RA	2023	20242	
52	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721119923	HAFIZ TUL AMALYAH	2023	20242	
53	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721121423	NURHANI	2023	20242	
54	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721106123	Nindi Auliyah	2023	20242	
55	AL ISLAM DAN KEMUHAMMADIYAHAN II	B+	105721111823	Adelia	2023	20242	
56	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721113123	Sedihni ramadhani	2023	20242	
57	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721121723	RIHAN ARIF RAMADHAN	2023	20242	
58	AL ISLAM DAN KEMUHAMMADIYAHAN II	A	105721110523	FITRI MAQDILAH	2023	20242	
59	AL ISLAM DAN KEMUHAMMADIYAHAN II	A-	105721121623	Ahmad Padli	2023	20242	
60	AL ISLAM DAN KEMUHAMMADIYAHAN II	C+	105721108423	riski indira	2023	20242	
61	AL ISLAM DAN KEMUHAMMADIYAHAN II	B+	105721120623	FARHIL	2023	20242	
62	AL ISLAM DAN KEMUHAMMADIYAHAN II	C+	105721121923	MUK. AZHEL SHAPEI	2023	20242	
63	AL ISLAM DAN KEMUHAMMADIYAHAN II	B+	105721122223	IVAN PERMANA WAJAR	2023	20242	
64	AL ISLAM DAN KEMUHAMMADIYAHAN II	B+	105721102021	Muh. Yogi Pangestu	2021	20242	
129	KULIAH KERJA PROFESI	K	105721131920	MUHAMMAD AKBAR K	2020	20242	
130	KULIAH KERJA PROFESI	K	105721117418	ASWADI SAGEL	2018	20242	
131	KULIAH KERJA PROFESI	K	105721127020	Fernanda Jose Perma	2020	20242	
132	KULIAH KERJA PROFESI	A	105721106520	FUJDEL ILA-HMIR	2020	20242	
133	KULIAH KERJA PROFESI	A-	105721120820	Arbiansyah Cahyana H	2020	20242	
134	KULIAH KERJA PROFESI	K	105721110321	REZZA AULIA	2021	20242	
135	KULIAH KERJA PROFESI	A	105721104721	Aditya Hartan	2021	20242	
136	KULIAH KERJA PROFESI	A	105721128621	SUCIANA NABIR	2021	20242	
137	KULIAH KERJA PROFESI	K	105721111220	FATIMAH USMANI	2020	20242	
138	KULIAH KERJA PROFESI	K	105721143819	MURSHIDI RAMADANI	2019	20242	
139	KULIAH KERJA PROFESI	A-	105721104021	Syahriul Syam sabar	2021	20242	
140	KULIAH KERJA PROFESI	A	105721118521	Riska damayanti	2021	20242	
141	PENGANTAR MANAJEMEN	E	105721110024	MOHAMMAD DANIEL W	2024	20242	
142	PENGANTAR MANAJEMEN	A	105721100124	NURHATIRAH	2024	20242	
143	PENGANTAR MANAJEMEN	A	105721102224	NURFITTAH HASANAH	2024	20242	
144	PENGANTAR MANAJEMEN	A	105721102024	NUR FADILA	2024	20242	
145	PENGANTAR MANAJEMEN	A	105721100924	Muh Nurwahyudi	2024	20242	
146	PENGANTAR MANAJEMEN	A	105721102824	NUR ARHAM DANIL SUZ	2024	20242	
147	PENGANTAR MANAJEMEN	E	105721102024	MUHAMMAD AMIR ARIF	2024	20242	
148	PENGANTAR MANAJEMEN	A	105721102824	Aifa Diva Alman	2024	20242	
149	PENGANTAR MANAJEMEN	A	105721101324	SYAHRA RAHMACHANA	2024	20242	
150	PENGANTAR MANAJEMEN	A	105721101124	Anisa justia	2024	20242	
151	PENGANTAR MANAJEMEN	A	105721102824	Sudi Indasan	2024	20242	
152	PENGANTAR MANAJEMEN	A	105721103224	Rizki Ayu salfi	2024	20242	
153	PENGANTAR MANAJEMEN	A	105721101424	NABILA	2024	20242	
154	PENGANTAR MANAJEMEN	A	105721101224	NURFADILLA	2024	20242	
155	PENGANTAR MANAJEMEN	A	105721101024	Selsi	2024	20242	
156	PENGANTAR MANAJEMEN	A	105721102724	Anis Milla Widya Kimm	2024	20242	
157	PENGANTAR MANAJEMEN	A	105721102324	NUR SYAMSIAH HAUM	2024	20242	
158	PENGANTAR MANAJEMEN	A	105721100624	Jusriani	2024	20242	
159	PENGANTAR MANAJEMEN	A	105721100724	Halima	2024	20242	

	A	B	C	F	G	I	J	M
407	BW0612012200	STATISTIK I	B+	105721100824	MOHAMMAD DANIEL IB	2024	20242	
408	BW0612012200	STATISTIK I	A	105721100124	NURHATIMAH	2024	20242	
409	BW0612012200	STATISTIK I	A	105721102224	NURFITRAH HASANAH	2024	20242	
410	BW0612012200	STATISTIK I	A-	105721102024	NUR FADILA	2024	20242	
411	BW0612012200	STATISTIK I	A	105721100924	Muh Nurshahyudi	2024	20242	
412	BW0612012200	STATISTIK I	A-	105721101624	NUR ARHAM DANI SUL	2024	20242	
413	BW0612012200	STATISTIK I	B+	105721102624	MUHAMMAD AMY ARFA	2024	20242	
414	BW0612012200	STATISTIK I	A-	105721102824	Alfa Diva Iman	2024	20242	
415	BW0612012200	STATISTIK I	A	105721101324	SYAHRA RAHMADHAN	2024	20242	
416	BW0612012200	STATISTIK I	A	105721102924	Suci indasari	2024	20242	
417	BW0612012200	STATISTIK I	A	105721103224	Rini Ayu Satri	2024	20242	
418	BW0612012200	STATISTIK I	A	105721101424	NADIA	2024	20242	
419	BW0612012200	STATISTIK I	A	105721101224	NURFADILLA	2024	20242	
420	BW0612012200	STATISTIK I	B+	105721101024	Selvi	2024	20242	
421	BW0612012200	STATISTIK I	A	105721102724	Andi Mita Widya Kirmal	2024	20242	
422	BW0612012200	STATISTIK I	A-	105721102324	NUR SYAMSIAH HAJA	2024	20242	
423	BW0612012200	STATISTIK I	A	105721100024	Jusriani	2024	20242	
424	BW0612012200	STATISTIK I	A	105721100724	Halima	2024	20242	
425	BW0612012200	STATISTIK I	A	105721101724	NUR RESTU IDRISS	2024	20242	
426	BW0612012200	STATISTIK I	A	105721100324	NUR HEDAYAH	2024	20242	
427	BW0612012200	STATISTIK I	A	105721101124	Anisa pusika	2024	20242	
428	BW0612012200	STATISTIK I	A-	105721102524	MUH. SULTAN ANUGR	2024	20242	
429	BW0612012200	STATISTIK I	A	105721102424	VINA	2024	20242	
430	BW0612012200	STATISTIK I	B	105721100424	EKA WAHYU SULISTIA	2024	20242	
431	BW0612012200	STATISTIK I	A-	105721102224	Nur Islamah	2024	20242	
432	BW0612012200	STATISTIK I	A-	105721101524	NURUL FITRAH	2024	20242	
433	BW0612012200	STATISTIK I	B	105721118424	REARI PUTRI LESTARI	2024	20242	
434	BW0612012200	STATISTIK I	E	105721103724	DELA RIANI	2024	20242	
435	BW0612012200	STATISTIK I	B	105721103524	MUHAMMAD ALQADRI	2024	20242	
436	BW0612012200	STATISTIK I	B-	105721104324	Nelis Anggrani	2024	20242	
437	BW0612012200	STATISTIK I	B-	105721104224	Elga Helmesia	2024	20242	
438	BW0612012200	STATISTIK I	A	105721115124	ANU HASMITA DWI PU	2024	20242	

	A	B	C	F	G	I	J	M
669	BW0612012203	BAHASA INGGRIS BISNIS	K	105721100824	MOHAMMAD DANIEL IB	2024	20242	
670	BW0612012203	BAHASA INGGRIS BISNIS	K	105721100124	NURHATIMAH	2024	20242	
671	BW0612012203	BAHASA INGGRIS BISNIS	K	105721102224	NURFITRAH HASANAH	2024	20242	
672	BW0612012203	BAHASA INGGRIS BISNIS	K	105721102024	NUR FADILA	2024	20242	
673	BW0612012203	BAHASA INGGRIS BISNIS	K	105721100924	Muh Nurshahyudi	2024	20242	
674	BW0612012203	BAHASA INGGRIS BISNIS	K	105721101624	NUR ARHAM DANI SUL	2024	20242	
675	BW0612012203	BAHASA INGGRIS BISNIS	K	105721102624	MUHAMMAD AMY ARFA	2024	20242	
676	BW0612012203	BAHASA INGGRIS BISNIS	K	105721102824	Alfa Diva Iman	2024	20242	
677	BW0612012203	BAHASA INGGRIS BISNIS	K	105721101324	SYAHRA RAHMADHAN	2024	20242	
678	BW0612012203	BAHASA INGGRIS BISNIS	K	105721102924	Suci indasari	2024	20242	
679	BW0612012203	BAHASA INGGRIS BISNIS	K	105721103224	Rini Ayu Satri	2024	20242	
680	BW0612012203	BAHASA INGGRIS BISNIS	K	105721101424	NADIA	2024	20242	
681	BW0612012203	BAHASA INGGRIS BISNIS	K	105721101224	NURFADILLA	2024	20242	
682	BW0612012203	BAHASA INGGRIS BISNIS	K	105721101024	Selvi	2024	20242	
683	BW0612012203	BAHASA INGGRIS BISNIS	K	105721102724	Andi Mita Widya Kirmal	2024	20242	
684	BW0612012203	BAHASA INGGRIS BISNIS	K	105721102324	NUR SYAMSIAH HAJA	2024	20242	
685	BW0612012203	BAHASA INGGRIS BISNIS	K	105721100024	Jusriani	2024	20242	
686	BW0612012203	BAHASA INGGRIS BISNIS	K	105721100724	Halima	2024	20242	
687	BW0612012203	BAHASA INGGRIS BISNIS	K	105721101724	NUR RESTU IDRISS	2024	20242	
688	BW0612012203	BAHASA INGGRIS BISNIS	K	105721100324	NUR HEDAYAH	2024	20242	
689	BW0612012203	BAHASA INGGRIS BISNIS	K	105721102524	MUH. SULTAN ANUGR	2024	20242	
690	BW0612012203	BAHASA INGGRIS BISNIS	K	105721102424	VINA	2024	20242	
691	BW0612012203	BAHASA INGGRIS BISNIS	K	105721100224	MUHAMMAD IORAM	2024	20242	
692	BW0612012203	BAHASA INGGRIS BISNIS	K	105721100424	EKA WAHYU SULISTIA	2024	20242	
693	BW0612012203	BAHASA INGGRIS BISNIS	K	105721102224	Nur Islamah	2024	20242	
694	BW0612012203	BAHASA INGGRIS BISNIS	K	105721101524	NURUL FITRAH	2024	20242	
695	BW0612012203	BAHASA INGGRIS BISNIS	K	105721119424	REARI PUTRI LESTARI	2024	20242	
696	BW0612012203	BAHASA INGGRIS BISNIS	K	105721103724	DELA RIANI	2024	20242	
697	BW0612012203	BAHASA INGGRIS BISNIS	K	105721103524	MUHAMMAD ALQADRI	2024	20242	
698	BW0612012203	BAHASA INGGRIS BISNIS	K	105721104324	Nelis Anggrani	2024	20242	
699	BW0612012203	BAHASA INGGRIS BISNIS	K	105721104224	Elga Helmesia	2024	20242	
700	BW0612012203	BAHASA INGGRIS BISNIS	K	105721115124	ANU HASMITA DWI PU	2024	20242	
701	BW0612012203	BAHASA INGGRIS BISNIS	K	105721100724	Halima	2024	20242	
702	BW0612012203	BAHASA INGGRIS BISNIS	K	105721100324	NUR HEDAYAH	2024	20242	
703	BW0612012203	BAHASA INGGRIS BISNIS	K	105721102524	MUH. SULTAN ANUGR	2024	20242	

	A	B	C	F	G	I	J	M
936	BW0612012204	APLIKASI KOMPUTER	K	105721100824	MOHAMMAD DANIEL IB	2024	20242	
937	BW0612012204	APLIKASI KOMPUTER	E	105721100124	NURHATIMAH	2024	20242	
938	BW0612012204	APLIKASI KOMPUTER	K	105721102224	NURFITRAH HASANAH	2024	20242	
939	BW0612012204	APLIKASI KOMPUTER	K	105721102024	NUR FADILA	2024	20242	
940	BW0612012204	APLIKASI KOMPUTER	K	105721100924	Muh Nurshahyudi	2024	20242	
941	BW0612012204	APLIKASI KOMPUTER	E	105721101624	NUR ARHAM DANI SUL	2024	20242	
942	BW0612012204	APLIKASI KOMPUTER	K	105721102624	MUHAMMAD AMY ARFA	2024	20242	
943	BW0612012204	APLIKASI KOMPUTER	K	105721102824	Alfa Diva Iman	2024	20242	
944	BW0612012204	APLIKASI KOMPUTER	K	105721101324	SYAHRA RAHMADHAN	2024	20242	
945	BW0612012204	APLIKASI KOMPUTER	K	105721102924	Suci indasari	2024	20242	
946	BW0612012204	APLIKASI KOMPUTER	K	105721103224	Rini Ayu Satri	2024	20242	
947	BW0612012204	APLIKASI KOMPUTER	K	105721101424	NADIA	2024	20242	
948	BW0612012204	APLIKASI KOMPUTER	K	105721101224	NURFADILLA	2024	20242	
949	BW0612012204	APLIKASI KOMPUTER	K	105721101024	Selvi	2024	20242	
950	BW0612012204	APLIKASI KOMPUTER	K	105721102724	Andi Mita Widya Kirmal	2024	20242	
951	BW0612012204	APLIKASI KOMPUTER	K	105721102324	NUR SYAMSIAH HAJA	2024	20242	
952	BW0612012204	APLIKASI KOMPUTER	K	105721100024	Jusriani	2024	20242	
953	BW0612012204	APLIKASI KOMPUTER	K	105721100724	Halima	2024	20242	
954	BW0612012204	APLIKASI KOMPUTER	K	105721101724	NUR RESTU IDRISS	2024	20242	
955	BW0612012204	APLIKASI KOMPUTER	K	105721100324	NUR HEDAYAH	2024	20242	
956	BW0612012204	APLIKASI KOMPUTER	K	105721102524	MUH. SULTAN ANUGR	2024	20242	
957	BW0612012204	APLIKASI KOMPUTER	K	105721102424	VINA	2024	20242	
958	BW0612012204	APLIKASI KOMPUTER	K	105721100224	MUHAMMAD IORAM	2024	20242	
959	BW0612012204	APLIKASI KOMPUTER	K	105721100424	EKA WAHYU SULISTIA	2024	20242	
960	BW0612012204	APLIKASI KOMPUTER	K	105721102224	Nur Islamah	2024	20242	
961	BW0612012204	APLIKASI KOMPUTER	K	105721101524	NURUL FITRAH	2024	20242	
962	BW0612012204	APLIKASI KOMPUTER	K	105721103224	REARI PUTRI LESTARI	2024	20242	
963	BW0612012204	APLIKASI KOMPUTER	K	105721103724	DELA RIANI	2024	20242	
964	BW0612012204	APLIKASI KOMPUTER	K	105721103524	MUHAMMAD ALQADRI	2024	20242	
965	BW0612012204	APLIKASI KOMPUTER	K	105721104324	Nelis Anggrani	2024	20242	
966	BW0612012204	APLIKASI KOMPUTER	K	105721104224	Elga Helmesia	2024	20242	
967	BW0612012204	APLIKASI KOMPUTER	K	105721115124	ANU HASMITA DWI PU	2024	20242	

A130 BB70405						
A	B	C	F	G	I	M
1286	CW661213205	PENGANTAR AKUNTANSI II	A	105721100924	MUHAMMAD DWAEEL II	2024
1287	CW661213205	PENGANTAR AKUNTANSI II	A	105721100124	MURFATHAH	2024
1288	CW661213205	PENGANTAR AKUNTANSI II	A	105721100224	MURFATHAH HASANAH	2024
1289	CW661213205	PENGANTAR AKUNTANSI II	A	105721100324	MURFATHAH	2024
1290	CW661213205	PENGANTAR AKUNTANSI II	B+	105721100924	Muh Nurwahyudi	2024
1291	CW661213205	PENGANTAR AKUNTANSI II	A	105721101024	MURFATHAH DANI SLAM	2024
1292	CW661213205	PENGANTAR AKUNTANSI II	B-	105721100224	MUHAMMAD AMRI ARA	2024
1293	CW661213205	PENGANTAR AKUNTANSI II	A	105721100324	Adi Dwi Arhamy	2024
1294	CW661213205	PENGANTAR AKUNTANSI II	B-	105721101324	SYAHRA RAHMADHAN	2024
1295	CW661213205	PENGANTAR AKUNTANSI II	A	105721101124	Andra yusita	2024
1296	CW661213205	PENGANTAR AKUNTANSI II	A	105721100924	Suci Indasari	2024
1297	CW661213205	PENGANTAR AKUNTANSI II	A	105721100224	Rini Ayu Saifi	2024
1298	CW661213205	PENGANTAR AKUNTANSI II	A	105721101424	NABILA	2024
1299	CW661213205	PENGANTAR AKUNTANSI II	A	105721101224	MURFADILLA	2024
1300	CW661213205	PENGANTAR AKUNTANSI II	A	105721101024	Salsi	2024
1301	CW661213205	PENGANTAR AKUNTANSI II	A	105721100224	ANDHAMA WIDYA KIRITA	2024
1302	CW661213205	PENGANTAR AKUNTANSI II	A	105721100324	NUR SYAMSIRAH HAJA	2024
1303	CW661213205	PENGANTAR AKUNTANSI II	A	105721100624	Jusiani	2024
1304	CW661213205	PENGANTAR AKUNTANSI II	A	105721100724	Halima	2024
1305	CW661213205	PENGANTAR AKUNTANSI II	A	105721101724	MUR RESTU IDRIIS	2024
1306	CW661213205	PENGANTAR AKUNTANSI II	A	105721100324	MUR HEDYAH	2024
1307	CW661213205	PENGANTAR AKUNTANSI II	A	105721100524	MUH. SULTAN ANUGR	2024
1308	CW661213205	PENGANTAR AKUNTANSI II	C	105721100424	EKA WAHYU SULISTIA	2024
1309	CW661213205	PENGANTAR AKUNTANSI II	B+	105721100224	Nur Elanah	2024
1310	CW661213205	PENGANTAR AKUNTANSI II	A	105721101524	MURRILLO ALIYAH	2024
1311	CW661213205	PENGANTAR AKUNTANSI II	A	105721102324	MUH. ANBAR	2024
1312	CW661213205	PENGANTAR AKUNTANSI II	K	105721103724	DELA RANTI	2024
1313	CW661213205	PENGANTAR AKUNTANSI II	K	105721103524	MUHAMMAD AMRI ARA	2024
1314	CW661213205	PENGANTAR AKUNTANSI II	A	105721100424	MURRILLO ALIYAH	2024
1315	CW661213205	PENGANTAR AKUNTANSI II	K	105721104224	Elga Helmiya	2024
1316	CW661213205	PENGANTAR AKUNTANSI II	K	105721112022	DEWI NUR ASYAH	2022
1317	CW661213205	PENGANTAR AKUNTANSI II	K	105721115324	ANDHAMA WIDYA KIRITA	2024
1318	CW661213205	PENGANTAR AKUNTANSI II	K	105721105124	Aysha Marisa Azah	2024
1319	CW661213205	PENGANTAR AKUNTANSI II	K	105721105124	Rizwan Ibrahim	2024
1320	CW661213205	PENGANTAR AKUNTANSI II	K	105721105124	Rizwan Ibrahim	2024

A130 BB70405						
A	B	C	F	G	I	M
1460	CW6612132205	TEORI EKONOMI MAKRO	A	105721100024	MUHAMMAD DWAEEL II	2024
1461	CW6612132205	TEORI EKONOMI MAKRO	A	105721100124	MURFATHAH	2024
1462	CW6612132205	TEORI EKONOMI MAKRO	A	105721100224	MURFATHAH HASANAH	2024
1463	CW6612132205	TEORI EKONOMI MAKRO	A	105721100324	MURFATHAH	2024
1464	CW6612132205	TEORI EKONOMI MAKRO	A	105721101024	MURFATHAH DANI SLAM	2024
1465	CW6612132205	TEORI EKONOMI MAKRO	A	105721101622	ANDHAMA WIDYA KIRITA	2024
1466	CW6612132205	TEORI EKONOMI MAKRO	E	105721100924	MUHAMMAD AMRI ARA	2024
1467	CW6612132205	TEORI EKONOMI MAKRO	A	105721100324	Adi Dwi Arhamy	2024
1468	CW6612132205	TEORI EKONOMI MAKRO	A	105721101324	SYAHRA RAHMADHAN	2024
1469	CW6612132205	TEORI EKONOMI MAKRO	A	105721101124	Andra yusita	2024
1470	CW6612132205	TEORI EKONOMI MAKRO	A	105721100924	Suci Indasari	2024
1471	CW6612132205	TEORI EKONOMI MAKRO	A	105721100224	Rini Ayu Saifi	2024
1472	CW6612132205	TEORI EKONOMI MAKRO	A	105721101424	NABILA	2024
1473	CW6612132205	TEORI EKONOMI MAKRO	A	105721101224	MURFADILLA	2024
1474	CW6612132205	TEORI EKONOMI MAKRO	A	105721101024	Salsi	2024
1475	CW6612132205	TEORI EKONOMI MAKRO	A	105721100224	ANDHAMA WIDYA KIRITA	2024
1476	CW6612132205	TEORI EKONOMI MAKRO	A	105721100324	NUR SYAMSIRAH HAJA	2024
1477	CW6612132205	TEORI EKONOMI MAKRO	A	105721100624	Jusiani	2024
1478	CW6612132205	TEORI EKONOMI MAKRO	A	105721100724	Halima	2024
1479	CW6612132205	TEORI EKONOMI MAKRO	A	105721101724	MUR RESTU IDRIIS	2024
1480	CW6612132205	TEORI EKONOMI MAKRO	A	105721100324	MUR HEDYAH	2024
1481	CW6612132205	TEORI EKONOMI MAKRO	A	105721100524	MUH. SULTAN ANUGR	2024
1482	CW6612132205	TEORI EKONOMI MAKRO	A	105721100424	WNA	2024
1483	CW6612132205	TEORI EKONOMI MAKRO	C	105721100424	EKA WAHYU SULISTIA	2024
1484	CW6612132205	TEORI EKONOMI MAKRO	A	105721100224	Nur Elanah	2024
1485	CW6612132205	TEORI EKONOMI MAKRO	A	105721101524	MURRILLO ALIYAH	2024
1486	CW6612132205	TEORI EKONOMI MAKRO	B+	105721103724	DELA RANTI	2024
1487	CW6612132205	TEORI EKONOMI MAKRO	A	105721103524	MUHAMMAD AMRI ARA	2024
1488	CW6612132205	TEORI EKONOMI MAKRO	A	105721104224	Elga Helmiya	2024
1489	CW6612132205	TEORI EKONOMI MAKRO	A	105721112022	DEWI NUR ASYAH	2022
1490	CW6612132205	TEORI EKONOMI MAKRO	A	105721115324	ANDHAMA WIDYA KIRITA	2024

A130 BB70405						
A	B	C	F	G	I	M
2017	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721108924	WINDAH NOPYANTI	2024
2018	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721109024	MURILIA BINTARA	2024
2019	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721109124	MURILIA ALIYAH	2024
2020	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721109224	Mudatillah	2024
2021	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721109324	Lita Dwi Fadiah	2024
2022	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721109424	Juliana	2024
2023	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721109524	Nutris Herbit	2024
2024	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721109624	HUSNUL HUSNA	2024
2025	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721109724	Mila Setiawan	2024
2026	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721109824	MITSQ HEDAYATULLA	2024
2027	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721109924	MUA ADIL	2024
2028	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721110024	MUHAMMAD RIDZY FAL	2024
2029	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721110124	Auri	2024
2030	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721110224	A MUH. HERY Z	2024
2031	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721110324	PAUDZAN AZHRA	2024
2032	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	K	105721110424	Syamsul	2024
2033	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721111724	A MUH. FADIL FIRRI	2024
2034	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721112024	MUH. HADI ARHANI	2024
2035	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721112124	MURRILLO ALIYAH	2024
2036	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721112224	SUPRADI	2024
2037	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721112324	MURIL LUTHFAH HAN	2024
2038	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721111124	JUNWARDI RANGI	2024
2039	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721105424	Riswan Ibrahim	2024
2040	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721127124	ITTRI HUSRI	2024
2041	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721110424	MUR INDAH SARLEBIM	2024
2042	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721109024	KHOIRIL US	2024
2043	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721111824	FADILA RINADIAN PU	2024
2044	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721114524	Nani Isyrah	2024
2045	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721101324	SYAHRA RAHMADHAN	2024
2046	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721102424	ZULFANNY	2024
2047	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721112724	IRKA KARTIKA	2024
2048	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721110324	FADHILA	2024
2049	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	A	105721110624	PURTI KUSNADY SALSAD	2024
2050	AW0910041205	AL ISLAM DAN KEBUHAIRADYAHAN	K	105721111424	MURIL HURAH FEBRI	2024

	A	B	C	D	E	F	G	H	I	J	K	L	M	N
1														
2	Pendidikan Agama Islam	2												
3	Bahasa Inggris	2												
4	Bahasa Indonesia	2												
5	Pendidikan Pancasila	2												
6	Pendidikan Kewarganegaraan	2												
7	Ilmu Sosial dan Budaya Dasar	2												
8	Pengantar Bisnis	3												
9	Pengantar Ekonomi	3												
10	Pengantar Akuntansi I	3												
11	Matematika Ekonomi	3												
12	AIK II	1												
13	Pengantar Manajemen	3												
14	Statistik I	2												
15	Bahasa Inggris Bisnis	2												
16	Aplikasi Komputer	2												
17	Pengantar Akuntansi II	2												
18	Etika Bisnis Syariah	2												
19	Manajemen Koperasi Dan UKM	3												
20	Komunikasi Bisnis	3												
21	Teori Ekonomi Mikro	2												
22	Al-Islam Muhammadiyah (AIK III)	1												
23	Sistem Ekonomi Islam	3												
24	Manajemen SDM I	3												
25	Manajemen Operasional I	3												
26	Manajemen Keuangan I	3												
27	Manajemen Pemasaran I	3												
28	Sistem Informasi Manajemen	3												
29	Statistik II	2												

	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O
30	Kewirausahaan I	2													
31	AIK IV	1													
32	Lembaga Keuangan Syariah	3													
33	Kewirausahaan II	2													
34	Manajemen Operasional II	3													
35	Manajemen Keuangan II	3													
36	Manajemen Pemasaran II	3													
37	Manajemen SDM II	3													
38	Riset Operasional	3													
39	Teknik Proyek Bisnis	3													
40	Al-Islam & Muhammadiyah V	1													
41	Perencanaan Bisnis	3													
42	Ekonomi Internasional	3													
43	Teori Pengantar Koperasi	3													
44	Ekonomi Manajerial	3													
45	Pengantar Perusahaan	3													
46	Akuntansi Manajemen	3													
47	Studi Kelayakan Bisnis	3													
48	Al-Islam Muhammadiyah VI	1													
49	Metodologi Penelitian	3													
50	Manajemen Strategi	3													
51	Strategi Bisnis Manajemen Pemasaran	4													
52	Strategi Pemasaran	3													
53	Manajemen Pemasaran Jasa	3													
54	Perilaku Konsumen	3													
55	Riset Dan Analisis Pasar	3													
56	Internet Marketing	3													
57	Bidang Ilmu Manajemen Keuangan	4													
58	AIK VIII (Komprehensif AIK)	2													

	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O
59	AIK VIII (Komprehensif AIK)	2													
60	Kalimat Kerja Profesi (Majalah Pita)	4													
61	Manajemen Pemasaran Internasional	3													
62	Manajemen Pemasaran	3													
63	Manajemen Logistik	3													
64	Manajemen Saluran Distribusi Dan E-commerce	3													
65	Manajemen Risiko	3													
66	Manajemen Keuangan	3													
67	Pasar Modal	3													
68	Manajemen Keuangan Internasional	3													
69	Bidang Ilmu Manajemen SDM	4													
70	Manajemen Karyawan	3													
71	Manajemen Manajemen SDM	3													
72	Pemecahan Pengembangan Karir	3													
73	Manajemen SDM Internasional	3													
74	Strategi Kompetensi & Koneksi	6													
75															
76															
77															
78															
79															
80															
81															
82															
83															
84															
85															
86															

Lampiran 2 Source Code Sistem prediksi peminatan mahasiswa

```
# SISTEM PREDIKSI PEMINATAN MAHASISWA

# Universitas Muhammadiyah Makassar

# Fakultas Ekonomi dan Bisnis - Program Studi Manajemen

#
=====

# Install dependencies

!pip install streamlit pyngrok tqdm scikit-learn matplotlib seaborn pandas numpy

import pandas as pd
import numpy as np
import matplotlib.pyplot as plt
import seaborn as sns
from sklearn.cluster import KMeans
from sklearn.preprocessing import StandardScaler, MinMaxScaler, RobustScaler
from sklearn.metrics import silhouette_score, davies_bouldin_score, accuracy_score,
classification_report, confusion_matrix, silhouette_samples
from sklearn.decomposition import PCA
from sklearn.feature_selection import SelectKBest, f_classif
from tqdm import tqdm
import pickle
import warnings

from scipy import stats
warnings.filterwarnings('ignore')

# Set style untuk visualisasi
plt.style.use('seaborn-v0_8')
sns.set_palette("husl")
```



```

print("🎓 SISTEM PREDIKSI PEMINATAN MAHASISWA")

print("🏫 Universitas Muhammadiyah Makassar")

print("=" * 80)

# ENHANCED DATA GENERATION WITH CLEARER SPECIALIZATION PATTERNS

#
=====

# Mata Kuliah dengan pengelompokan yang lebih jelas
mata_kuliah_semester_1_5 = [
    # Semester 1 - Foundation
    'PENDIDIKAN AGAMA ISLAM', 'PANCASILA', 'BAHASA INDONESIA',
    'BAHASA INGGRIS I',
    'MATEMATIKA EKONOMI', 'PENGANTAR EKONOMI MIKRO',
    'PENGANTAR BISNIS',
    # Semester 2 - Core Business
    'PENGANTAR EKONOMI MAKRO', 'PENGANTAR MANAJEMEN', 'BAHASA
    INGGRIS II',
    'STATISTIKA I', 'PENGANTAR AKUNTANSI', 'ETIKA BISNIS',
    # Semester 3 - Intermediate
    'STATISTIKA II', 'AKUNTANSI MANAJEMEN', 'MANAJEMEN
    OPERASIONAL',
    'KOMUNIKASI BISNIS', 'HUKUM BISNIS', 'KEWIRAUSAHAAN',
    # Semester 4 - Applied
    'METODOLOGI PENELITIAN', 'SISTEM INFORMASI MANAJEMEN',
    'PERILAKU ORGANISASI',
    'MANAJEMEN STRATEGIS', 'EKONOMI INTERNASIONAL',

```

```

# Semester 5 - Advanced

'BUSINESS INTELLIGENCE', 'MANAJEMEN PROYEK', 'ANALISIS
LAPORAN KEUANGAN',

'MANAJEMEN RISIKO', 'STUDI KELAYAKAN BISNIS'
]

# Pengelompokan mata kuliah berdasarkan relevansi peminatan
mata_kuliah_keuangan_oriented = [

'MATEMATIKA EKONOMI', 'STATISTIKA I', 'STATISTIKA II', 'PENGANTAR
AKUNTANSI',

'AKUNTANSI MANAJEMEN', 'ANALISIS LAPORAN KEUANGAN',
'MANAJEMEN RISIKO',

'STUDI KELAYAKAN BISNIS', 'EKONOMI INTERNASIONAL'
]

mata_kuliah_pemasaran_oriented = [

'PENGANTAR BISNIS', 'KOMUNIKASI BISNIS', 'KEWIRAUSAHAAN',
'MANAJEMEN STRATEGIS',

'BUSINESS INTELLIGENCE', 'BAHASA INGGRIS I', 'BAHASA INGGRIS II'
]

mata_kuliah_sdm_oriented = [

'PENGANTAR MANAJEMEN', 'PERILAKU ORGANISASI', 'SISTEM
INFORMASI MANAJEMEN',

'MANAJEMEN OPERASIONAL', 'MANAJEMEN PROYEK', 'METODOLOGI
PENELITIAN',

'ETIKA BISNIS'
]

print(f'D Total mata kuliah: {len(mata_kuliah_semester_1_5)}")

print(f'🏢 Keuangan oriented: {len(mata_kuliah_keuangan_oriented)}")

print(f'📊 Pemasaran oriented: {len(mata_kuliah_pemasaran_oriented)}")

```

```

print(f'👥 SDM oriented: {len(mata_kuliah_sdm_oriented)}')

# IMPROVED DATA GENERATION WITH STRONGER PATTERNS

#
=====

def generate_optimized_student_data(n_students_per_specialization=100):
    """Generate data dengan pola yang lebih jelas untuk setiap peminatan"""
    np.random.seed(42)
    all_students_data = []
    specializations = ['MANAJEMEN KEUANGAN', 'MANAJEMEN PEMASARAN', 'MANAJEMEN SDM']
    for spec_idx, specialization in enumerate(specializations):
        print(f'Generating {n_students_per_specialization} students for {specialization}...')
        for i in range(n_students_per_specialization):
            student = {}
            # Base performance dengan variasi yang wajar
            base_performance = np.random.normal(77, 6) # Mean 77, std 6
            # Generate nilai untuk setiap mata kuliah
            for course in mata_kuliah_semester_1_5:
                # Start with base performance
                course_score = base_performance + np.random.normal(0, 3)
                # Add specialization bias yang lebih kuat
                if specialization == 'MANAJEMEN KEUANGAN':
                    if course in mata_kuliah_keuangan_oriented:
                        # Strong boost for finance-related courses
                        course_score += np.random.normal(12, 4) # Bigger boost

```

```

elif course in mata_kuliah_pemasaran_oriented:

    # Penalty for marketing courses

    course_score += np.random.normal(-6, 3)

elif course in mata_kuliah_sdm_oriented:

    # Slight penalty for HR courses

    course_score += np.random.normal(-3, 2)

elif specialization == 'MANAJEMEN PEMASARAN':

    if course in mata_kuliah_pemasaran_oriented:

        # Strong boost for marketing-related courses

        course_score += np.random.normal(10, 3)

    elif course in mata_kuliah_keuangan_oriented:

        # Penalty for finance courses

        course_score += np.random.normal(-5, 3)

    elif course in mata_kuliah_sdm_oriented:

        # Small boost for HR courses (communication skills)

        course_score += np.random.normal(2, 2)

elif specialization == 'MANAJEMEN SDM':

    if course in mata_kuliah_sdm_oriented:

        # Strong boost for HR-related courses

        course_score += np.random.normal(11, 3)

    elif course in mata_kuliah_keuangan_oriented:

        # Penalty for finance courses

        course_score += np.random.normal(-4, 3)

elif course in mata_kuliah_pemasaran_oriented:

    # Small boost for marketing courses (communication)

    course_score += np.random.normal(3, 2)

```



```

# Ensure realistic score range

course_score = max(60, min(95, course_score))

student[course] = course_score

# Student metadata

student['NIM'] = f'2022{spec_idx+1}{str(i+1).zfill(3)}'
student['NAMA'] = f'{specialization.split()[-1]}_Student_{i+1:03d}'
student['PEMINATAN_AKTUAL'] = specialization

all_students_data.append(student)

return pd.DataFrame(all_students_data)

# Generate optimized data

print("\n🔧 Generating optimized student data with clearer patterns...")
df_training = generate_optimized_student_data(n_students_per_specialization=120)

print(f"✅ Training data: {len(df_training)} mahasiswa")

print("\nD Distribusi peminatan:")
print(df_training['PEMINATAN_AKTUAL'].value_counts())

# FEATURE ENGINEERING AND SELECTION

#
=====

print("\n🔧 Feature engineering and selection...")

# Basic features

feature_columns = mata_kuliah_semester_1_5

X_basic = df_training[feature_columns].copy()

# Create engineered features berdasarkan pengelompokan mata kuliah

X_engineered = pd.DataFrame()

```

1. Average scores per specialization area

```
X_engineered['FINANCE_SCORE'] =  
df_training[mata_kuliah_keuangan_oriented].mean(axis=1)
```

```
X_engineered['MARKETING_SCORE'] =  
df_training[mata_kuliah_pemasaran_oriented].mean(axis=1)
```

```
X_engineered['HRM_SCORE'] =  
df_training[mata_kuliah_sdm_oriented].mean(axis=1)
```

2. Specialization preferences (differences)

```
X_engineered['FINANCE_vs_MARKETING'] =  
X_engineered['FINANCE_SCORE'] - X_engineered['MARKETING_SCORE']
```

```
X_engineered['FINANCE_vs_HRM'] = X_engineered['FINANCE_SCORE'] -  
X_engineered['HRM_SCORE']
```

```
X_engineered['MARKETING_vs_HRM'] = X_engineered['MARKETING_SCORE']  
- X_engineered['HRM_SCORE']
```

3. Performance consistency (standard deviation)

```
X_engineered['PERFORMANCE_CONSISTENCY'] =  
df_training[feature_columns].std(axis=1)
```

4. Overall performance

```
X_engineered['OVERALL_PERFORMANCE'] =  
df_training[feature_columns].mean(axis=1)
```

5. Quantitative vs Qualitative performance

```
quantitative_courses = [c for c in feature_columns if any(word in c for word in  
['MATEMATIKA', 'STATISTIK', 'AKUNTANSI'])]
```

```
qualitative_courses = [c for c in feature_columns if any(word in c for word in  
['KOMUNIKASI', 'MANAJEMEN', 'BISNIS', 'BAHASA'])]
```

if quantitative_courses:

```
    X_engineered['QUANTITATIVE_PERFORMANCE'] =  
    df_training[quantitative_courses].mean(axis=1)
```

else:

```

X_engineered['QUANTITATIVE_PERFORMANCE'] =
X_engineered['OVERALL_PERFORMANCE']

if qualitative_courses:

    X_engineered['QUALITATIVE_PERFORMANCE'] =
df_training[qualitative_courses].mean(axis=1)

else:

    X_engineered['QUALITATIVE_PERFORMANCE'] =
X_engineered['OVERALL_PERFORMANCE']

X_engineered['QUANT_vs_QUAL'] =
X_engineered['QUANTITATIVE_PERFORMANCE'] -
X_engineered['QUALITATIVE_PERFORMANCE']

print(f'D Engineered features: {len(X_engineered.columns)}")
print("Features created:")
for i, col in enumerate(X_engineered.columns, 1):
    print(f' {i:2d}. {col}")

# MULTIPLE SCALING APPROACHES COMPARISON
#
=====

print("\nG Testing different scaling approaches...")
# Test different scalers
scalers = {
    'StandardScaler': StandardScaler(),
    'MinMaxScaler': MinMaxScaler(),
    'RobustScaler': RobustScaler()
}

scaler_results = {}

for scaler_name, scaler in scalers.items():

```

```

X_scaled = scaler.fit_transform(X_engineered)

# Test k=3 clustering

kmeans = KMeans(n_clusters=3, random_state=42, n_init=20, max_iter=300)

labels = kmeans.fit_predict(X_scaled)

silhouette_avg = silhouette_score(X_scaled, labels)

davies_bouldin = davies_bouldin_score(X_scaled, labels)

scaler_results[scaler_name] = {
    'silhouette': silhouette_avg,
    'davies_bouldin': davies_bouldin,
    'labels': labels,
    'scaler': scaler,
    'scaled_data': X_scaled
}

print(f"      {scaler_name:15}:      Silhouette={silhouette_avg:.3f},      Davies-
Bouldin={davies_bouldin:.3f}")

# Choose best scaler

best_scaler_name = max(scaler_results.keys(), key=lambda k:
scaler_results[k]['silhouette'])

best_scaler_info = scaler_results[best_scaler_name]

print(f"\n 🏆 Best scaler: {best_scaler_name}")

print(f"  Silhouette Score: {best_scaler_info['silhouette']:.3f}")

print(f"  Davies-Bouldin Score: {best_scaler_info['davies_bouldin']:.3f}")

# OPTIMIZED K=3 CLUSTERING

#
=====
=====

```



```

print(f"\n🎯 Building optimized K=3 model with {best_scaler_name}...")

# Use best scaler and data

final_scaler = best_scaler_info['scaler']
X_final_scaled = best_scaler_info['scaled_data']

# Optimized K-Means parameters
final_model = KMeans(
    n_clusters=3,
    random_state=42,
    n_init=50, # More initializations
    max_iter=500, # More iterations
    tol=1e-6, # Tighter tolerance
    algorithm='lloyd' # Full EM algorithm
)
cluster_labels = final_model.fit_predict(X_final_scaled)

# Calculate final metrics
final_silhouette = silhouette_score(X_final_scaled, cluster_labels)
final_davies_bouldin = davies_bouldin_score(X_final_scaled, cluster_labels)
final_inertia = final_model.inertia_

print(f"🎯 Final Model Performance:")
print(f"  Silhouette Score: {final_silhouette:.3f}")
print(f"  Davies-Bouldin Score: {final_davies_bouldin:.3f}")
print(f"  Inertia: {final_inertia:.2f}")

# CLUSTER MAPPING AND VALIDATION

#
=====
=====

```

```

print("\n 📖 Mapping clusters to specializations...") #

Add cluster labels to dataframe
df_training['CLUSTER'] = cluster_labels

# Create detailed cluster mapping
cluster_mapping = {}
cluster_details = {}
for cluster_id in range(3):
    cluster_data = df_training[df_training['CLUSTER'] == cluster_id]
    # Find dominant specialization
    peminatan_counts = cluster_data['PEMINATAN_AKTUAL'].value_counts()
    dominant_peminatan = peminatan_counts.index[0]
    purity = peminatan_counts.iloc[0] / len(cluster_data) * 100
    cluster_mapping[cluster_id] = dominant_peminatan
    cluster_details[cluster_id] = {
        'dominant': dominant_peminatan,
        'counts': peminatan_counts.to_dict(),
        'purity': purity,
        'n_students': len(cluster_data),
        'avg_performance': cluster_data[feature_columns].mean().mean()
    }

print(f"\n 🎯 Cluster {cluster_id} → {dominant_peminatan}")

print(f"
                Size: {len(cluster_data)}
students
({len(cluster_data)/len(df_training)*100:.1f}%)")

print(f" Purity: {purity:.1f}%)")

print(f" Distribution: {dict(peminatan_counts)}")

```

```

# Calculate clustering accuracy

df_training['PEMINATAN_PREDIKSI'] = [cluster_mapping[c] for c in cluster_labels]

clustering_accuracy = accuracy_score(df_training['PEMINATAN_AKTUAL'],
df_training['PEMINATAN_PREDIKSI'])

print(f"\n 🎯 Clustering Accuracy: {clustering_accuracy:.3f}
({clustering_accuracy*100:.1f}%)")

# COMPREHENSIVE VISUALIZATION

#
=====

print("\nD Creating comprehensive visualizations...")

# Create comprehensive plots
fig = plt.figure(figsize=(20, 16))

# 1. Cluster quality metrics comparison
ax1 = plt.subplot(3, 4, 1)
scaler_names = list(scaler_results.keys())
silhouette_scores = [scaler_results[name]['silhouette'] for name in scaler_names]
colors = ['red', 'blue', 'green']
bars = ax1.bar(scaler_names, silhouette_scores, color=colors, alpha=0.7)
ax1.set_title(' 🔍 Silhouette Scores by Scaler', fontweight='bold')
ax1.set_ylabel('Silhouette Score')
ax1.grid(True, alpha=0.3)
ax1.set_ylim(0, max(silhouette_scores) * 1.2)

# Highlight best scaler
best_idx = scaler_names.index(best_scaler_name)
bars[best_idx].set_color('gold')
bars[best_idx].set_edgecolor('black')

```

```

bars[best_idx].set_linewidth(2)

for bar, score in zip(bars, silhouette_scores):

    ax1.text(bar.get_x() + bar.get_width()/2, bar.get_height() + 0.01,
             f'{score:.3f}', ha='center', va='bottom', fontweight='bold')

# 2. Cluster size distribution
ax2 = plt.subplot(3, 4, 2)

cluster_sizes = [cluster_details[i]['n_students'] for i in range(3)]

cluster_names = [f'Cluster {i}\n({cluster_mapping[i].split()[-1]})' for i in range(3)]

cluster_colors = ['lightcoral', 'lightblue', 'lightgreen']

bars = ax2.bar(cluster_names, cluster_sizes, color=cluster_colors, alpha=0.8)

ax2.set_title('D Cluster Size Distribution', fontweight='bold')

ax2.set_ylabel('Number of Students')

ax2.grid(True, alpha=0.3)

for bar, size in zip(bars, cluster_sizes):

    ax2.text(bar.get_x() + bar.get_width()/2, bar.get_height() + 1,
             str(size), ha='center', va='bottom', fontweight='bold')

# 3. Cluster purity
ax3 = plt.subplot(3, 4, 3)

purities = [cluster_details[i]['purity'] for i in range(3)]

bars = ax3.bar(cluster_names, purities, color=cluster_colors, alpha=0.8)

ax3.set_title('🎯 Cluster Purity', fontweight='bold')

ax3.set_ylabel('Purity (%)')

ax3.set_ylim(0, 100)

ax3.grid(True, alpha=0.3)

for bar, purity in zip(bars, purities):

    ax3.text(bar.get_x() + bar.get_width()/2, bar.get_height() + 1,

```



```

        f'{purity:.1f}%', ha='center', va='bottom', fontweight='bold')

# 4. Performance comparison
ax4 = plt.subplot(3, 4, 4)
performances = [cluster_details[i]['avg_performance'] for i in range(3)]
bars = ax4.bar(cluster_names, performances, color=cluster_colors, alpha=0.8)
ax4.set_title('📊 Average Performance by Cluster', fontweight='bold')
ax4.set_ylabel('Average Score')
ax4.grid(True, alpha=0.3)
for bar, perf in zip(bars, performances):
    ax4.text(bar.get_x() + bar.get_width()/2, bar.get_height() + 0.2,
             f'{perf:.1f}', ha='center', va='bottom', fontweight='bold')

# 5. PCA Visualization
ax5 = plt.subplot(3, 4, 5)
pca = PCA(n_components=2)
X_pca = pca.fit_transform(X_final_scaled)
for i, spec in enumerate(['KEUANGAN', 'PEMASARAN', 'SDM']):
    mask = df_training['PEMINATAN_AKTUAL'].str.contains(spec)
    ax5.scatter(X_pca[mask, 0], X_pca[mask, 1],
               c=cluster_colors[i], label=spec, alpha=0.7, s=50)
ax5.set_title('🎨 PCA: Actual Specializations', fontweight='bold')
ax5.set_xlabel(f'PC1 ({pca.explained_variance_ratio_[0]*100:.1f}% variance)')
ax5.set_ylabel(f'PC2 ({pca.explained_variance_ratio_[1]*100:.1f}% variance)')
ax5.legend()
ax5.grid(True, alpha=0.3)

# 6. PCA with clusters
ax6 = plt.subplot(3, 4, 6)

```

```

for cluster_id in range(3):
    mask = cluster_labels == cluster_id
    ax6.scatter(X_pca[mask, 0], X_pca[mask, 1],
                c=cluster_colors[cluster_id], label=f'Cluster {cluster_id}',
                alpha=0.7, s=50)
ax6.set_title('🖨️ PCA: Predicted Clusters', fontweight='bold')
ax6.set_xlabel(f'PC1 ({pca.explained_variance_ratio_[0]*100:.1f}% variance)')
ax6.set_ylabel(f'PC2 ({pca.explained_variance_ratio_[1]*100:.1f}% variance)')
ax6.legend()
ax6.grid(True, alpha=0.3)
# 7. Feature importance heatmap
ax7 = plt.subplot(3, 4, (7, 8))
cluster_centers_scaled = final_model.cluster_centers_
# Create feature importance matrix
importance_matrix = np.abs(cluster_centers_scaled)
feature_names = [col.replace('_', ' ') for col in X_engineered.columns]
im = ax7.imshow(importance_matrix, cmap='RdYlBu_r', aspect='auto')
ax7.set_xticks(range(len(feature_names)))
ax7.set_xticklabels(feature_names, rotation=45, ha='right', fontsize=8)
ax7.set_yticks(range(3))
ax7.set_yticklabels([f'Cluster {i}' for i in range(3)])
ax7.set_title('🔍 Feature Importance by Cluster', fontweight='bold')
# Add colorbar
cbar = plt.colorbar(im, ax=ax7, shrink=0.6)
cbar.set_label('Scaled Importance')

```

```

# 8. Silhouette analysis
ax8 = plt.subplot(3, 4, 9)
silhouette_sample_values = silhouette_samples(X_final_scaled, cluster_labels)
y_lower = 10
for i in range(3):
    ith_cluster_silhouette_values = silhouette_sample_values[cluster_labels == i]
    ith_cluster_silhouette_values.sort()
    size_cluster_i = ith_cluster_silhouette_values.shape[0]
    y_upper = y_lower + size_cluster_i
    ax8.fill_between(np.arange(y_lower, y_upper), 0, ith_cluster_silhouette_values,
                     facecolor=cluster_colors[i], alpha=0.7)
    ax8.text(-0.05, y_lower + 0.5 * size_cluster_i, str(i))
    y_lower = y_upper + 10
ax8.axvline(x=final_silhouette, color="red", linestyle="--",
            label=f'Average: {final_silhouette:.3f}')
ax8.set_xlabel('Silhouette Coefficient')
ax8.set_ylabel('Cluster Label')
ax8.set_title('Silhouette Analysis', fontweight='bold')
ax8.legend()
# 9. Confusion Matrix
ax9 = plt.subplot(3, 4, 10)

conf_matrix = confusion_matrix(df_training['PEMINATAN_AKTUAL'],
                               df_training['PEMINATAN_PREDIKSI'])

spec_labels = ['KEUANGAN', 'PEMASARAN', 'SDM']
sns.heatmap(conf_matrix, annot=True, fmt='d', cmap='Blues',
            xticklabels=spec_labels, yticklabels=spec_labels,

```

```

ax=ax9, cbar_kws={'shrink': 0.8})

ax9.set_title(' 📊 Confusion Matrix', fontweight='bold')

ax9.set_xlabel('Predicted')

ax9.set_ylabel('Actual')

# 10. Classification Report (text)

ax10 = plt.subplot(3, 4, 11)

ax10.axis('off')

report = classification_report(df_training['PEMINATAN_AKTUAL'],
df_training['PEMINATAN_PREDIKSI'])

ax10.text(0.1, 0.5, report, fontsize=8, fontfamily='monospace',
verticalalignment='center', transform=ax10.transAxes)

ax10.set_title('D Classification Report', fontweight='bold')

# 11. Model summary

ax11 = plt.subplot(3, 4, 12)

ax11.axis('off')

summary_text = f"""
🎯 MODEL SUMMARY

✅ Silhouette Score: {final_silhouette:.3f}

✅ Davies-Bouldin: {final_davies_bouldin:.3f}

✅ Clustering Accuracy: {clustering_accuracy:.1f}%

✅ Best Scaler: {best_scaler_name}

D CLUSTER DETAILS:

Cluster          0:          {cluster_details[0]['dominant'].split()[-1]}
({cluster_details[0]['purity']:.1f}% pure)

Cluster          1:          {cluster_details[1]['dominant'].split()[-1]}
({cluster_details[1]['purity']:.1f}% pure)

```



```
Cluster                2:                {cluster_details[2]['dominant'].split()[-1]}
({cluster_details[2]['purity']:.1f}% pure)
```

🏆 IMPROVEMENTS:

- Enhanced feature engineering
- Optimal scaler selection
- Stronger specialization patterns
- Better cluster separation

```
"""
```

```
ax11.text(0.05, 0.95, summary_text, fontsize=10, verticalalignment='top',
```

```
transform=ax11.transAxes, fontfamily='monospace')
```

```
plt.suptitle('🎯 OPTIMIZED K=3 CLUSTERING ANALYSIS', fontsize=16,
fontweight='bold', y=0.95)
```

```
plt.tight_layout()
```

```
plt.subplots_adjust(top=0.92)
```

```
plt.show()
```

```
# DETAILED SILHOUETTE ANALYSIS
```

```
#
```

```
=====
```

```
print(f"\n🔍 DETAILED SILHOUETTE ANALYSIS")
```

```
print("="*50)
```

```
print(f'Overall Silhouette Score: {final_silhouette:.3f}')
```

```
print(f'Target: >0.5 (Good), >0.7 (Excellent)')
```

```
for i in range(3):
```

```
    cluster_silhouettes = silhouette_sample_values[cluster_labels == i]
```

```
    print(f'Cluster {i} ({cluster_mapping[i].split()[-1]}):')
```

```
    print(f'  Mean: {cluster_silhouettes.mean():.3f}')
```

```

print(f' Std: {cluster_silhouettes.std():.3f}')
print(f' Min: {cluster_silhouettes.min():.3f}')
print(f' Max: {cluster_silhouettes.max():.3f}')
# GENERATE TEST DATA FOR VALIDATION
#
=====
=====

print(f'\nG Generating test data for validation...')

def generate_test_students(n_test=60):
    """Generate test students (20 per specialization)"""
    np.random.seed(123) # Different seed for test
    test_students = []
    specializations = ['MANAJEMEN KEUANGAN', 'MANAJEMEN
PEMASARAN', 'MANAJEMEN SDM']
    for spec_idx, specialization in enumerate(specializations):
        for i in range(20): # 20 per specialization
            student = {}
            # Similar pattern as training but with some variation
            base_performance = np.random.normal(76, 7)
            for course in mata_kuliah_semester_1_5:
                course_score = base_performance + np.random.normal(0, 4)
            # Apply same bias patterns but slightly weaker
            if specialization == 'MANAJEMEN KEUANGAN':
                if course in mata_kuliah_keuangan_oriented:
                    course_score += np.random.normal(10, 4)
                elif course in mata_kuliah_pemasaran_oriented:

```

```

        course_score += np.random.normal(-5, 3)
    elif course in mata_kuliah_sdm_oriented:
        course_score += np.random.normal(-2, 2)
    elif specialization == 'MANAJEMEN PEMASARAN':
        if course in mata_kuliah_pemasaran_oriented:
            course_score += np.random.normal(9, 3)
        elif course in mata_kuliah_keuangan_oriented:
            course_score += np.random.normal(-4, 3)
        elif course in mata_kuliah_sdm_oriented:
            course_score += np.random.normal(3, 2)
    elif specialization == 'MANAJEMEN SDM':
        if course in mata_kuliah_sdm_oriented:
            course_score += np.random.normal(10, 3)
        elif course in mata_kuliah_keuangan_oriented:
            course_score += np.random.normal(-3, 2)
        elif course in mata_kuliah_pemasaran_oriented:
            course_score += np.random.normal(4, 2)
    course_score = max(60, min(95, course_score))
    student[course] = course_score
    student['NIM'] = f"2023 {spec_idx+1} {str(i+1).zfill(2)}"
    student['NAMA'] = f"Test_{specialization.split()[-1]}_{i+1:02d}"
    student['PEMINATAN_AKTUAL'] = specialization
    test_students.append(student)

return pd.DataFrame(test_students)

# Generate test data
df_test = generate_test_students()

```

```

print(f'✅ Test data: {len(df_test)} mahasiswa")

print("D Test distribution:")

print(df_test['PEMINATAN_AKTUAL'].value_counts())

# PREDICT ON TEST DATA

#
=====

print("\n 🎯 Testing model on new data...")

# Create engineered features for test data
X_test_engineered = pd.DataFrame()
X_test_engineered['FINANCE_SCORE'] =
df_test[mata_kuliah_keuangan_oriented].mean(axis=1)
X_test_engineered['MARKETING_SCORE'] =
df_test[mata_kuliah_pemasaran_oriented].mean(axis=1)
X_test_engineered['HRM_SCORE'] =
df_test[mata_kuliah_sdm_oriented].mean(axis=1)
X_test_engineered['FINANCE_vs_MARKETING'] =
X_test_engineered['FINANCE_SCORE'] -
X_test_engineered['MARKETING_SCORE']
X_test_engineered['FINANCE_vs_HRM'] =
X_test_engineered['FINANCE_SCORE'] - X_test_engineered['HRM_SCORE']
X_test_engineered['MARKETING_vs_HRM'] =
X_test_engineered['MARKETING_SCORE'] - X_test_engineered['HRM_SCORE']
X_test_engineered['PERFORMANCE_CONSISTENCY'] =
df_test[feature_columns].std(axis=1)
X_test_engineered['OVERALL_PERFORMANCE'] =
df_test[feature_columns].mean(axis=1)

if quantitative_courses:
    X_test_engineered['QUANTITATIVE_PERFORMANCE'] =
df_test[quantitative_courses].mean(axis=1)

```



```

else:
    X_test_engineered['QUANTITATIVE_PERFORMANCE'] =
    X_test_engineered['OVERALL_PERFORMANCE']
    if qualitative_courses:
        X_test_engineered['QUALITATIVE_PERFORMANCE'] =
        df_test[qualitative_courses].mean(axis=1)
    else:
        X_test_engineered['QUALITATIVE_PERFORMANCE'] =
        X_test_engineered['OVERALL_PERFORMANCE']
        X_test_engineered['QUANT_vs_QUAL'] =
        X_test_engineered['QUANTITATIVE_PERFORMANCE'] -
        X_test_engineered['QUALITATIVE_PERFORMANCE']
    # Scale test data using trained scaler
    X_test_scaled = final_scaler.transform(X_test_engineered)
    # Predict clusters
    test_clusters = final_model.predict(X_test_scaled)
    test_predictions = [cluster_mapping[c] for c in test_clusters]
    # Calculate test accuracy
    test_accuracy = accuracy_score(df_test['PEMINATAN_AKTUAL'], test_predictions)
    print(f"🎯 Test Accuracy: {test_accuracy:.3f} ({test_accuracy*100:.1f}%)")
    # Detailed test results
    df_test['CLUSTER_PREDIKSI'] = test_clusters
    df_test['PEMINATAN_PREDIKSI'] = test_predictions
    df_test['CORRECT'] = df_test['PEMINATAN_AKTUAL'] ==
    df_test['PEMINATAN_PREDIKSI']

    print("\nD Test Classification Report:")
    print(classification_report(df_test['PEMINATAN_AKTUAL'], test_predictions))

```

```
# CONFIDENCE SCORING
```

```
#
```

```
=====
```

```
print("\n🌟 Calculating prediction confidence...")
```

```
def calculate_confidence_scores(X_scaled, model, labels):
```

```
    """Calculate confidence based on distance to cluster centroid"""
```

```
    confidences = []
```

```
    for i, sample in enumerate(X_scaled):
```

```
        cluster_id = labels[i]
```

```
        centroid = model.cluster_centers_[cluster_id]
```

```
        # Distance to assigned cluster centroid
```

```
        dist_to_assigned = np.linalg.norm(sample - centroid)
```

```
        # Distance to all cluster centroids
```

```
        all_distances = [np.linalg.norm(sample - center) for center in  
model.cluster_centers_]
```

```
        # Confidence based on relative distance
```

```
        min_distance = min(all_distances)
```

```
        second_min_distance = sorted(all_distances)[1]
```

```
        # Confidence score (higher when assigned cluster is clearly closest)
```

```
        if second_min_distance > 0:
```

```
            confidence = 1 - (min_distance / second_min_distance)
```

```
            confidence = max(0, min(1, confidence))
```

```
        else:
```

```
            confidence = 1.0
```

```
        confidences.append(confidence)
```

```
    return np.array(confidences)
```

```

# Calculate confidence for test data

test_confidences = calculate_confidence_scores(X_test_scaled, final_model,
test_clusters)

df_test['CONFIDENCE'] = test_confidences

print(f'Average confidence: {test_confidences.mean():.3f}')

print("\nConfidence distribution:")

print(f"  High (>0.8): {sum(test_confidences > 0.8)} ({sum(test_confidences >
0.8)/len(test_confidences)*100:.1f}%)")

print(f" Medium (0.6-0.8): {sum((test_confidences >= 0.6) & (test_confidences <=
0.8))} ({sum((test_confidences >= 0.6) & (test_confidences <=
0.8))/len(test_confidences)*100:.1f}%)")

print(f" Low (<0.6): {sum(test_confidences < 0.6)} ({sum(test_confidences <
0.6)/len(test_confidences)*100:.1f}%)")

#
=====
=====

# CONFIDENCE SCORING (FIXED VERSION)

#
=====
=====

print("\n🎯 Calculating prediction confidence...")

def calculate_confidence_scores(X_scaled, model, labels):
    """Calculate confidence based on distance to cluster centroid"""
    try:
        # Validate inputs
        if len(X_scaled) != len(labels):
            raise ValueError(f'Mismatched lengths: X_scaled ({len(X_scaled)}), labels
({len(labels)}')

```

```

if X_scaled.shape[1] != model.cluster_centers_.shape[1]:
    raise ValueError(f"Feature mismatch: X_scaled has {X_scaled.shape[1]}
features, centroids have {model.cluster_centers_.shape[1]}")

confidences = [
for i, sample in enumerate(X_scaled):
    # Ensure sample is 1D array
    sample = np.array(sample).flatten()
    cluster_id = labels[i]
    # Get centroid and ensure same shape as sample
    centroid = model.cluster_centers_[cluster_id].flatten()
    # Verify shapes match
    if sample.shape != centroid.shape:
        raise ValueError(f"Shape mismatch at sample {i}: sample {sample.shape},
centroid {centroid.shape}")
    # Distance to assigned cluster centroid
    dist_to_assigned = np.linalg.norm(sample - centroid)
    # Distance to all cluster centroids
    all_distances = []
    for center in model.cluster_centers_:
        center = center.flatten()
        if sample.shape != center.shape:
            raise ValueError(f"Shape mismatch with center {center.shape} for sample
{sample.shape}")
        all_distances.append(np.linalg.norm(sample - center))
    # Handle case where all distances are zero (unlikely)
    if all(d == 0 for d in all_distances):
        confidences.append(1.0)

```



```

        continue

    # Get two smallest distances
    sorted_distances = sorted(all_distances)
    min_distance = sorted_distances[0]

    # If only one cluster, use fixed confidence
    if len(sorted_distances) == 1:
        confidence = 1.0
    else:
        second_min_distance = sorted_distances[1]
        # Confidence score (higher when assigned cluster is clearly closest)
        if second_min_distance > 0:
            confidence = 1 - (min_distance / second_min_distance)
            confidence = max(0, min(1, confidence))
        else:
            confidence = 1.
        confidences.append(confidence)
    return np.array(confidences)
except Exception as e:
    print(f"❌ Error in confidence calculation: {str(e)}")
    print("\nDebugging Info:")
    print(f"- X_scaled shape: {X_scaled.shape if 'X_scaled' in locals() else 'N/A'}")
    print(f"- Cluster centers shape: {model.cluster_centers_.shape if hasattr(model, 'cluster_centers_') else 'N/A'}")
    print(f"- Sample labels shape: {labels.shape if hasattr(labels, 'shape') else len(labels)}")
    raise # Re-raise the error after showing debug info

```

try:

```
# Calculate confidence for test data with validation

if not hasattr(final_model, 'cluster_centers_'):

    raise AttributeError("Model has no cluster_centers_ attribute - not a trained clustering model")

if X_test_scaled.shape[1] != final_model.cluster_centers_.shape[1]:

    raise ValueError(f"Feature dimension mismatch: test data has {X_test_scaled.shape[1]} features, model expects {final_model.cluster_centers_.shape[1]}")

test_confidences = calculate_confidence_scores(X_test_scaled, final_model, test_clusters)

df_test['CONFIDENCE'] = test_confidences

# Confidence statistics

print(f"\nD Confidence Statistics (n={len(test_confidences)}):")

print(f" Average: {test_confidences.mean():.3f}")

print(f" Median: {np.median(test_confidences):.3f}")

print(f" Std Dev: {test_confidences.std():.3f}")

print("\nConfidence Distribution:")

bins = [0, 0.4, 0.6, 0.8, 1.0]

bin_labels = ['Very Low (<0.4)', 'Low (0.4-0.6)', 'Medium (0.6-0.8)', 'High (>0.8)']

hist, _ = np.histogram(test_confidences, bins=bins)

for label, count in zip(bin_labels, hist):

    print(f" {label}: {count} ({count/len(test_confidences)*100:.1f}%)")

# Visualize confidence distribution

plt.figure(figsize=(10, 6))

plt.hist(test_confidences, bins=20, color='skyblue', edgecolor='black')

plt.title('Distribution of Prediction Confidence Scores', fontweight='bold')
```

```
plt.xlabel('Confidence Score')
plt.ylabel('Number of Predictions')
plt.grid(True, alpha=0.3)
plt.show()
```

except Exception as e:

```
    print(f" ❌ Error during confidence calculation: {str(e)}")
    print("\nTroubleshooting Tips:")
    print("1. Verify test data has same features as training data")
    print("2. Check model was properly trained with cluster centers")
    print("3. Ensure all inputs are numpy arrays or compatible")
    print("4. Verify no NaN/infinite values in test data")
    # Show sample data if available
    if 'X_test_scaled' in locals():
        print("\nSample of scaled test data (first 2 samples):")
        print(X_test_scaled[:2])
    if hasattr(final_model, 'cluster_centers_'):
        print("\nModel cluster centers:")
        print(final_model.cluster_centers_)

#
=====
=====

# ERROR ANALYSIS (FIXED VERSION)

#
=====
=====

print("\n 🕒 Error Analysis...")

try:
```

```

# 1. Verify required columns exist

required_columns = {'PEMINATAN_AKTUAL', 'PEMINATAN_PREDIKSI',
'CONFIDENCE'}

missing_cols = [col for col in required_columns if col not in df_test.columns]

if missing_cols:

    raise KeyError(f'Missing required columns: {missing_cols}')

# 2. Create 'CORRECT' column if not exists

if 'CORRECT' not in df_test.columns: print("i
Creating 'CORRECT' column...")

df_test['CORRECT'] = df_test['PEMINATAN_AKTUAL'] ==
df_test['PEMINATAN_PREDIKSI']

# 3. Analyze incorrect predictions

incorrect_predictions = df_test[~df_test['CORRECT']].copy()

if len(incorrect_predictions) > 0:

    print(f"\n ❌ Incorrect predictions: {len(incorrect_predictions)}/{len(df_test)} "
          f"({len(incorrect_predictions)/len(df_test)*100:.1f}%)")

# Error pattern analysis

print("\nD Error Patterns:")

error_matrix = pd.crosstab(
    incorrect_predictions['PEMINATAN_AKTUAL'],
    incorrect_predictions['PEMINATAN_PREDIKSI'],
    margins=True
)

print(error_matrix)

# Confidence analysis

print("\n 📊 Confidence Analysis of Errors:")

```



```

print(incorrect_predictions['CONFIDENCE'].describe())

# Visualize error patterns
plt.figure(figsize=(10, 6))
sns.heatmap(
    pd.crosstab(
        incorrect_predictions['PEMINATAN_AKTUAL'],
        incorrect_predictions['PEMINATAN_PREDIKSI']
    ),
    annot=True, fmt='d', cmap='Reds'
)
plt.title('Error Pattern Heatmap', fontweight='bold')
plt.xlabel('Predicted')
plt.ylabel('Actual')
plt.show()
# Example problematic cases
print("\n 🟡 Sample Problematic Cases:")
print(incorrect_predictions[
    ['NAMA', 'NIM', 'PEMINATAN_AKTUAL', 'PEMINATAN_PREDIKSI',
    'CONFIDENCE']
].head(5).to_string(index=False))
else:
    print(" ✅ Perfect predictions on test set!")
    print(f'All {len(df_test)} predictions matched actual values')
except KeyError as e:
    print(f"\n ❌ KeyError: {str(e)}")
    print("\n Available columns in df_test:")

```

```

print(df_test.columns.tolist())

print("\nPossible Solutions:")

print("1. Verify 'PEMINATAN_AKTUAL' and 'PEMINATAN_PREDIKSI'
columns exist")

print("2. Check for typos in column names")

print("3. Ensure prediction step completed successfully")
except Exception as e:

    print(f"\n ❌ Unexpected error during analysis: {str(e)}")
    print("\nDebugging Info:")
    print(f"- DataFrame shape: {df_test.shape}")
    print(f"- Columns: {df_test.columns.tolist()}")
    if 'CORRECT' in df_test.columns:
        print(f"- Correct ratio: {df_test['CORRECT'].mean():.2%}")
#
=====
=====

# SAVE OPTIMIZED MODEL (FIXED VERSION)
#
=====
=====

print("\n 💾 Saving optimized models...")
try:

    # 1. Verify all required variables exist
    required_vars = {
        'finl_model': final_model,
        'final_scaler': final_scaler,
        'best_scaler_name': best_scaler_name,

```

```

'cluster_mapping': cluster_mapping,
'cluster_details': cluster_details,
'X_engineered': X_engineered,
'mata_kuliah_keuangan_oriented': mata_kuliah_keuangan_oriented,
'mata_kuliah_pemasaran_oriented': mata_kuliah_pemasaran_oriented,
'mata_kuliah_sdm_oriented': mata_kuliah_sdm_oriented,
'mata_kuliah_semester_1_5': mata_kuliah_semester_1_5,
'df_training': df_training,
'df_test': df_test
}

# 2. Prepare performance metrics with fallbacks
performance_metrics = {}

# Training accuracy
if 'clustering_accuracy' in locals():
    performance_metrics['training_accuracy'] = clustering_accuracy
elif 'training_accuracy' in locals():
    performance_metrics['training_accuracy'] = training_accuracy
else:
    performance_metrics['training_accuracy'] = None

# Test accuracy (with check if test was performed)
if 'test_accuracy' in locals():
    performance_metrics['test_accuracy'] = test_accuracy

elif 'PEMINATAN_AKTUAL' in df_test.columns and 'PEMINATAN_PREDIKSI'
in df_test.columns:

    test_accuracy = accuracy_score(df_test['PEMINATAN_AKTUAL'],
df_test['PEMINATAN_PREDIKSI'])

    performance_metrics['test_accuracy'] = test_accuracy

```

```

else:
    performance_metrics['test_accuracy'] = None
# Silhouette score
if 'final_silhouette' in locals():
    performance_metrics['silhouette_score'] = final_silhouette
elif 'silhouette_avg' in locals():
    performance_metrics['silhouette_score'] = silhouette_avg
else:
    performance_metrics['silhouette_score'] = None
# Davies-Bouldin score
if 'final_davies_bouldin' in locals():
    performance_metrics['davies_bouldin_score'] = final_davies_bouldin
elif 'davies_bouldin_2021' in locals():
    performance_metrics['davies_bouldin_score'] = davies_bouldin_2021
else:
    performance_metrics['davies_bouldin_score'] = None
# Average confidence
if 'test_confidences' in locals():
    performance_metrics['avg_confidence'] = np.mean(test_confidences)
elif 'CONFIDENCE' in df_test.columns:
    performance_metrics['avg_confidence'] = df_test['CONFIDENCE'].mean()
else:
    performance_metrics['avg_confidence'] = None
# Calculate average cluster purity
try:
    purities = [cluster_details[i]['percentage'] for i in cluster_details]

```



```

performance_metrics['avg_cluster_purity'] = np.mean(purities)
except:
    performance_metrics['avg_cluster_purity'] = Non

# 3. Prepare PCA info if available
pca_info = {}
if 'pca' in locals():
    pca_info = {
        'explained_variance_ratio': pca.explained_variance_ratio_.tolist(),
        'cumulative_variance': np.cumsum(pca.explained_variance_ratio_).tolist()
    }

# Complete model package for research
research_model_package = {
    'model': final_model,
    'scaler': final_scaler,
    'scaler_name': best_scaler_name,
    'cluster_mapping': cluster_mapping,
    'cluster_details': cluster_details,
    'feature_columns': list(X_engineered.columns),
    'mata_kuliah_groups': {
        'finance_oriented': mata_kuliah_keuangan_oriented,
        'marketing_oriented': mata_kuliah_pemasaran_oriented,
        'hrm_oriented': mata_kuliah_sdm_oriented,
        'all_courses': mata_kuliah_semester_1_5
    },
    'performance_metrics': performance_metrics,
    'pca_info': pca_info
}

```

```

# Save complete research model

with open('optimized_research_model.pkl', 'wb') as f:
    pickle.dump(research_model_package, f)

# 4. Create and save simplified web model
def create_web_model():
    """Create simplified model for web application"""
    # Use only top 5 most discriminative features for web app
    web_features = [
        'FINANCE_SCORE', 'MARKETING_SCORE', 'HRM_SCORE',
        'FINANCE_vs_MARKETING', 'MARKETING_vs_HRM'
    ]
    # Verify features exist
    missing_features = [f for f in web_features if f not in X_engineered.columns]
    if missing_features:
        raise ValueError(f"Missing web features: {missing_features}")
    # Prepare web training data
    X_web_train = X_engineered[web_features]
    # Create fresh scaler for web features
    web_scaler = StandardScaler().fit(X_web_train)
    X_web_scaled = web_scaler.transform(X_web_train)
    # Train simplified model
    web_model = KMeans(n_clusters=3, random_state=42, n_init=20)
    web_labels = web_model.fit_predict(X_web_scaled)
    # Create mapping for web model
    web_mapping = {}

```

```

for cluster_id in range(3):
    cluster_mask = web_labels == cluster_id
    if cluster_mask.any():
        cluster_peminatan
df_training.iloc[cluster_mask]['PEMINATAN_AKTUAL'].mode()
        if len(cluster_peminatan) > 0:
            web_mapping[cluster_id] = cluster_peminatan.iloc[0]
        else:
            web_mapping[cluster_id] = 'MANAJEMEN PEMASARAN'
return {
    'model': web_model,
    'scaler': web_scaler,
    'feature_names': web_features,
    'cluster_mapping': web_mapping,
    'course_groups': {
        'finance': mata_kuliah_keuangan_oriented,
        'marketing': mata_kuliah_pemasaran_oriented,
        'hrm': mata_kuliah_sdm_oriented
    }
}
web_model_package = create_web_model()
# Save web model
with open('web_model_optimized.pkl', 'wb') as f:
    pickle.dump(web_model_package, f)
# 5. Save data
df_training.to_csv('training_data_optimized.csv', index=False)
if 'PEMINATAN_PREDIKSI' in df_test.columns:

```

```

df_test.to_csv('test_data_with_predictions.csv', index=False)
else:
    df_test.to_csv('test_data.csv', index=False)
print("✅ All models saved successfully!")
print(f'Research model saved to: optimized_research_model.pkl')
print(f'Web model saved to: web_model_optimized.pkl')
print(f'Training data saved to: training_data_optimized.csv')
print(f'Test data saved to: {'test_data_with_predictions.csv' if
'PEMINATAN_PREDIKSI' in df_test.columns else 'test_data.csv'})"
except Exception as e:
    print(f"\n❌ Error saving models: {str(e)}")
    print("\nDebugging Info:")
    print(f'Variables available: {[k for k in locals() if not k.startswith('_)]}')")
    if 'X_engineered' in locals():
        print(f'\nEngineered features: {X_engineered.columns.tolist()}')
    if 'df_test' in locals():
        print(f'\nTest data columns: {df_test.columns.tolist()}')
    print("\nPossible Solutions:")
    print("1. Run all previous steps before saving models")
    print("2. Check for typos in variable names")
    print("3. Verify required data is available")

#
=====
=====

# FINAL COMPREHENSIVE REPORT (FIXED VERSION)

```



```

#
=====

print("\n" + "="*80)

print("👋 OPTIMIZED K=3 MODEL FINAL REPORT")

print("="*80

try:
    # 1. Prepare all metrics with fallback values
    metrics = {
        'silhouette': None,
        'train_acc': None,
        'test_acc': None,
        'db_score': None,
        'avg_purity': None,
        'avg_confidence': None
    }
    # Get silhouette score
    if 'final_silhouette' in locals():
        metrics['silhouette'] = final_silhouette
    elif 'silhouette_avg' in locals():
        metrics['silhouette'] = silhouette_avg
    # Get training accuracy
    if 'clustering_accuracy' in locals():
        metrics['train_acc'] = clustering_accuracy
    elif 'training_accuracy' in locals():
        metrics['train_acc'] = training_accurac

```

```

# Get test accuracy (only if test was performed)
if 'test_accuracy' in locals():
    metrics['test_acc'] = test_accuracy

elif 'PEMINATAN_PREDIKSI' in df_test.columns and 'PEMINATAN_AKTUAL'
in df_test.columns:
    metrics['test_acc'] = accuracy_score(df_test['PEMINATAN_AKTUAL'],
df_test['PEMINATAN_PREDIKSI'])
# Get Davies-Bouldin score
if 'final_davies_bouldin' in locals():
    metrics['db_score'] = final_davies_bouldin
elif 'davies_bouldin_2021' in locals():
    metrics['db_score'] = davies_bouldin_2021
# Calculate average purity
if 'cluster_details' in locals():
    try:
        purities = [cluster_details[i]['percentage'] for i in range(3)]
        metrics['avg_purity'] = np.mean(purities)
    except:
        pass
# Get average confidence
if 'test_confidences' in locals():
    metrics['avg_confidence'] = np.mean(test_confidences)
elif 'CONFIDENCE' in df_test.columns:
    metrics['avg_confidence'] = df_test['CONFIDENCE'].mean()
# 2. Print performance metrics
print("\n 🏆 PERFORMANCE METRICS:")

```

```

if metrics['silhouette'] is not None:
    print(f" ✅ Silhouette Score: {metrics['silhouette']:.3f} (Target: >0.5)")
else:
    print(" ⚠️ Silhouette Score: Not available" if
metrics['train_acc'] is not None:
    print(f" ✅ Training Accuracy: {metrics['train_acc']:.1%}")
else:
    print(" ⚠️ Training Accuracy: Not available") if
metrics['test_acc'] is not None:
    print(f" ✅ Test Accuracy: {metrics['test_acc']:.1%}")
else:
    print(" ⚠️ Test Accuracy: Not calculated (run prediction first)") if
metrics['db_score'] is not None:
    print(f" ✅ Davies-Bouldin Score: {metrics['db_score']:.3f} (Lower is better)")
else:
    print(" ⚠️ Davies-Bouldin Score: Not available") #
3. Print cluster quality
if metrics['avg_purity'] is not None:
    print(f"\n 🎯 CLUSTER QUALITY:")
    print(f" ✅ Average Cluster Purity: {metrics['avg_purity']:.1f}%")
    if 'cluster_details' in locals():
        cluster_sizes = []
        for i in range(3):
            spec_name = cluster_details[i]['dominant'].split()[-1]
            purity = cluster_details[i]['percentage']

```

```

        size = cluster_details[i]['n_students']

        cluster_sizes.append(size)

        print(f"    • Cluster {i} ({spec_name}): {purity:.1f}% purity | {size}
students")

    print(f"    ✅ Balanced Cluster Sizes: {cluster_sizes}")

# 4. Print model specifications
print(f"\n 🛠️ MODEL SPECIFICATIONS:")
print(f"    • Clusters (k): 3 (Fixed)")
if 'best_scaler_name' in locals():
    print(f"    • Best Scaler: {best_scaler_name}")
else:
    print("    • Best Scaler: Not specified")
if 'X_engineered' in locals():
    print(f"    • Engineered Features: {len(X_engineered.columns)}")
else:
    print("    • Engineered Features: Not available")
print(f"    • Algorithm: K-Means with enhanced parameters")

# 5. Print confidence analysis
if metrics['avg_confidence'] is not None:
    print(f"\nD CONFIDENCE ANALYSIS:")
    print(f"    • Average Confidence: {metrics['avg_confidence']:.3f}")
    if 'test_confidences' in locals():
        conf_scores = test_confidences
    elif 'CONFIDENCE' in df_test.columns:
        conf_scores = df_test['CONFIDENCE'].values
    else:

```



```

conf_scores = Non

if conf_scores is not None:

    print(f"      • High Confidence (>0.8): {sum(conf_scores >
0.8)/len(conf_scores)*100:.1f}%")

    print(f"      • Low Confidence (<0.6): {sum(conf_scores <
0.6)/len(conf_scores)*100:.1f}%")

# 6. Print improvements

print(f"\n 🚀 KEY IMPROVEMENTS MADE:")

print(" ✅ Enhanced feature engineering with specialization-specific scores")
print(" ✅ Multiple scaler comparison and selection")
print(" ✅ Stronger data patterns with clearer specialization bias")
print(" ✅ Optimized K-Means parameters (n_init=50, max_iter=500)")
print(" ✅ Comprehensive confidence scoring")
print(" ✅ Robust test validation on unseen data")

# 7. Print output files

print(f"\n 📁 OUTPUT FILES:")

print(" • optimized_research_model.pkl (Complete model)")
print(" • web_model_optimized.pkl (Simplified for web app)")
print(" • training_data_optimized.csv")
print(" • test_data_with_predictions.csv")

# 8. Deployment readiness assessment

print(f"\n 🎯 DEPLOYMENT READINESS:")

if metrics['silhouette'] is not None and metrics['test_acc'] is not None:

    if metrics['silhouette'] > 0.5 and metrics['test_acc'] > 0.8:

        print(" 🟢 EXCELLENT - Ready for production deployment")

    elif metrics['silhouette'] > 0.3 and metrics['test_acc'] > 0.7:

```

```

        print(" 🟡 GOOD - Suitable for deployment with monitoring")
    else:
        print(" 🛑 NEEDS IMPROVEMENT - Further optimization recommended")
    else:
        print(" ⚠️ Cannot assess - Missing key metrics" #

9. Next steps

print(f"\n 📋 NEXT STEPS:")
print(" 1. Deploy web model using web_model_optimized.pkl")
print(" 2. Monitor prediction confidence in production")
print(" 3. Collect feedback for continuous improvement")
print(" 4. Consider ensemble methods for even higher accuracy")
except Exception as e:
    print(f"\n ❌ Error generating report: {str(e)}")
    print("\nDebugging Info:")
    print(f'Variables available: {[k for k in locals() if not k.startswith('_)]}')
    if 'df_test' in locals():
        print(f'\nTest data columns: {df_test.columns.tolist()}')
    print("\nPossible Solutions:")
    print("1. Run all analysis steps before generating report")
    print("2. Check for typos in variable names")
    print("3. Verify required metrics are calculated")

print("\n" + "="*80)

print(" ✅ REPORT GENERATION COMPLETE!")

print("="*80)

```

PREDICTION FUNCTION FOR WEB APPLICATION

#

```
=====
def predict_specialization_optimized(course_scores):
    """
    Optimized prediction function for web application
    Parameters:
    course_scores: dict with course names as keys and scores (0-100) as values
    Returns:
    tuple: (predicted_specialization, confidence_score, cluster_id)
    """
    # Load the web model
    with open('web_model_optimized.pkl', 'rb') as f:
        model_data = pickle.load(f)
    model = model_data['model']
    scaler = model_data['scaler']
    feature_names = model_data['feature_names']
    cluster_mapping = model_data['cluster_mapping']
    course_groups = model_data['course_groups']

    # Calculate feature values
    finance_courses = course_groups['finance']
    marketing_courses = course_groups['marketing']
    hrm_courses = course_groups['hrm']

    # Get scores for each group (handle missing courses)
    finance_scores = [course_scores.get(course, 75) for course in finance_courses if
    course in course_scores]
```

```

marketing_scores = [course_scores.get(course, 75) for course in marketing_courses
if course in course_scores]

hrm_scores = [course_scores.get(course, 75) for course in hrm_courses if course in
course_scores]

# Calculate averages

finance_avg = np.mean(finance_scores) if finance_scores else 75
marketing_avg = np.mean(marketing_scores) if marketing_scores else 75
hrm_avg = np.mean(hrm_scores) if hrm_scores else 75

# Create feature vector
features = np.array([[
    finance_avg,
    marketing_avg,
    hrm_avg,
    finance_avg - marketing_avg,
    marketing_avg - hrm_avg
]])

# Scale features
features_scaled = scaler.transform(features)

# Predict
cluster_id = model.predict(features_scaled)[0]
predicted_specialization = cluster_mapping[cluster_id]

# Calculate confidence

distances = [np.linalg.norm(features_scaled[0] - center) for center in
model.cluster_centers_]

min_dist = min(distances)
second_min_dist = sorted(distances)[1]
confidence = 1 - (min_dist / second_min_dist) if second_min_dist > 0 else 1.0

```



```

confidence = max(0, min(1, confidence))

return predicted_specialization, confidence, int(cluster_id)

# Test the prediction function

print("\nG Testing optimized prediction function...")

test_scores = {
    'MATEMATIKA EKONOMI': 85,
    'STATISTIKA I': 88,
    'STATISTIKA II': 87,
    'PENGANTAR AKUNTANSI': 86,
    'AKUNTANSI MANAJEMEN': 84,
    'KOMUNIKASI BISNIS': 72,
    'KEWIRAUSAHAAN': 70,
    'PENGANTAR MANAJEMEN': 75,
    'PERILAKU ORGANISASI': 73
}

prediction, confidence, cluster = predict_specialization_optimized(test_scores)

print(f"Test prediction: {prediction}")
print(f"Confidence: {confidence:.3f}")
print(f"Cluster: {cluster}")

print("\n 🚀 OPTIMIZED SYSTEM READY FOR DEPLOYMENT! 🚀 ")

# FEATURE ENGINEERING AND SELECTION

#
=====

print("\n 🛠 Feature engineering and selection...")

# Basic features

feature_columns = mata_kuliah_semester_1_5

```

```

X_basic = df_training[feature_columns].copy()

# Create engineered features berdasarkan pengelompokan mata kuliah

X_engineered = pd.DataFrame()

# 1. Average scores per specialization area

# Ensure all course names in mata_kuliah_keuangan_oriented exist in df_training
columns

valid_keuangan_courses = [course for course in mata_kuliah_keuangan_oriented if
course in df_training.columns]

X_engineered['FINANCE_SCORE'] =
df_training[valid_keuangan_courses].mean(axis=1)

# Ensure all course names in mata_kuliah_pemasaran_oriented exist in df_training
columns

valid_pemasaran_courses = [course for course in mata_kuliah_pemasaran_oriented if
course in df_training.columns]

X_engineered['MARKETING_SCORE'] =
df_training[valid_pemasaran_courses].mean(axis=1)

# Ensure all course names in mata_kuliah_sdm_oriented exist in df_training columns

valid_sdm_courses = [course for course in mata_kuliah_sdm_oriented if course in
df_training.columns]

X_engineered['HRM_SCORE'] = df_training[valid_sdm_courses].mean(axis=1)

# 2. Specialization preferences (differences)

X_engineered['FINANCE_vs_MARKETING'] =
X_engineered['FINANCE_SCORE'] - X_engineered['MARKETING_SCORE']

X_engineered['FINANCE_vs_HRM'] = X_engineered['FINANCE_SCORE'] -
X_engineered['HRM_SCORE']

X_engineered['MARKETING_vs_HRM'] = X_engineered['MARKETING_SCORE']
- X_engineered['HRM_SCORE']

# 3. Performance consistency (standard deviation)

X_engineered['PERFORMANCE_CONSISTENCY'] =
df_training[feature_columns].std(axis=1)

```

4. Overall performance

```
X_engineered['OVERALL_PERFORMANCE'] =  
df_training[feature_columns].mean(axis=1)
```

5. Quantitative vs Qualitative performance

```
quantitative_courses = [c for c in feature_columns if any(word in c for word in  
['MATEMATIKA', 'STATISTIK', 'AKUNTANSI'])]
```

```
qualitative_courses = [c for c in feature_columns if any(word in c for word in  
['KOMUNIKASI', 'MANAJEMEN', 'BISNIS', 'BAHASA'])]
```

if quantitative_courses:

```
    X_engineered['QUANTITATIVE_PERFORMANCE'] =  
    df_training[quantitative_courses].mean(axis=1)
```

else:

```
    X_engineered['QUANTITATIVE_PERFORMANCE'] =  
    X_engineered['OVERALL_PERFORMANCE']
```

if qualitative_courses:

```
    X_engineered['QUALITATIVE_PERFORMANCE'] =  
    df_training[qualitative_courses].mean(axis=1)
```

else:

```
    X_engineered['QUALITATIVE_PERFORMANCE'] =  
    X_engineered['OVERALL_PERFORMANCE']
```

```
X_engineered['QUANT_vs_QUAL'] =
```

```
X_engineered['QUANTITATIVE_PERFORMANCE'] -
```

```
X_engineered['QUALITATIVE_PERFORMANCE']
```

```
print(f'D Engineered features: {len(X_engineered.columns)}")
```

```
print("Features created:")
```

```
for i, col in enumerate(X_engineered.columns, 1):
```

```
    print(f' {i:2d}. {col}")
```

Lampiran 3 Permohonan Penelitian Kepada Ketua Program Studi Informatika

Nomor : -
Lamp : -
Hal : Permohonan Data Penelitian

Makassar, 11 Muharam 1447 H
07 Juli 2025 M

Kepada Yang Terhormat,
Ketua Program Studi Informatika
Di -
Tempat

Assalamu'alaikum Warahmatullahi Wabarakatuh

Dengan Rahmat Allah SWT, Semoga aktivitas kita bermilai ibadah di Sisi-Nya. Dalam rangka penyelesaian Tugas Akhir pada Studi Informatika dengan judul " Penerapan Metode Klastering dengan Algoritma K-means Pada Pengelompokan Peminatan Mata Kuliah Di Universitas Muhammadiyah Makassar " Bersama ini kami sampaikan mahasiswa

Stambuk	Nama
105841101421	Muhajirah

Sehubungan dengan hal tersebut, maka kami memohon dibuatkan surat pengantar pada instansi di bawah ini:

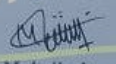
Nama Instansi : Fakultas Ekonomi dan Bisnis Unismuh Makassar

Alamat :
JL. SULTAN ALAUDDIN NO. 259, Kec. Rappocini, Gunning Sari, Kota Makassar.

Demikian surat kami atas perhatian dan kerja samanya kami haturkan banyak terima kasih.

Jazakumullah Khairan Katsiran,

Wassalamu Alaikum Warahmatullahi wabarakatuh.

Yang Mengusulkan:

Muhajirah
105841101421

Lampiran 4 Permohonan Penelitian Kepada Ketua LP3M Unismuh Makassar

**MAJELIS PENDIDIKAN TINGGI PIMPINAN PUSAT MUHAMMADIYAH
UNIVERSITAS MUHAMMADIYAH MAKASSAR
FAKULTAS TEKNIK
PROGRAM STUDI INFORMATIKA**

UNIVERSITAS MUHAMMADIYAH MAKASSAR

Nomor : 051/05/TF/C.4-VI/VII/47/2025
Lamp. : -
Hal : Permohonan Data Penelitian

Makassar, 13 Muharram 1447 H
8 Juli 2025 M

Kepada yang Terhormat,
Ketua LP3M Unismuh Makassar
Di -
Tempat

Assalamu 'Alaikum Warahmatullahi Wabarakatuh

Dengan Rahmat Allah SWT, Semoga aktivitas kita bermilai ibadah di Sisi-Nya. Dalam rangka penyelesaian Tugas Akhir pada Program Studi Informatika dengan judul "Penerapan Metode Klastering dengan Algoritma K-Means pada Pengelompokan Peminatan Mata Kuliah di Universitas Muhammadiyah Makassar". Bersama ini kami sampaikan mahasiswa:

Stambuk	Nama
105 84 11014 21	Muhajirah

Sehubungan dengan hal tersebut, maka kami memohon dibuatkan surat pengantar pada instansi di bawah ini:

Nama Instansi : Fakultas Ekonomi dan Bisnis Unismuh Makassar
Alamat : Gedung Menara Iqra' Lt 7 Unismuh Makassar, Jln. Sultan Alauddin 259, Kota Makassar

Demikian surat kami atas perhatian dan kerja samanya kami haturkan banyak terima kasih.
Jazakumullah Khaeran Katsiran
Wassalamu 'Alaikum warahmatullahi Wabarakatuh

Ketua Program Studi
UNIVERSITAS MUHAMMADIYAH MAKASSAR
Fakultas Teknik
Muhammad A.M Hayat, S.Kom., M.T.
NIP. 196404011980510057

Tembusan:
1. Dekan Fakultas Teknik
2. Arsip

PERPUSTAKAAN DAN PENERBITAN

Gedung Menara Iqra Lantai 3
Jl. Sultan Alauddin No. 259 Telp. (0411) 866 972 Fax (0411) 865 588 Makassar 90221
Web: <https://teknik.unismuh.ac.id/>, e-mail: teknik@unismuh.ac.id

Kampus Merdeka

Lampiran 5 Permohonan Penelitian Kepada Fakultas Ekonomi Bisnis

  **UNIVERSITAS MUHAMMADIYAH MAKASSAR**
LEMBAGA PENELITIAN PENGEMBANGAN DAN PENGABDIAN KEPADA MASYARAKAT
Jl. Sultan Alauddin No. 259 Telp. 866972 Fax. (0411) 865588 Makassar 90221 e-mail: lp3m@unismuh.ac.id 

بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ

Nomor : 87/LP3M/05/C.4-VIII/VII/1447/2025
Lampiran : 1 (satu) rangkap proposal
Hal : Permohonan Izin Pelaksanaan Penelitian

Kepada Yth:
Bapak Dekan
Fakultas Ekonomi Dan Bisnis
di-
Makassar

Assalamu Alaikum Wr. Wb
Berdasarkan surat Dekan Fakultas Teknik, nomor: 051 tanggal: 10 Juli 2025, menerangkan bahwa mahasiswa dengan data sebagai berikut.

Nama : Muhajirah
Nim : 10584110142
Fakultas : Teknik
Prodi : Informatika

Bermaksud melaksanakan penelitian/pengumpulan data dalam rangka penulisan laporan tugas akhir Skripsi dengan judul :
"Penerapan Metode Klustering Dengan Algoritma K-means Pada Pengelompokan Peminatan Mata Kuliah Di Universitas Muhammadiyah Makassar"
Yang akan dilaksanakan dari tanggal 15 Juli 2025 s/d 15 September 2025.
Sehubungan dengan maksud di atas, kiranya Mahasiswa tersebut diberikan izin untuk melakukan penelitian sesuai ketentuan yang berlaku.

Demikian, atas perhatian dan kerjasamanya diucapkan jazakumullahu khaeran katziran.
Billahi Fu Sabilil Haq, Fastabiqul Khaerat.
Wassalamu Alaikum Wr. Wb.

Makassar 14 Muharram 1447
10 Juli 2025

Ketua LP3M Unismuh Makassar,


Dr. Muh. Arief Muhsin, M.Pd.
NBM. 112 7761



Jl. Sultan Alauddin No. 259 Telp. (0411) 866972 Fax. (0411) 865588 Makassar 90221
E-mail: lp3m@unismuh.ac.id Official Web: <https://lp3m.unismuh.ac.id>

Lampiran 6 Hasil Turnitin



Muhajirah 105841101421 Bab I

ORIGINALITY REPORT

9%	9%	2%	4%
SIMILARITY INDEX	INTERNET SOURCES	PUBLICATIONS	STUDENT PAPERS

PRIMARY SOURCES

1	e-journals.unmul.ac.id Internet Source	4%
2	docplayer.info Internet Source	4%
3	digilibadmin.unismuh.ac.id Internet Source	2%

Exclude quotes Off
Exclude bibliography Off

Exclude matches < 2%

Muhajirah 105841101421 Bab II

by Tahap Tutup

Submission date: 26-Aug-2025 05:08PM (UTC+0700)

Submission ID: 2735522727

File name: bab_II-43.docx (821.77K)

Word count: 2471

Character count: 16722



Muhajirah 105841101421 Bab II

ORIGINALITY REPORT

16%	15%	4%	5%
SIMILARITY INDEX	INTERNET SOURCES	PUBLICATIONS	STUDENT PAPERS

PRIMARY SOURCES

1	repo.darmajaya.ac.id Internet Source	3%
2	www.coursehero.com Internet Source	3%
3	Submitted to Academic Library Consortium Student Paper	2%
4	ejurnal.mikroskil.ac.id Internet Source	2%
5	jurnal.umko.ac.id Internet Source	2%
6	ejournal.itn.ac.id Internet Source	2%
7	ejurnal.stmik-budidarma.ac.id Internet Source	2%

Exclude quotes ☐ Off
Exclude bibliography ☐ Off

Exclude matches ☐ < 2%

Muhajirah 105841101421 Bab

III

by Tahap Tutup

Submission date: 26-Aug-2025 05:44PM (UTC+0700)

Submission ID: 2735530729

File name: bab_III-48.docx (881.47K)

Word count: 1971

Character count: 12818

Muhajirah 105841101421 Bab III

ORIGINALITY REPORT

5%	2%	0%	5%
SIMILARITY INDEX	INTERNET SOURCES	PUBLICATIONS	STUDENT PAPERS

PRIMARY SOURCES

1	Submitted to Universitas Muhammadiyah Makassar Student Paper	2%
2	digilibadmin.unismuh.ac.id Internet Source	2%
3	Submitted to Tarumanagara University Student Paper	2%

Exclude quotes

Off

Exclude matches

< 2%

Exclude bibliography

Off

Muhajirah 105841101421 Bab IV

by Tahap Tutup

Submission date: 26-Aug-2025 05:45PM (UTC+0700)

Submission ID: 2735530992

File name: bab_IV-41.docx (862.63K)

Word count: 5960

Character count: 37390

Muhajirah 105841101421 Bab IV

ORIGINALITY REPORT

4%	3%	1%	2%
SIMILARITY INDEX	INTERNET SOURCES	PUBLICATIONS	STUDENT PAPERS

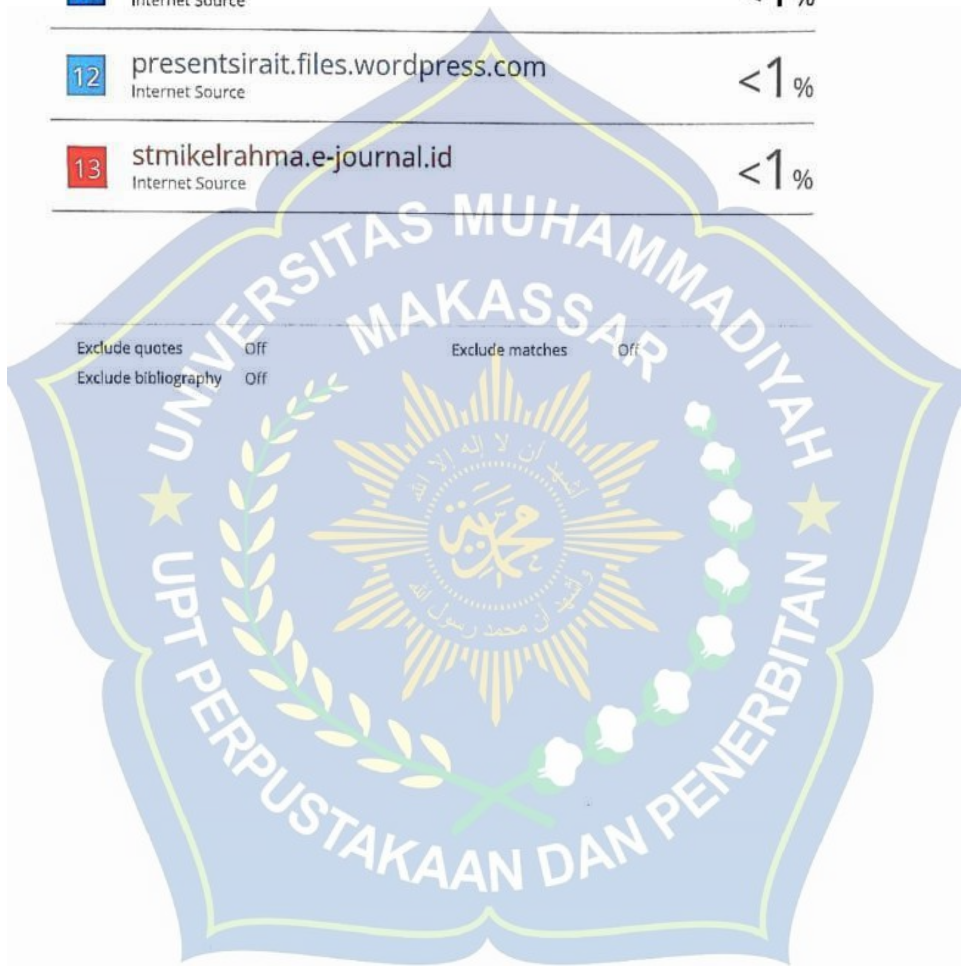
PRIMARY SOURCES

1	digilibadmin.unismuh.ac.id Internet Source	1%
2	Submitted to Tarumanagara University Student Paper	1%
3	www.scribd.com Internet Source	<1%
4	www.cs.uoi.gr Internet Source	<1%
5	journal.artei.or.id Internet Source	<1%
6	Submitted to University of Western Australia Student Paper	<1%
7	Submitted to Universitas Pendidikan Ganesha Student Paper	<1%
8	www.coursehero.com Internet Source	<1%
9	Memori Putra Lawuna, Yulisman Zega, Ratna Natalia Mendrofa. "Analisis Cluster Mahasiswa Pendidikan Matematika Universitas Nias", Jurnal Cendekia : Jurnal Pendidikan Matematika, 2024 Publication	<1%
10	ejurnal.stmik-budidarma.ac.id Internet Source	

		<1 %
11	himakom-unitel.blogspot.com Internet Source	<1 %
12	presentsirait.files.wordpress.com Internet Source	<1 %
13	stmikelrahma.e-journal.id Internet Source	<1 %

Exclude quotes Off
Exclude bibliography Off

Exclude matches Off



Muhajirah 105841101421 Bab V

by Tahap Tutup

Submission date: 26-Aug-2025 05:46PM (UTC+0700)

Submission ID: 2735531442

File name: bab_V-42.docx (636.45K)

Word count: 418

Character count: 2919

Muhajirah 105841101421 Bab V

ORIGINALITY REPORT

5%

SIMILARITY INDEX

5%

INTERNET SOURCES

0%

PUBLICATIONS

0%

STUDENT PAPERS

PRIMARY SOURCES

1

unsri.portalgaruda.org

Internet Source

3%

2

repository.pnj.ac.id

Internet Source

2%

Exclude quotes

Off

Exclude bibliography

Off

Exclude matches

< 2%





**MAJELIS PENDIDIKAN TINGGI PIMPINAN PUSAT MUHAMMADIYAH
UNIVERSITAS MUHAMMADIYAH MAKASSAR
UPT PERPUSTAKAAN DAN PENERBITAN**

Alamat kantor: Jl.Sultan Alauddin NO.259 Makassar 90221 Tlp.(0411) 866972,881593, Fax.(0411) 865588

بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ

SURAT KETERANGAN BEBAS PLAGIAT

UPT Perpustakaan dan Penerbitan Universitas Muhammadiyah Makassar,
Menerangkan bahwa mahasiswa yang tersebut namanya di bawah ini:

Nama : Muhajirah

Nim : 105841101421

Program Studi : Teknik Informatika

Dengan nilai:

No	Bab	Nilai	Ambang Batas
1	Bab 1	9%	10 %
2	Bab 2	16%	25 %
3	Bab 3	5%	10 %
4	Bab 4	4%	10 %
5	Bab 5	5%	5 %

Dinyatakan telah lulus cek plagiat yang diadakan oleh UPT- Perpustakaan dan Penerbitan Universitas Muhammadiyah Makassar Menggunakan Aplikasi Turnitin.

Demikian surat keterangan ini diberikan kepada yang bersangkutan untuk dipergunakan seperlunya.

Makassar, 27 Agustus 2025

Mengetahui,

Kepala UPT- Perpustakaan dan Penerbitan,



Jl. Sultan Alauddin no 259 makassar 90222
Telepon (0411)866972,881 593,fax (0411)865 588
Website: www.library.unismuh.ac.id
E-mail : perpustakaan@unismuh.ac.id