

BAB II

TINJAUAN PUSTAKA DAN DASAR TEORI

2.1 Tinjauan Pustaka

Analisis terhadap perilaku mahasiswa pada LMS merupakan isu utama dalam ranah *learning analytics* dan *educational data mining* (EDM). Data aktivitas yang tercatat di LMS seperti frekuensi login, akses ke materi pembelajaran, partisipasi dalam forum diskusi, dan perilaku pengumpulan tugas memberikan wawasan menyeluruh mengenai tingkat keterlibatan serta pola belajar mahasiswa. Papamitsiou dan Economides (2014) menemukan bahwa teknik EDM seperti klasifikasi dan klusterisasi mampu mengidentifikasi mahasiswa berisiko gagal secara dini, sehingga memungkinkan intervensi yang lebih tepat sasaran.

Al-Shabandar et al (2019) mengaplikasikan EDM dan *machine learning* untuk memprediksi hasil belajar pada kursus daring masif menunjukkan bahwa algoritma *gradient boosting* dan *support vector machine* memberikan akurasi prediksi yang tinggi terhadap tingkat kelulusan peserta. Selain itu, penelitian ini turut menyoroti pentingnya pemilihan fitur perilaku yang relevan, seperti frekuensi partisipasi dalam forum, interaksi dengan materi pembelajaran, serta keaktifan dalam pengumpulan tugas sebagai faktor yang berpengaruh terhadap hasil akademik. Performa akademik dapat diprediksi dengan klasifikasi *decision tree*, di mana hasilnya membantu institusi pendidikan dalam mendesain program remedial yang lebih personal (Dutt, et al, 2017).

Studi oleh Bakhshinategh et al (2018) memanfaatkan *text mining* pada data forum diskusi LMS, menghasilkan temuan bahwa tingkat keterlibatan mahasiswa dalam forum sangat memengaruhi ketercapaian kompetensi pembelajaran. Dengan menggunakan teknik analisis teks, penelitian ini berhasil mengidentifikasi pola partisipasi dan interaksi mahasiswa dalam LMS. Hasilnya menunjukkan bahwa mahasiswa yang aktif dalam forum diskusi lebih cepat memahami materi mahasiswaan, dan cenderung memperoleh nilai akademik yang lebih baik. Keterlibatan yang tinggi ini mencerminkan adanya budaya belajar kolaboratif, di mana mahasiswa saling

berbagi pengetahuan dan pengalaman, sekaligus memperkuat daya ingat serta pemahaman terhadap konsep yang dimahasiswainya. Penelitian ini sekaligus menegaskan bahwa optimalisasi fitur forum pada LMS mampu menjadi strategi efektif untuk meningkatkan ketercapaian hasil belajar dan mendukung terciptanya ekosistem pembelajaran digital yang lebih inklusif serta responsif terhadap keperluan mahasiswa. Romero et al (2019) mengaplikasikan *ensemble learning* untuk analisis prediktif, dan hasilnya menunjukkan peningkatan akurasi dalam mengklasifikasikan pola belajar mahasiswa serta mendukung pengambilan keputusan berbasis data di institusi pendidikan.

Hasil-hasil penelitian tersebut memperlihatkan bagaimana EDM dapat dimanfaatkan untuk deteksi dini risiko akademik, pengembangan strategi pembelajaran adaptif, hingga optimalisasi layanan pendidikan berbasis data perilaku mahasiswa di dalam LMS. Dengan demikian, EDM menjadi salah satu pendekatan utama dalam inovasi pengelolaan pendidikan digital di era modern.

Berbagai penelitian telah memanfaatkan LMS sebagai sumber data utama dalam analisis perilaku mahasiswa, dengan tujuan utama untuk meningkatkan efektivitas pembelajaran daring. Segmentasi mahasiswa berdasarkan perilaku pada LMS umumnya dilakukan menggunakan metode klusterisasi yang mengelompokkan mahasiswa ke dalam segmen perilaku tertentu. Namun, metode tersebut sering kali kurang optimal pada data pendidikan yang cenderung tidak seimbang (*imbalanced*) dan dinamis. Klusterisasi pada data yang tidak seimbang membuat kelompok kecil misalnya mahasiswa berisiko tinggi cenderung tenggelam dalam mayoritas data seperti klusterisasi menggunakan K-Means sensitif ke *centroid* dan mengabaikan kluster kecil serta *outlier* (Shaikh, et al, 2024 dan Chen, 2025).

Ouassif et al (2024) menemukan bahwa K-means adalah klusterisasi efektif untuk memetakan heterogenitas perilaku belajar siswa menjadi beberapa profil yang bermakna sehingga intervensi bisa lebih presisi. Studi pada konteks LMS menjadi beberapa kluster yang berbeda berdasarkan keterlibatan yang rendah hingga sangat tinggi. Hasil klusterisasi kemudian diterjemahkan menjadi rekomendasi strategi

pembelajaran yang dipersonalisasi sesuai dengan kebutuhan tiap klaster yang sudah terbentuk. Studi lain dari Jin (2025) menekankan pentingnya data multidimensi (waktu belajar, partisipasi kursus/aktivitas, diskusi online, interaksi, proyek kelompok, dsb.) harus mampu melakukan *feature selection* yang tepat sebelum klasterisasi dapat menemukan 4 kelompok perilaku (aktif, pasif, sosial-terlibat, fokus akademik) dengan kualitas klaster yang tinggi sehingga relevan untuk perancangan intervensi yang lebih tepat.

Twin SVM menawarkan solusi yang lebih efisien dalam menangani permasalahan data yang besar dan tidak seimbang dengan membangun dua *hyperplane non-parallel* yang mampu memisahkan data secara lebih stabil. Penelitian oleh Nasiri et al (2019) menunjukkan Twin SVM mampu meningkatkan akurasi klasifikasi perilaku mahasiswa, khususnya dalam konteks data LMS yang distribusinya tidak seimbang. Dengan demikian, penerapan Twin SVM pada segmentasi perilaku mahasiswa di LMS diharapkan dapat menghasilkan kelompok mahasiswa yang lebih representatif dan relevan untuk pengembangan strategi pembelajaran adaptif.

Integrasi Twin SVM dalam analisis *learning analytics* memperkuat proses pengambilan keputusan berbasis data di institusi pendidikan tinggi. Studi-studi mutakhir juga menyoroti pentingnya validasi pedagogis terhadap hasil segmentasi, guna memastikan bahwa hasil pengelompokan benar-benar bermanfaat untuk intervensi pembelajaran yang lebih efektif dan personal (Cantabella, et al, 2019). Dengan pendekatan ini, penelitian segmentasi mahasiswa berbasis perilaku pada LMS dengan Twin SVM memberikan kontribusi signifikan dalam mendukung transformasi digital pendidikan di pendidikan tinggi.

2.1.1 Analisis Perilaku Mahasiswa di LMS

Romero dan Ventura (2020) menyatakan bahwa data log aktivitas LMS dapat dimanfaatkan untuk mengidentifikasi pola perilaku belajar mahasiswa yang berkorelasi dengan performa akademik. Indikator seperti frekuensi login, partisipasi forum, dan ketepatan pengumpulan tugas terbukti memiliki hubungan yang signifikan dengan

keberhasilan belajar mahasiswa. Temuan ini memperkuat pentingnya analisis perilaku mahasiswa berbasis data aktivitas LMS.

Henrie et al (2015) juga menegaskan bahwa perilaku keterlibatan (*engagement*) mahasiswa dalam pembelajaran daring dapat diukur melalui aktivitas digital yang tercatat dalam LMS, sehingga analisis perilaku menjadi komponen penting dalam evaluasi pembelajaran daring.

Cenka et al (2022) menemukan berdasarkan analisis log aktivitas LMS, perilaku mahasiswa cenderung terkonsentrasi pada fitur inti seperti akses materi atau konten, pengerjaan dan pengumpulan tugas, dan aktivitas forum, jumlah akses LMS juga paling sering hanya terjadi pada hari perkuliahan, dan intensitas aktivitas mahasiswa menurun ketika instruksi dosen berkurang (misalnya pasca ujian tengah) sehingga menunjukkan kuatnya pengaruh strategi instruksional terhadap keterlibatan mahasiswa. Sedangkan temuan lain dari proses *mining* data LMS menunjukkan bahwa mahasiswa berprestasi tinggi lebih sering mengakses LMS dan memiliki alur proses belajar yang lebih kompleks dibanding mahasiswa lain. Strategi mengajar yang sistematis dan beragam berdampak kuat pada perilaku belajar dan dapat dipakai untuk membuat estimasi performa serta penyelesaian mata kuliah.

Studi lain pada LMS Moodle oleh Estacio et al (2017) menunjukkan bahwa *action logs* dapat diagregasi menjadi satu nilai numerik untuk memvisualisasikan level aktivitas mahasiswa, tetapi hubungan antara level aktivitas dan nilai akhir bersifat sangat bervariasi antar mata kuliah. Bahkan, ada indikasi bahwa lebih lama atau seringnya aktivitas di LMS Moodle tidak selalu berbanding lurus dengan capaian yang kemungkinan dipengaruhi perbedaan pendekatan pengajaran dan faktor penilaian di luar aktivitas *online*. Visualisasi *log* ini berguna untuk mengidentifikasi kebutuhan intervensi pedagogis yang lebih tepat.

Cantabella et al (2019) dan Lin (2025) menegaskan bahwa *log* aktivitas mahasiswa di dalam LMS dapat dipakai untuk membaca perilaku belajar mahasiswa bahkan dengan menerapkan fokus analitik yang berbeda. Pemanfaatan *log* aktivitas ini memberikan peluang besar bagi institusi pendidikan untuk memperoleh gambaran

komprehensif mengenai dinamika pembelajaran mahasiswa. Melalui analisis *log* yang terstruktur, kampus dapat mengidentifikasi pola interaksi, tingkat keterlibatan, hingga potensi permasalahan yang mungkin dihadapi oleh mahasiswa dalam proses pembelajaran daring.

Analisis perilaku mahasiswa di LMS memungkinkan institusi pendidikan untuk memetakan kebutuhan dan tantangan belajar yang dihadapi mahasiswa secara lebih spesifik. Dengan memahami kecenderungan interaksi mahasiswa terhadap fitur-fitur LMS, seperti penggunaan forum diskusi, akses ke modul pembelajaran, serta pola pengumpulan tugas, pihak kampus dapat merancang intervensi yang lebih efektif dan personal. Pendekatan ini tidak hanya meningkatkan kualitas pembelajaran daring, tetapi juga mendorong terciptanya lingkungan belajar yang adaptif dan responsif terhadap dinamika mahasiswa di era digital.

2.1.2 Segmentasi Mahasiswa Berbasis Aktivitas pada LMS

Segmentasi mahasiswa bertujuan untuk mengelompokkan mahasiswa berdasarkan kesamaan pola perilaku belajar. Waspada et al (2019) menerapkan metode K-Means *clustering* untuk mengelompokkan mahasiswa berdasarkan aktivitas mereka di LMS Moodle. Hasil penelitian menunjukkan bahwa segmentasi dapat membantu dosen dalam memahami variasi perilaku mahasiswa, namun metode *clustering* murni memiliki keterbatasan dalam menangani data dengan distribusi tidak seimbang serta kesulitan dalam mengklasifikasikan data mahasiswa baru.

Costa et al (2022) dalam kajiannya menyatakan bahwa sebagian besar penelitian segmentasi dalam *learning analytics* masih mengandalkan metode klasterisasi, yang cenderung menghasilkan segmen yang tidak stabil dan sulit ditafsirkan secara pedagogis jika tidak dikombinasikan dengan pendekatan lain.

Sejumlah penelitian menunjukkan bahwa segmentasi mahasiswa berbasis aktivitas pada LMS dapat dilakukan secara efektif dengan pendekatan *unsupervised clustering* pada jejak digital yang merekam interaksi seperti kuis, tugas, forum atau diskusi, dan akses pada materi. Model *weekly engagement profiling* berbasis K-means

memanfaatkan indikator keterlibatan dari data aktivitas LMS selama beberapa minggu untuk memetakan mahasiswa ke profil *high moderate low engagement*, dan temuan pentingnya adalah keterlibatan bersifat dinamis (Alzahrani, et al, 2025). Profil yang terbentuk juga berasosiasi dengan capaian akademik sehingga klaster rendah dapat diperlakukan sebagai sinyal awal risiko.

Studi lain yang membandingkan beberapa algoritma klaster yaitu BIRCH, DBSCAN, dan GMM pada data *log* platform pendidikan menemukan struktur klaster yang konsisten dan memisahkan kelompok dengan rentang aktivitas yang beragam dibandingkan kelompok dengan pola aktivitas yang lebih sempit dan spesifik, sehingga berguna untuk memahami strategi belajar yang berbeda dan merancang dukungan yang lebih personal. Dari sisi kualitas klaster, evaluasi internal menunjukkan beberapa algoritma dapat menghasilkan pemisahan klaster yang lebih jelas pada data pendidikan berdimensi tinggi (Pecuchova, et al, 2024).

Bessadok et al (2023) menunjukkan bahwa segmentasi mahasiswa berbasis aktivitas pada LMS dapat dilakukan dengan pendekatan EDM menggunakan *log* aktivitas pada LMS Blackboard. Aktivitas digital mahasiswa diklaster menjadi tiga profil yang merepresentasikan perbedaan pola keterlibatan di LMS, lalu peneliti menguji keterkaitan profil tersebut dengan capaian akademik menggunakan uji korelasi dan menemukan hubungan yang signifikan secara statistik, sehingga profil perilaku di LMS berpengaruh terhadap performa. Analisis juga menegaskan bahwa mahasiswa dengan profil aktivitas yang berbeda masih dapat mencapai performa akademik yang mirip sehingga tinggi atau rendah aktivitas di LMS tidak selalu linier dengan nilai mahasiswa. Secara segmentasi ini membantu dosen memahami cara mahasiswa memanfaatkan konten LMS.

2.1.3 Twin Support Vector Machines dalam Analisis Perilaku Mahasiswa

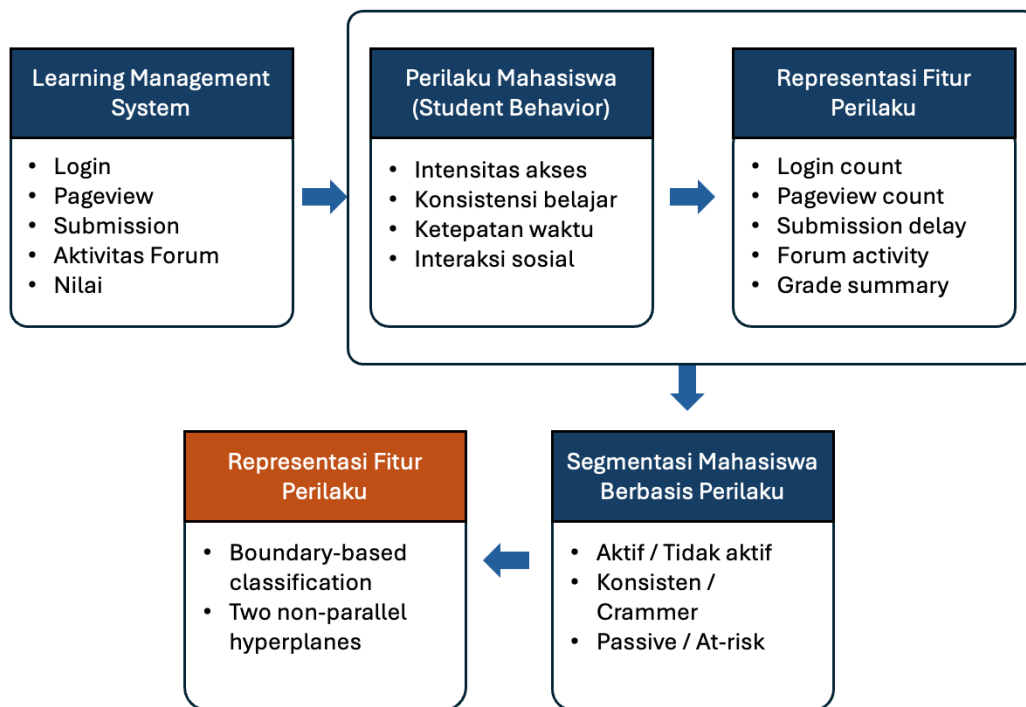
Twin SVM diperkenalkan oleh Jayadeva et al (2007) sebagai pengembangan dari Support Vector Machine (SVM) dengan membangun dua *hyperplane non-parallel*.

Pendekatan ini menghasilkan efisiensi komputasi yang lebih baik dan kinerja yang lebih stabil pada data dengan distribusi tidak seimbang.

Tanveer et al (2021) menunjukkan bahwa Twin SVM mampu menghasilkan performa klasifikasi yang lebih baik dibandingkan metode konvensional dalam memprediksi data yang bersifat *imbalanced* termasuk di dalamnya adalah perilaku dan performa mahasiswa, khususnya pada data pendidikan. Hal ini menunjukkan potensi Twin SVM untuk diterapkan dalam konteks segmentasi perilaku mahasiswa pada LMS. Tanveer et al (2025) melanjutkan penelitiannya memperoleh hasil bahwa kelebihan Twin SVM terutama ada pada cara belajar yang lebih efisien dan bentuk batas keputusan yang lebih adaptif dibanding SVM klasik. Twin SVM membangun dua *hyperplane non-parallel* yang masing-masing mendekat data dari satu kelas sambil menjaga jarak terhadap kelas lain, Sehingga pemisahan kelas bisa lebih fleksibel daripada satu *hyperplane* tunggal. Dari sisi komputasi, Twin SVM menyelesaikan dua masalah *Quadratic Programming Problem* (QPP) yang lebih kecil bukan satu QPP besar seperti SVM, sehingga proses pelatihannya dilaporkan mendekat empat kali lebih cepat daripada SVM.

Secara konseptual hubungan antara LMS dan Twin SVM adalah data perilaku mahasiswa di dalam LMS dilakukan segmentasi menggunakan Twin SVM. Hubungan konseptual tersebut dapat dilihat dalam grafik berikut.

SEKOLAH PASCASARJANA



Gambar 2. 1 Diagram Konseptual

LMS berfungsi sebagai lingkungan digital yang mencatat seluruh aktivitas pembelajaran mahasiswa secara terstruktur. Aktivitas tersebut, seperti login, akses materi, partisipasi forum, dan pengumpulan tugas, merepresentasikan interaksi mahasiswa dengan sistem pembelajaran daring. Aktivitas yang terekam dalam LMS mencerminkan perilaku belajar mahasiswa, yang meliputi intensitas keterlibatan, konsistensi belajar, ketepatan waktu, serta interaksi akademik. Perilaku ini tidak dapat diamati secara langsung, tetapi dapat direpresentasikan melalui indikator kuantitatif yang diekstraksi dari data log LMS.

Representasi perilaku mahasiswa selanjutnya dibentuk dalam bentuk fitur perilaku, yang digunakan sebagai dasar dalam proses segmentasi mahasiswa berbasis perilaku. Segmentasi bertujuan untuk mengelompokkan mahasiswa ke dalam segmen-segmen dengan karakteristik perilaku belajar yang relatif homogen.

Dalam penelitian ini, segmentasi mahasiswa dipandang sebagai permasalahan klasifikasi perilaku, sehingga digunakan metode *Twin Support Vector Machines* (Twin SVM) sebagai pendekatan pemodelan batas antar segmen. Twin SVM dipilih karena kemampuannya membangun dua *hyperplane non-parallel* yang lebih fleksibel dalam memodelkan perbedaan perilaku antar segmen.

2.1.4 Gap Penelitian

Berdasarkan tinjauan pustaka yang telah dilakukan, dapat disimpulkan bahwa:

1. Segmentasi mahasiswa berbasis data aktivitas LMS memiliki peran penting dalam mendukung pembelajaran adaptif.
2. Metode klatering murni masih mendominasi penelitian segmentasi LMS, meskipun memiliki keterbatasan.
3. Penerapan Twin SVM dalam segmentasi perilaku mahasiswa pada LMS masih relatif terbatas, khususnya dalam konteks pendidikan tinggi di Indonesia.

Oleh karena itu, penelitian ini berupaya mengisi celah penelitian dengan mengusulkan pendekatan segmentasi perilaku mahasiswa berbasis klasifikasi menggunakan Twin SVM, yang diharapkan mampu menghasilkan segmen mahasiswa yang lebih stabil, *interpretable*, dan relevan secara pedagogis.

Pemilihan Twin SVM sebagai metode segmentasi dalam penelitian ini didasarkan pada keunggulannya dalam membangun dua *hyperplane non-parallel* yang lebih fleksibel dan adaptif untuk memisahkan data dengan karakteristik perilaku yang kompleks. Menurut Kumar dan Gopal (2009), Twin SVM mampu menghasilkan klasifikasi yang lebih efisien dan akurat dibandingkan SVM konvensional, terutama pada data yang memiliki distribusi tidak linier. Selain itu, Twin SVM memberikan peningkatan pada stabilitas komputasi dan efisiensi waktu proses, sehingga sangat sesuai untuk analisis data *log* LMS yang bersifat besar dan dinamis.

Keunggulan tersebut diperkuat oleh hasil studi sebelumnya yang menunjukkan bahwa Twin SVM dapat mengatasi keterbatasan metode *clustering* murni, khususnya dalam konteks segmentasi berbasis perilaku mahasiswa, seperti yang diungkapkan oleh

Jayadeva et al (2007). Dengan demikian, penggunaan Twin SVM pada penelitian ini diharapkan mampu memberikan model segmentasi yang lebih representatif dan relevan dengan kebutuhan pembelajaran adaptif di lingkungan LMS.

2.2 Dasar Teori

2.2.1 Learning Management System (LMS)

Learning Management System (LMS) merupakan platform digital yang digunakan untuk mengelola, menyampaikan, dan mengevaluasi proses pembelajaran. LMS mencatat berbagai aktivitas pengguna, seperti login, akses materi, diskusi forum, pengumpulan tugas, dan asesmen, yang dapat dimanfaatkan untuk analisis pembelajaran (*learning analytics*). Berbagai jenis LMS digunakan di seluruh dunia, baik yang bersifat *open source* maupun komersial antara lain (1) Moodle adalah LMS *open source* yang sangat banyak digunakan di institusi pendidikan dan pelatihan. Moodle dikenal karena fleksibilitasnya, komunitas pengembang yang besar, serta dukungan untuk berbagai *plugin* sehingga mudah disesuaikan dengan kebutuhan pengguna. (2) Canvas merupakan LMS yang dikembangkan oleh Instructure dan banyak diadopsi oleh universitas di Amerika Serikat maupun negara lain. Canvas menawarkan antarmuka pengguna yang modern, fitur kolaborasi, serta integrasi yang baik dengan aplikasi pihak ketiga. (3) Blackboard adalah salah satu LMS komersial tertua dan paling banyak digunakan di dunia. Blackboard menyediakan fitur manajemen pembelajaran lengkap, analitik pembelajaran, dan sistem keamanan yang solid untuk institusi pendidikan menengah hingga perguruan tinggi. (4) Google Classroom merupakan LMS berbasis *cloud* yang terintegrasi dengan ekosistem Google. Google Classroom populer di kalangan sekolah dasar dan menengah karena kemudahan penggunaan, akses gratis, serta integrasi dengan Google Drive dan Google Docs. Setiap LMS memiliki karakteristik, kelebihan, dan kekurangan tersendiri, sehingga pemilihannya biasanya disesuaikan dengan kebutuhan institusi, skala pengguna, serta fitur yang diinginkan.

Moodle merupakan salah satu LMS yang paling populer dan banyak digunakan di dunia pendidikan karena sifatnya yang *open source* dan fleksibel. Keunggulan Moodle terletak pada struktur basis data yang terbuka serta kemampuannya untuk mencatat aktivitas pengguna secara detail, mulai dari *log* masuk, akses materi, hingga partisipasi dalam forum diskusi dan penilaian. Hal ini menjadikan Moodle sangat ramah untuk keperluan *data mining*, karena data aktivitas yang terekam dapat diekstraksi dan dianalisis dengan mudah untuk mendukung penelitian perilaku pembelajaran, prediksi performa mahasiswa, serta pengembangan sistem pembelajaran adaptif sesuai kebutuhan institusi pendidikan. Selain itu, Moodle mendukung berbagai *plugin* analitik yang dapat diintegrasikan langsung untuk mengoptimalkan proses *data mining* tanpa harus memodifikasi sistem utama secara kompleks.

2.2.2 Perilaku Mahasiswa dalam LMS

Perilaku mahasiswa dalam LMS mencakup aktivitas berupa (1) frekuensi akses baik secara harian atau mingguan, (2) kerapatan jarak antar login, (3) jumlah interaksi di forum diskusi, (4) ketepatan waktu pengumpulan tugas, (5) partisipasi dalam kuis, forum, dan asesmen, serta (6) jumlah tugas yang selesai dikumpulkan. Variabel-variabel tersebut merupakan indikator keterlibatan (*engagement*) yang penting dalam *e-learning* (Henrie, et al, 2015). Dengan kata lain, semakin tinggi frekuensi akses dan interaksi mahasiswa di dalam platform LMS, semakin besar pula kemungkinan mereka memahami dan mengikuti perkembangan pembelajaran yang disediakan secara bertahap. Frekuensi login yang konsisten menunjukkan aktivitas dan komitmen mahasiswa dalam mengikuti perkuliahan, sementara keterlibatan dalam forum dan diskusi mencerminkan sikap proaktif dalam bertukar pandangan, mengajukan pertanyaan, dan memberikan umpan balik kepada teman sekelas dan dosen. Ketepatan waktu dalam pengumpulan tugas mencerminkan disiplin diri dan manajemen waktu yang baik, sementara partisipasi dalam kuis dan penilaian lainnya menunjukkan tekad mahasiswa untuk mengevaluasi dan meningkatkan pemahaman mereka tentang konsep yang diajarkan. Jumlah tugas yang berhasil diselesaikan tidak hanya berfungsi sebagai

tolok ukur prestasi akademik, tetapi juga memberikan wawasan tentang ketahanan dan motivasi mahasiswa dalam menyelesaikan suatu mata kuliah. Semua aspek ini, jika dianalisis secara komprehensif, dapat membantu lembaga pendidikan mengidentifikasi tingkat keterlibatan mahasiswa dan merancang intervensi yang tepat untuk meningkatkan kualitas pembelajaran secara adaptif dan inklusif.

Learning analytics di dalam Moodle yang hanya menggunakan *simple rule* pada dasarnya memanfaatkan data aktivitas dasar, dengan *simple rule*, institusi dapat menetapkan aturan sederhana, misalnya mahasiswa dikategorikan aktif jika login lebih dari tiga kali seminggu, atau dianggap berisiko jika tidak mengakses materi selama tujuh hari berturut-turut. Pendekatan ini tidak membutuhkan algoritma kompleks, melainkan hanya mengandalkan logika *if-then* sederhana untuk mengevaluasi perilaku mahasiswa berdasarkan indikator dasar yang tersedia. Hasil analisis ini dapat langsung digunakan untuk mengidentifikasi mahasiswa yang membutuhkan intervensi, seperti pengingat atau dukungan tambahan, sehingga proses monitoring menjadi efisien dan mudah diterapkan di lingkungan pendidikan yang skalanya besar.

2.2.3 Pseudo-Labeling Klasterisasi

Klasterisasi sebagai sumber *pseudo-label* untuk melatih model klasifikasi *supervised* dan *semisupervised* terbukti efektif di banyak domain dengan syarat kualitas dan pemilihan *pseudo-label* dikelola dengan hati-hati untuk meminimalkan *noise* dan *error propagation*. *Pseudo labeling* berbasis klasterisasi adalah teknik yang mengubah hasil klasterisasi data tak berlabel menjadi label semu (*pseudo-label*) yang kemudian digunakan untuk melatih model klasifikasi secara *supervised* atau *semi-supervised*. Pendekatan ini sangat berguna ketika data berlabel sangat terbatas, seperti pada klasifikasi citra medis, *hyperspectral*, teks, atau data graf. Berbagai studi menunjukkan bahwa *pseudo-label* dari klasterisasi dapat secara signifikan meningkatkan akurasi model *supervised*, bahkan mendekati performa *fully-supervised* jika strategi seleksi label dan penanganan *noise* dilakukan dengan baik (Surbakti, et al, 2025) dan (Shahade, et al, 2025).

Teknik *pseudo-label* juga banyak digunakan dalam domain *few-shot learning*, *domain adaptation*, *federated learning*, dan *continual learning* (Wang, et al, 2025). Tantangan utama adalah kualitas *pseudo-label* yang dihasilkan, karena label yang salah dapat menurunkan performa model. Oleh karena itu, banyak penelitian mengembangkan metode seleksi label, *threshold adaptif*, *curriculum learning*, dan *loss* khusus untuk mengurangi dampak *noise*.

Fu et al (2024) memposisikan *pseudo-labeling* sebagai strategi *self-supervised* berbasis label klaster dengan *node* terlebih dulu diklaster menggunakan K-means untuk menghasilkan klasterisasi label yang kemudian dipakai sebagai *pseudo-label* untuk melatih model. Namun dalam termuannya ditegaskan bahwa memakai label klaster secara langsung berisiko karena reliabilitas label klaster tidak selalu baik sehingga dapat membawa *noise* ke proses belajar. Untuk menguatkan kualitas *pseudo-label*, perlu dilakukan *selection enhancement*, menghitung *confident score* dan menyaring *node* yang kurang terpercaya agar *pseudo-label* yang dipakai untuk tugas *self-supervised* lebih representatif.

2.2.4 Segmentasi Berbasis Klasifikasi

Segmentasi perilaku mahasiswa dalam penelitian ini dipandang sebagai permasalahan klasifikasi, di mana mahasiswa dikelompokkan ke dalam segmen perilaku tertentu berdasarkan karakteristik aktivitas mereka. Pendekatan klasifikasi memungkinkan pemodelan batas antar segmen secara eksplisit dan mendukung penerapan segmentasi pada data mahasiswa baru. Menurut Henrie et al (2015) bahwa keterlibatan siswa dalam lingkungan pembelajaran yang dimediasi teknologi dapat diukur menggunakan berbagai indikator perilaku yang dicatat oleh LMS. Indikator-indikator yang tercatat dalam *log* di dalam LMS dapat digunakan untuk mengelompokkan siswa ke dalam kelompok perilaku yang berbeda melalui pendekatan klasifikasi, sehingga memungkinkan intervensi yang tepat sasaran dan peningkatan hasil pembelajaran.

Dalam penelitian Luo, et al. (2024), segmentasi dilakukan melalui pendekatan klusterisasi lalu klasifikasi berbasis pada jejak aktivitas mahasiswa di LMS dengan skala besar. Pertama, data perilaku online mahasiswa dilakukan klusterisasi menggunakan algoritma *Expectation–Maximization* (EM) dan menghasilkan 5 profil perilaku yaitu *inactive*, *low-activity*, *assignment-focused*, *video-focused*, dan *high-activity*. Selanjutnya, level segmentasi dinaikkan dari mahasiswa ke level mata kuliah yaitu setiap *blended course* diklasifikasikan berdasarkan kluster mahasiswa yang proporsinya paling dominan di mata kuliah tersebut, sehingga terbentuk beberapa tipe perilaku kuliah. Setelah *course-type* terbentuk, peneliti membangun model klasifikasi prediksi hasil belajar dari klusterisasi menggunakan Random Forest dan menunjukkan bahwa akurasi prediksi meningkat ketika prediksi dilakukan per kategori. Implikasi dari segmentasi tersebut adalah kombinasi profil perilaku dengan tipe mata kuliah memungkinkan sistem melakukan intervensi lebih terarah dan penelitian ini juga menegaskan bahwa tidak ada satu indikator perilaku tunggal yang selalu dominan. Akurasi yang baik muncul ketika mahasiswa terlibat pada berbagai jenis aktivitas di LMS.

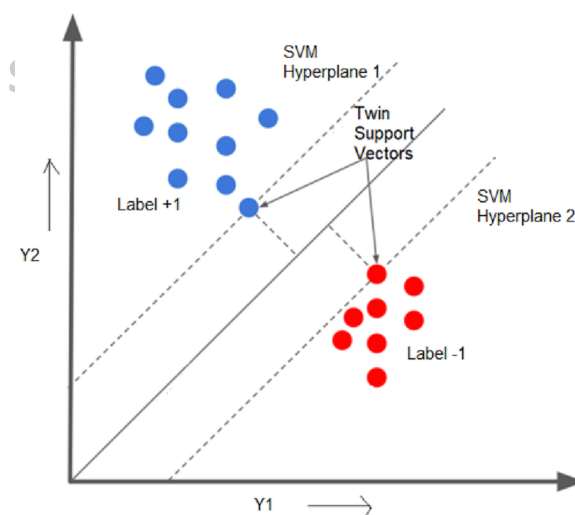
2.2.5 Twin Support Vector Machines (Twin SVM)

Twin SVM adalah metode klasifikasi yang menggunakan dua *hyperplane*, masing-masing dekat dengan satu kelas dan sejauh mungkin dari kelas lainnya. Keunggulannya terletak pada efisiensi komputasi dan kinerjanya dalam menangani data tidak seimbang. TWIN SVM menyelesaikan dua masalah optimisasi kuadratik yang lebih kecil dibandingkan SVM konvensional, sehingga lebih cepat dalam pelatihan (Jayadeva, et al, 2007). Metode ini memanfaatkan konsep *Generalized Eigenvalues Proximal Support Vector Machine* (GEPSVM) dan menemukan dua bidang tidak sejajar untuk setiap kelas dengan menyelesaikan sepasang permasalahan *Quadratic Programming*. Twin SVM meningkatkan kecepatan komputasi dibandingkan dengan Support Vector Machine (SVM) tradisional.

Pada awalnya, Twin SVM dikembangkan untuk menyelesaikan permasalahan klasifikasi biner, selanjutnya para peneliti berhasil memperluas penerapannya ke domain permasalahan multi-kelas. Twin SVM secara konsisten memberikan hasil empiris yang menjanjikan, sehingga memiliki banyak fitur menarik yang meningkatkan tingkat aplikabilitasnya.

Konsep utama Twin SVM adalah membangun dua *hyperplane* secara simultan, di mana setiap *hyperplane* dirancang agar berada dekat dengan satu kelas (kelompok data) dan pada saat yang sama sejauh mungkin dari kelas lainnya. Berbeda dengan SVM tradisional yang hanya membangun satu *hyperplane* optimal untuk memisahkan dua kelas, Twin SVM membagi permasalahan menjadi dua submasalah, yang masing-masing menghasilkan satu *hyperplane* yang lebih sesuai untuk data yang tidak seimbang dan berskala besar.

Setiap *hyperplane* dalam Twin SVM dibentuk melalui proses optimasi kuadratik yang lebih kecil, sehingga memungkinkan model ini beroperasi secara lebih efisien dari segi waktu pelatihan dan sumber daya komputasi. Dengan pendekatan ini, Twin SVM dapat mengurangi waktu komputasi dan memberikan pemisahan kelas yang lebih baik dalam situasi di mana data tidak seimbang atau memiliki distribusi yang kompleks.



Gambar 2. 2 Ilustrasi Dua *Hyperplane* dalam Twin SVM

Twin SVM dan SVM tradisional sama-sama merupakan metode klasifikasi berbasis *hyperplane*, namun terdapat beberapa perbedaan utama di antara keduanya. SVM tradisional membangun satu *hyperplane* optimal untuk memisahkan dua kelas data, sedangkan Twin SVM membangun dua *hyperplane* secara simultan, masing-masing dekat dengan satu kelas dan sejauh mungkin dari kelas lainnya. Twin SVM menyelesaikan dua masalah optimasi kuadratik yang lebih kecil, sehingga proses pelatihan menjadi lebih cepat dan efisien, terutama pada data yang tidak seimbang atau berukuran besar. Sementara itu, SVM tradisional cenderung lebih berat secara komputasi karena hanya mengandalkan satu masalah optimasi utama. Dengan demikian, Twin SVM menawarkan keunggulan dalam hal kecepatan komputasi dan performa pada data tidak seimbang dibandingkan SVM konvensional.

Optimasi penggunaan Twin SVM dilakukan dengan cara membagi permasalahan klasifikasi menjadi dua submasalah optimasi kuadratik yang lebih kecil. Setiap submasalah bertujuan untuk mencari satu bidang pemisah (*hyperplane*) yang paling dekat dengan data dari satu kelas namun sejauh mungkin dari data kelas lainnya. Dengan demikian, tiap *hyperplane* dioptimalkan agar memberikan margin yang optimal khusus untuk kelasnya sendiri, bukan secara keseluruhan seperti pada SVM tradisional. Pendekatan ini tidak hanya mempercepat proses pelatihan, tetapi juga membuat model lebih adaptif terhadap distribusi data yang tidak seimbang, sehingga mampu menghasilkan segmentasi yang lebih akurat dan efisien dalam konteks data dunia nyata.

2.2.6 Implementasi Twin Support Vector Machines

Sebagai sebuah metode klasifikasi Twin SVM memiliki aturan dan struktur berupa notasi dasar, bentuk *hyperplane*, fungsi objektif, dan aturan klasifikasi.

1. Notasi Dasar

Misalnya tersedia *dataset* pelatihan:

$$\{(x_i, y_i)\}_{i=1}^N, x_i \in \mathbb{R}^d, y_i \in \{+1, -1\} \quad (2.1)$$

Data dipisahkan menjadi dua matriks:

[1] $A \in \mathbb{R}^{m \times d} \rightarrow$ data kelas +1

[2] $B \in \mathbb{R}^{n \times d} \rightarrow$ data kelas -1

dengan:

- 1) $\{(x_i, y_i)\}_{i=1}^N$ adalah himpunan pasangan data berlabel sebanyak N sampel
- 2) Indeks $i = 1, \dots, N$ artinya menomori sampel berlabel dari 1 sampai N
- 3) x_i adalah vektor fitur (input) untuk sampel ke- i
- 4) $\mathbb{R}^d$ adalah ruang vektor berdimensi d (artinya tiap x_i memiliki d fitur). Jadi secara konsep $x_i \in \mathbb{R}^d$
- 5) y_i adalah label kelas untuk sampel ke- i
- 6) $y_i \in \{+1, -1\}$ adalah masalah klasifikasi biner, dengan dua kelas yang dikodekan sebagai +1 dan -1
- 7) m dan n adalah jumlah data pada masing-masing kelas
- 8) d adalah jumlah fitur.

2. Bentuk Hyperplane pada Twin SVM

Twin SVM membangun dua *hyperplane non-paralel*:

$$w_1^T x + b_1 = 0 \text{ dan } w_2^T x + b_2 = 0 \quad (2.2)$$

Dengan:

- 1) x adalah vektor data fitur sebuah sampel input
- 2) w_1, w_2 adalah vektor bobot untuk *hyperplane* pertama dan kedua
- 3) w_k adalah vektor normal yang tegak lurus terhadap *hyperplane* ke- k
- 4) w_1^T, w_2^T adalah tanda transpose, dipakai supaya operasi $w_k^T x$ menjadi perkalian *dot product*.
- 5) $w_k^T x$ adalah *dot product* antara bobot dan fitur: nilai proyeksi x ke arah w_k

6) b_1, b_2 adalah bias untuk *hyperplane* pertama dan kedua, fungsinya menggeser posisi *hyperplane*

7) $= 0$ menyatakan himpunan titik x yang tepat berada di *hyperplane* tersebut.

Maka:

Jika $w_k^T x + b_k > 0$: titik berada di satu sisi *hyperplane* ke- k dan jika $w_k^T x + b_k < 0$: titik berada di sisi lainnya.

Formulasi di atas merupakan ciri khas Twin SVM yaitu membangun dua *hyperplane*:

- 1) *Hyperplane* 1 ($w_1^T x + b_1 = 0$) biasanya dibuat dekat dengan kelas +1 dan sejauh mungkin dari kelas -1.
- 2) *Hyperplane* 2 ($w_2^T x + b_2 = 0$) biasanya dibuat dekat dengan kelas -1 dan sejauh mungkin dari kelas +1.

3. Formulasi Optimasi TSVM

Hyperplane untuk Kelas +1

Twin SVM meminimalkan jarak data kelas +1 ke *hyperplane* pertama dan memaksimalkan jarak data kelas -1 dari *hyperplane* tersebut.

$$\min_{w_1, b_1, \xi} \frac{1}{2} \|Aw_1 + e_1 b_1\|^2 + c_1 \|\xi\|^2 \quad (2.3)$$

Dengan:

- 1) $\|Aw_1 + e_1 b_1\|^2$ adalah vektor $Aw_1 + e_1 b_1$ berisi nilai $(w_1^T x_i + b_1)$ untuk semua titik kelas A. Meminimalkan norm kuadratnya berarti: *hyperplane* 1 dibuat sedekat mungkin ke titik-titik kelas +1
- 2) ξ adalah vektor *slack* untuk sampel kelas B (ukuran m_2).
- 3) c_1 mengatur *trade-off* yaitu semakin besar c_1 maka semakin keras hukuman pelanggaran kendala (*slack* kecil).

Kendala:

$$Bw_1 + e_2 b_1 + \xi \geq e_2 \quad (2.4)$$

Ini ekuivalen per-baris (untuk tiap sampel x_j di kelas B):

$$w_1^T x_j + b_1 \geq 1 - \xi_j \quad (2.5)$$

Makna dari titik-titik kelas -1 (B) dipaksa berada di sisi tertentu *hyperplane* 1 dengan margin 1 tetapi boleh melanggar jika perlu lewat slack ξ_j (umumnya $\xi_j \geq 0$ secara implisit dalam formulasi standar). Intuisinya *hyperplane* 1 menempel pada kelas +1, sambil mendorong kelas -1 berada di sisi lain dengan margin.

Hyperplane untuk Kelas -1

$$\min_{w_2, b_2, \eta} \frac{1}{2} \|Bw_2 + e_2 b_2\|^2 + c_{12} \|\eta\|^2 \quad (2.6)$$

Dengan:

- 1) $\|Bw_2 + e_2 b_2\|^2$ memaksa *hyperplane* 2 dekat ke titik-titik kelas -1
- 2) η adalah slack untuk sampel kelas A (ukuran m_1).
- 3) c_2 adalah bobot penalti pelanggaran kendala untuk *hyperplane* kelas -1

Kendala:

$$Aw_2 + e_1 b_2 - \eta \leq -e_1 \quad (2.7)$$

Per baris (untuk tiap sampel x_i di kelas A):

$$w_2^T x_i + b_2 \leq -1 + \eta_i \quad (2.8)$$

Maknanya dari titik-titik kelas +1 (A) dipaksa berada di sisi lain *hyperplane* 2 dengan margin 1, tapi boleh melanggar via slack η_i (umumnya $\eta_i \geq 0$).

4. Aturan Klasifikasi (*Decision Rule*)

Untuk data uji x , jarak terhadap masing-masing hyperplane dihitung sebagai:

$$d_1(x) = \frac{|w_1^T x + b_1|}{\|w_1\|}, d_2(x) = \frac{|w_2^T x + b_2|}{\|w_2\|} \quad (2.9)$$

Keputusan kelas:

$$\hat{y}(x) = \begin{cases} +1, & \text{jika } d_1(x) < d_2(x) \\ -1, & \text{jika } d_2(x) < d_1(x) \end{cases}$$

Dengan:

- 1) $d_1(x)$ adalah arak (tegak lurus) dari titik x ke *hyperplane* 1
- 2) $d_2(x)$ adalah arak (tegak lurus) dari titik x ke *hyperplane* 2
- 3) $w_1^T x + b_1$ adalah nilai fungsi *hyperplane* 1 untuk titik x jika nilainya 0 maka x tepat di *hyperplane* 1 dan jika positif atau negatif, berarti x ada di salah satu sisi *hyperplane* 1
- 4) $w_2^T x + b_2$ adalah nilai fungsi *hyperplane* 2 untuk titik x jika nilainya 0 maka x tepat di *hyperplane* 2 dan jika positif atau negatif, berarti x ada di salah satu sisi *hyperplane* 2
- 5) Pembagi $\|w_1\|$ dan $\|w_2\|$ menormalkan supaya hasilnya benar-benar jarak geometris, jika tidak dibagi, nilainya hanya sekor bukan jarak.

5. Ringkasan Konseptual TSVM

Twin SVM membangun dua *hyperplane non-paralel* yang masing-masing merepresentasikan satu kelas. Dengan memecah permasalahan klasifikasi menjadi dua masalah optimasi kuadratik yang lebih kecil, Twin SVM menawarkan efisiensi komputasi yang lebih baik dibandingkan SVM klasik, serta fleksibilitas dalam memodelkan batas keputusan antar kelas. Twisn SVM merupakan pengembangan dari

metode SVM yang digunakan sebagai model klasifikasi, termasuk untuk permasalahan klasifikasi QPP. Berbeda dengan SVM klasik yang membangun satu *hyperplane* pemisah, TSVM membangun dua *hyperplane* secara simultan, masing-masing lebih dekat ke satu kelas dan sejauh mungkin dari kelas lain.

2.2.7 Evaluasi Model Segmentasi

Segmentasi mahasiswa berbasis perilaku dalam penelitian ini dimodelkan sebagai permasalahan klasifikasi, di mana setiap mahasiswa diklasifikasikan ke dalam segmen perilaku tertentu berdasarkan hasil pemodelan Twin SVM. Oleh karena itu, evaluasi kinerja model segmentasi dilakukan menggunakan pendekatan evaluasi model klasifikasi.

Salah satu metode evaluasi yang umum digunakan untuk menilai kinerja model klasifikasi adalah *confusion matrix*. *Confusion matrix* memberikan gambaran mengenai kesesuaian antara label aktual dan label hasil prediksi model, sehingga memungkinkan analisis kesalahan klasifikasi secara lebih rinci. Confusion matrix merupakan tabel yang menyajikan jumlah prediksi benar dan salah dari suatu model klasifikasi. Pada kasus klasifikasi biner, confusion matrix terdiri dari empat komponen utama, yaitu:

1. *True Positive* (TP): jumlah data yang diprediksi sebagai kelas positif dan benar-benar termasuk kelas positif.
2. *True Negative* (TN): jumlah data yang diprediksi sebagai kelas negatif dan benar-benar termasuk kelas negatif.
3. *False Positive* (FP): jumlah data yang diprediksi sebagai kelas positif tetapi sebenarnya termasuk kelas negatif.
4. *False Negative* (FN): jumlah data yang diprediksi sebagai kelas negatif tetapi sebenarnya termasuk kelas positif.

Struktur umum *confusion matrix* untuk klasifikasi biner ditunjukkan pada tabel berikut.

Tabel 2. 1 Klasifikasi Biner

Aktual \ Prediksi	Positif	Negatif
Positif	TP	FN
Negatif	FP	TN

2.2.8 Metrik Evaluasi Berbasis *Confusion Matrix*

Berdasarkan *confusion matrix*, beberapa metrik evaluasi dapat dihitung untuk menilai kinerja model segmentasi, antara lain:

1. Akurasi (*Accuracy*)

Mengukur proporsi prediksi yang benar terhadap seluruh data.

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (2.10)$$

2. Presisi (*Precision*)

Mengukur tingkat ketepatan model dalam memprediksi kelas positif.

$$\text{Precision} = \frac{TP}{TP + FP} \quad (2.11)$$

3. Recall (*Sensitivity*)

Mengukur kemampuan model dalam mendeteksi seluruh data kelas positif.

$$\text{Recall} = \frac{TP}{TP + FN} \quad (2.12)$$

4. F1-Score

Merupakan rata-rata harmonik antara presisi dan recall.

$$\text{F1-Score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (2.13)$$

2.2.9 Peran Confusion Matrix dalam Evaluasi Hasil Segmentasi

Dalam konteks segmentasi mahasiswa berbasis perilaku, *confusion matrix* digunakan untuk mengevaluasi sejauh mana model Twin SVM dan LSTSVM mampu membedakan mahasiswa ke dalam segmen perilaku yang telah didefinisikan. Evaluasi ini penting terutama pada segmen-segmen tertentu, seperti mahasiswa berisiko akademik (*at-risk*), di mana kesalahan klasifikasi berupa *false negative* dapat berdampak signifikan terhadap efektivitas intervensi akademik.

Selain itu, *confusion matrix* memungkinkan analisis performa model pada setiap tahap klasifikasi biner dalam skema segmentasi hierarkis, misalnya pada pemisahan mahasiswa aktif dan tidak aktif, atau mahasiswa *passive* dan *at-risk*. Dengan demikian, evaluasi tidak hanya dilakukan secara global, tetapi juga secara spesifik pada masing-masing keputusan segmentasi.

Tabriz (2025) menjelaskan *confusion matrix* dipakai untuk menganalisis detail kesalahan klasifikasi (*misclassification*) pada model prediksi tingkat *engagement* dan performa akademik mahasiswa pada 3 segmen yaitu tinggi, sedang, dan rendah. *Confusion matrix* tidak hanya menunjukkan akurasi, tetapi juga mengukur seberapa sering model salah menempatkan kelas pada tempat yang tidak sesuai. Dalam evaluasi hasil model klasifikasi *confusion matrix* mampu menunjukkan kelas mana yang paling sulit dibedakan dan membantu memahami pola *overlap* antar kelas dan kelas transisi yang tidak stabil.

SEKOLAH PASCASARJANA

2.2.10 *Balanced Accuracy*

Balanced Accuracy merupakan metrik evaluasi yang dirancang untuk mengukur performa model klasifikasi pada data yang tidak seimbang (*imbalanced data*). Berbeda dengan akurasi tradisional yang hanya menghitung proporsi prediksi benar dari total keseluruhan data, *balanced accuracy* memperhitungkan sensitivitas (*recall*) pada masing-masing kelas, sehingga memberikan bobot yang adil antara kelas mayoritas dan minoritas. Dengan demikian, *balanced accuracy* sangat relevan ketika distribusi

data antar kelas tidak seimbang, seperti pada kasus segmentasi perilaku mahasiswa di mana jumlah mahasiswa dalam masing-masing segmen bisa sangat berbeda.

Secara matematis, *balanced accuracy* dihitung sebagai rata-rata dari *recall* pada setiap kelas.

$$\text{Balanced Accuracy} = \frac{\text{Recall Kelas 1} + \dots + \text{Recall Kelas } n}{n} \quad (2.14)$$

Pendekatan ini membantu menghindari bias evaluasi yang biasanya terjadi jika hanya menggunakan akurasi biasa, yang cenderung dipengaruhi oleh kelas yang jumlahnya lebih banyak. Dalam konteks segmentasi mahasiswa, *balanced accuracy* memastikan bahwa performa model dalam mengidentifikasi mahasiswa pada kelas minoritas, seperti kelompok *at-risk*, tetap terukur dengan baik dan tidak terabaikan oleh dominasi kelas mayoritas.

Penggunaan *balanced accuracy* sangat penting dalam penelitian yang fokus pada deteksi kelompok minoritas atau kelas yang memiliki dampak signifikan terhadap pengambilan keputusan, seperti mahasiswa dengan risiko akademik. Dengan *balanced accuracy*, model yang dihasilkan akan lebih adil dalam menilai performa pada seluruh kelas, sehingga intervensi atau kebijakan berbasis hasil segmentasi dapat diterapkan secara lebih tepat sasaran. Oleh karena itu, *balanced accuracy* menjadi salah satu metrik utama yang wajib dipertimbangkan dalam evaluasi model klasifikasi pada data yang tidak seimbang dalam ranah pendidikan maupun bidang lain yang serupa.

2.2.11 Cohen's Kappa

Cohen's Kappa adalah sebuah metrik statistik yang digunakan untuk mengukur tingkat kesepakatan antara dua penilai atau pengklasifikasi terhadap kategori yang sama, dengan memperhitungkan kemungkinan kesepakatan yang terjadi secara kebetulan. Tidak seperti akurasi sederhana yang hanya menghitung proporsi kesepakatan dari total klasifikasi, Cohen's Kappa memberikan penyesuaian terhadap

peluang kesepakatan acak, sehingga hasil yang diperoleh merefleksikan tingkat konsistensi yang lebih objektif antara penilai.

Nilai *Kappa* berkisar dari -1 hingga 1, di mana nilai 1 menunjukkan kesepakatan sempurna, 0 berarti kesepakatan yang terjadi hanya sebesar peluang acak, dan nilai negatif menunjukkan adanya ketidaksepakatan yang lebih besar daripada yang diharapkan secara kebetulan. Interpretasi tingkat kesepakatan umumnya mengikuti kategori berikut, kurang dari 0 (tidak ada kesepakatan), 0–0,20 (kesepakatan sangat lemah), 0,21–0,40 (lemah), 0,41–0,60 (moderat), 0,61–0,80 (kuat), dan 0,81–1,00 (sangat kuat). Dengan demikian, Cohen's Kappa tidak hanya sekadar menghitung jumlah prediksi yang benar, namun juga memberikan gambaran kualitas kesepakatan yang lebih mendalam.

Penggunaan Cohen's Kappa sangat cocok dalam evaluasi model klasifikasi, terutama pada kasus di mana distribusi kelas tidak seimbang dan ketika pengambilan keputusan dilakukan oleh lebih dari satu penilai. Dalam konteks segmentasi perilaku mahasiswa, Cohen's Kappa membantu memastikan bahwa model atau penilai mampu mengklasifikasikan mahasiswa ke dalam segmen yang tepat secara konsisten, tanpa bias akibat dominasi kelas tertentu. Hal ini penting agar intervensi atau tindak lanjut berbasis hasil segmentasi benar-benar akurat dan adil bagi setiap kelompok mahasiswa.

Secara matematis, rumus Cohen's Kappa dapat dituliskan sebagai berikut:

SEKOLAH PASCASARJANA

$$Kappa = \frac{Po - P}{1 - Pe} \quad (2.15)$$

Dengan:

Po adalah proporsi kesepakatan yang diamati antara dua penilai,

Pe adalah proporsi kesepakatan yang diharapkan terjadi secara kebetulan.

Dengan menggunakan rumus ini, tingkat kesepakatan yang dihasilkan menjadi lebih objektif karena memperhitungkan peluang acak dalam proses klasifikasi.

2.2.12 Matthews Correlation Coefficient (MCC)

Matthews Correlation Coefficient (MCC) adalah salah satu metrik evaluasi performa model klasifikasi yang memperhitungkan keempat elemen pada confusion matrix, yaitu *True Positive* (TP), *True Negative* (TN), *False Positive* (FP), dan *False Negative* (FN). MCC menghasilkan nilai antara -1 hingga 1, di mana nilai 1 menunjukkan klasifikasi sempurna, 0 berarti performa model setara dengan tebakan acak, dan -1 menunjukkan klasifikasi yang sepenuhnya salah. Metrik ini dikembangkan oleh Brian W. Matthews pada tahun 1975 dan kini banyak digunakan dalam berbagai bidang, khususnya ketika data yang dianalisis memiliki distribusi kelas yang tidak seimbang.

Keunggulan utama MCC terletak pada kemampuannya memberikan penilaian yang adil bahkan ketika proporsi kelas mayoritas dan minoritas sangat timpang. Berbeda dengan akurasi biasa yang dapat menyesatkan pada data tidak seimbang, MCC tetap memperhitungkan kesalahan klasifikasi pada kedua kelas secara proporsional. Dengan demikian, MCC mampu menggambarkan kualitas model secara menyeluruh tanpa bias dominasi kelas tertentu, sehingga sangat relevan untuk aplikasi di dunia nyata seperti deteksi kelompok minoritas atau penilaian risiko.

Bidang pendidikan khususnya dalam segmentasi perilaku mahasiswa atau deteksi kelompok *at-risk*, MCC sangat bermanfaat karena mampu mengukur seberapa baik model dalam membedakan mahasiswa yang benar-benar membutuhkan intervensi dengan yang tidak. Metrik ini memastikan bahwa evaluasi performa model tidak hanya mengandalkan keberhasilan pada kelompok mayoritas, melainkan juga memperhatikan kemampuan model dalam mendeteksi kelompok minoritas yang seringkali menjadi fokus utama dalam pengambilan keputusan berbasis data. Selain itu, MCC juga sering digunakan dalam penelitian biomedis, bioinformatika, serta berbagai studi ilmiah lain yang melibatkan data klasifikasi dengan distribusi tak seimbang. Dengan karakteristiknya yang robust terhadap ketimpangan kelas, MCC menjadi pilihan utama bagi peneliti dan praktisi yang membutuhkan evaluasi objektif dan komprehensif atas

performa model klasifikasi mereka. Oleh karena itu, memahami definisi dan keunggulan MCC sangat penting untuk memastikan kualitas dan keadilan dalam proses evaluasi model, baik dalam konteks akademik maupun industri.

Secara matematis, rumus Matthews Correlation Coefficient (MCC) dapat dituliskan sebagai berikut:

$$MCC = \frac{TP \times TN - FP \times FN}{\sqrt{(TP + FP)(TP + FN)(TN + FP)(TN + FN)}} \quad (2.16)$$

dengan:

- 1) TP adalah *True Positive*
- 2) TN adalah *True Negative*
- 3) FP adalah *False Positive*
- 4) FN adalah *False Negative*.

Rumus ini memastikan bahwa seluruh elemen *confusion matrix* digunakan dalam perhitungan, sehingga menghasilkan penilaian yang lebih komprehensif terhadap performa model klasifikasi.

2.2.13 Principal Component Analysis

Principal Component Analysis (PCA) adalah suatu teknik statistik yang digunakan untuk mereduksi dimensi data dengan cara mengubah variabel-variabel asli yang saling berkorelasi menjadi sejumlah variabel baru yang disebut *principal components*. Komponen utama ini merupakan kombinasi linier dari variabel-variabel awal dan disusun sedemikian rupa sehingga komponen pertama memuat variasi terbesar, diikuti oleh komponen berikutnya secara berturut-turut. PCA berfungsi untuk menyederhanakan data berdimensi tinggi tanpa kehilangan informasi penting, sehingga dapat mempermudah proses analisis dan interpretasi data. Selain itu, PCA juga membantu dalam mengurangi *noise* atau redundansi antar variabel, serta meningkatkan

efisiensi komputasi dalam proses analisis selanjutnya, seperti klasterisasi atau klasifikasi.

PCA sangat bermanfaat dalam memvisualisasikan hasil klasterisasi dan klasifikasi, khususnya ketika data memiliki banyak variabel. Dengan mereduksi data ke dalam dua atau tiga *principal components* utama, kita dapat memproyeksikan data ke dalam ruang berdimensi rendah yang mudah divisualisasikan. Hasil klaster yang terbentuk dapat digambarkan secara visual, sehingga pola-pola antar klaster maupun keterpisahan antar kelompok menjadi lebih jelas terlihat. Visualisasi ini sangat membantu dalam evaluasi hasil klasterisasi dan dalam komunikasi hasil analisis kepada pihak lain secara lebih intuitif dan informatif.

Dalam konteks proyek *Hierarchical Student Segmentation*, PCA digunakan untuk:

1. Visualisasi 2D untuk mereduksi fitur multidimensi menjadi 2 komponen untuk *scatter plot*,
2. Interpretasi klaster untuk memvisualisasikan pemisahan antar segmen mahasiswa, dan
3. *Validasi Decision Boundary* untuk menampilkan *hyperplane* Twin SVM dalam ruang 2D.

Rumus utama dalam PCA adalah proses dekomposisi matriks kovarians atau korelasi dari data, yang dilakukan melalui *eigen decomposition*. Secara matematis, rumus PCA dapat dituliskan sebagai berikut:

Diberikan matriks data:

$$X \in \mathbb{R}^{(n \times p)} \quad (2.17)$$

dengan n sampel dan p fitur:

Langkah 1: Standarisasi Data

$$X_{std} = (X - \mu) / \sigma \quad (2.18)$$

dimana μ adalah vektor mean dan σ adalah vektor standar deviasi.

Langkah 2: Menghitung Matriks Kovarians

$$C = \frac{1}{n-1} X_{\text{std}}^T X_{\text{std}} \quad (2.19)$$

Matriks kovarians $C \in \mathbb{R}^{(p \times p)}$ menangkap hubungan linear antar fitur.

Langkah 3: Dekomposisi *Eigen*

$$C v_i = \lambda_i v_i \quad (2.20)$$

dengan:

1. λ_i adalah eigenvalue ke-i (variansi yang dijelaskan oleh komponen ke-i)
2. v_i adalah eigenvector ke-i (arah komponen utama ke-i)

Eigenvalue diurutkan secara menurun: $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_p$

Langkah 4: Proyeksi ke Komponen Utama

$$Z = X_{\text{std}} x V_K \quad (2.21)$$

dimana $V_K = [V_1, V_2, \dots, V_K]$ adalah matriks yang berisi k eigenvector teratas.

Untuk visualisasi 2D, $k = 2$, sehingga:

$$Z = X_{\text{std}} x [V_1 \ V_2] \quad (2.21)$$

$$PC_1 = \sum_{j=1}^p w_{1j} x_j$$

Kontribusi fitur ke *Principal Components* adalah setiap *Principal Component* merupakan kombinasi linear dari fitur-fitur asli:

$$PC_1 = \sum_{j=1}^p w_{1j} x_j \quad (2.22)$$

$$PC_1 = \sum_{j=1}^p w_{1j} x_j \quad (2.22)$$

dengan w_{1j} adalah bobot fitur j pada komponen i , yang merupakan elemen dari eigenvector V_1 .

Secara langkah kerja inti dari PCA meliputi: 1) melakukan normalisasi data, 2) menghitung matriks kovarians, 3) mencari nilai dan vektor *eigen*, lalu 4) memilih sejumlah komponen utama berdasarkan nilai *eigen* terbesar untuk merepresentasikan data secara lebih ringkas tanpa kehilangan informasi penting. Dengan demikian, PCA mengubah data berdimensi tinggi menjadi beberapa komponen utama yang saling tidak berkorelasi.



SEKOLAH PASCASARJANA