

MODEL PENGESANAN MASALAH KESIHATAN MENTAL MENGUNAKAN PENDEKATAN PEMBELAJARAN MESIN TRADISIONAL DAN MENDALAM

¹Vilaasini A/P Kumar, ¹Assoc. Prof. Dr. Sabrina Tiun

¹Fakulti Teknologi & Sains Maklumat
43600 Universiti Kebangsaan Malaysia

ABSTRAK

Masalah kesihatan mental, termasuk kebimbangan, kemurungan, dan kecenderungan bunuh diri, semakin menjadi perkara biasa, terutamanya di ruang dalam talian di mana individu sering meluahkan perjuangan mereka melalui teks tidak berstruktur. Pengesanan awal dan pengelasan keadaan kesihatan mental ini daripada perbincangan dalam talian adalah penting untuk mediasi dan sokongan yang tepat pada masanya. Projek ini bertujuan membangunkan sistem pengesanan awal masalah kesihatan mental berdasarkan teks tidak berstruktur yang dikongsi dalam talian, seperti di media sosial dan forum. Masalah utama yang dikenal pasti ialah kesukaran untuk mengklasifikasikan status kesihatan mental secara automatik kerana teks pengguna bersifat subjektif, menggunakan pelbagai gaya bahasa, dan mengandungi emosi yang tidak dinyatakan secara jelas, menjadikannya sukar membezakan antara kategori seperti kebimbangan, kemurungan, dan kecenderungan bunuh diri. Kekurangan sistem klasifikasi yang cekap boleh melambatkan intervensi dan sokongan yang diperlukan oleh individu yang berisiko. Untuk menyelesaikan masalah ini, sistem klasifikasi dibangunkan menggunakan pendekatan gabungan pembelajaran mesin (ML) dan pembelajaran mendalam (DL). Model ML seperti Support Vector Machine (SVM) dan Light Gradient Boosting Machine (LGBM) digunakan bersama teknik pengekstrakan ciri seperti TF-IDF dan statistik teks bagi membentuk ciri gabungan. Pendekatan DL pula melibatkan model Long Short-Term Memory (LSTM), Convolutional Neural Networks (CNN), dan Extreme Gradient Boosting with Bidirectional Encoder Representations from Transformers (XGB-BERT) untuk menangkap sifat berurutan dan konteks semantik dalam data teks. Proses pembangunan sistem melibatkan latihan dan penilaian model menggunakan set data seimbang daripada perbincangan kesihatan mental dalam talian, dengan penggunaan metrik prestasi seperti kejituan, ketepatan, dapatan semula, dan skor-F1. Sistem turut dibangunkan bersama antara muka pengguna berasaskan Gradio yang membolehkan pengguna memasukkan teks dan menerima maklum balas segera mengenai status kesihatan mental mereka. Hasil akhir menunjukkan bahawa model LGBM dengan ciri gabungan memberikan prestasi klasifikasi terbaik dalam mengenal pasti keempat-empat kategori status mental iaitu normal, kebimbangan, kemurungan, dan kecenderungan bunuh diri. Sistem ini bukan sahaja menyokong proses pemantauan kesihatan mental secara awal, tetapi juga menyediakan mesej sokongan dan aktiviti interaktif yang mesra pengguna. Projek ini berpotensi besar untuk diperluas ke dalam sistem chatbot atau platform kesihatan digital bagi meningkatkan akses kepada sokongan mental secara proaktif dan berskala.

Kata kunci: Kesihatan Mental, Pembelajaran Mesin, Pembelajaran Mendalam, Pengelasan Teks, Gradio.

ABSTRACT

Mental health issues, including anxiety, depression, and suicidal tendencies, are becoming increasingly common, especially in online spaces where individuals often express their struggles

through unstructured text. Early detection and classification of these mental health conditions from online discussions are crucial for timely mediation and support. This project aims to develop an early detection system for mental health issues based on unstructured text shared online, such as on social media and forums. The main challenge identified is the difficulty in automatically classifying mental health status due to the subjective nature of user-generated content, the diverse use of language, and the presence of emotional expressions that are implicit and not directly stated. These factors make it hard to distinguish between categories such as anxiety, depression, and suicidal tendencies. The lack of an efficient classification system may delay interventions and support for individuals at risk. To address this issue, a classification system is developed using a combination of machine learning (ML) and deep learning (DL) approaches. ML models such as Support Vector Machine (SVM) and Light Gradient Boosting Machine (LGBM) are employed alongside feature extraction techniques like TF-IDF and text statistics to form combined features. The DL approach includes models such as Long Short-Term Memory (LSTM), Convolutional Neural Networks (CNN), and Extreme Gradient Boosting with Bidirectional Encoder Representations from Transformers (XGB-BERT) to capture the sequential and contextual nature of online text data. The system development involves training and evaluating models using a balanced dataset derived from online mental health discussions, with performance metrics such as accuracy, precision, recall, and F1-score. A user-friendly interface is built using Gradio, allowing users to input text and receive immediate feedback on their mental health status. The final result shows that the LGBM model with combined features achieved the best classification performance in identifying all four mental health status categories: normal, anxiety, depression, and suicidal thoughts. The system not only supports early mental health monitoring but also provides support messages and interactive wellness activities. This project has great potential to be expanded into chatbot systems or digital health platforms to enhance mental health support accessibility in a proactive and scalable manner.

Keywords: Mental Health, Machine Learning, Deep Learning, Text Classification, Gradio

1.0 PENGENALAN

Isu kesihatan mental menjadi kebimbangan global yang kritikal, terutamanya semakin kritikal dengan peningkatan penggunaan platform digital di mana individu meluahkan emosi, tekanan, dan perjuangan mental mereka. Menurut Pertubuhan Kesihatan Sedunia (WHO), keadaan kesihatan mental seperti kemurungan dan kebimbangan telah menjejaskan lebih daripada 970 juta orang di seluruh dunia yang menjadikannya salah satu penyebab utama bunuh diri (World Health Organization: WHO, 2022). Selain itu, pandemik COVID-19 telah memperburuk krisis kesihatan mental dengan anggaran peningkatan sebanyak 25% dalam isu kebimbangan dan kemurungan (Mental Health, Brain Health and Substance Use (MSD), 2022). Namun, akses kepada penjagaan kesihatan mental masih terhad, di mana diagnosis awal dan rawatan sering terjejas. Hal ini disebabkan oleh stigma yang berkaitan dengan isu kesihatan mental, kekurangan sumber, serta akses yang terhad kepada profesional.

Selain itu, kajian menunjukkan bahawa platform digital telah menjadi ruang yang penting bagi individu untuk mendedahkan isu kesihatan mental mereka. Dalam konteks pandemik

Covid-19, seperti yang disorot oleh kajian mengenai rakyat Malaysia, 94% responden mencari maklumat di platform media sosial seperti *Facebook*, *Instagram*, dan X (sebelumnya dikenali sebagai *Twitter*) untuk mengurangkan kebimbangan (Nor Shafrin Ahmad et al., 2021). Walaupun platform ini menyediakan ruang untuk sokongan emosi dan perkongsian maklumat, banyak kes tekanan mental mungkin tidak disedari tanpa mekanisme pengesanan yang betul. Hal ini menekankan keperluan untuk model pembelajaran mesin yang lebih canggih yang dapat memantau perbincangan dalam talian untuk mengenal pasti tanda-tanda kebimbangan kesihatan mental dan membolehkan pencegahan awal.

Bagi menangani cabaran yang semakin meningkat ini, alat digital yang memanfaatkan pembelajaran mesin telah menarik perhatian kerana potensinya untuk membantu dalam pengesanan masalah kesihatan mental. Platform media sosial, forum dalam talian, dan saluran komunikasi digital lain menyediakan sumber data yang boleh dianalisis untuk mengesan kebimbangan kesihatan mental. Walaupun teknik pembelajaran mesin tradisional seperti SVM telah digunakan dalam analisis sentimen dan klasifikasi teks, kaedah ini sering gagal memahami nuansa kompleks bahasa manusia (Sattar et al., 2021). Dengan kemunculan pembelajaran mendalam, terutamanya model berasaskan LSTM dan CNN serta model transformer seperti XGB-BERT terdapat peningkatan ketara dalam kemampuan untuk menangkap makna kontekstual dan emosi dalam teks (Sayyida Tabinda Kokab et al., 2022). Ini membolehkan model memahami konteks sebelum dan selepas sesuatu perkataan dalam ayat dengan menjadikan model tersebut sangat berguna bagi pengesanan kesihatan mental.

Kajian ini bertujuan untuk membandingkan keberkesanan pendekatan pembelajaran mesin tradisional dan pembelajaran mendalam dalam pengesanan isu kesihatan mental berdasarkan perbincangan dalam talian. Dengan fokus terhadap model seperti SVM, LGBM, LSTM, CNN dan XGB-BERT, kajian ini akan menilai sejauh mana setiap pendekatan dapat mengenal pasti kebimbangan mental dengan tepat, termasuk keupayaan mereka untuk memahami konteks dan nuansa bahasa.

2.0 KAJIAN LITERATUR

Kajian literatur dalam bidang pengesanan masalah kesihatan mental menunjukkan pelbagai pendekatan telah diterapkan menggunakan pembelajaran mesin tradisional dan pembelajaran mendalam. Dalam pembelajaran mesin tradisional, algoritma seperti Support Vector Machine (SVM), Random Forest, dan Naïve Bayes telah digunakan secara meluas.

Kajian oleh Mohamed et al. (2023) menunjukkan bahawa SVM memberikan prestasi tinggi dalam klasifikasi biner dengan ketepatan sehingga 97.67% bagi data Reddit, manakala Random Forest mencapai sehingga 98.14%, menunjukkan kestabilan tinggi dalam menangani overfitting melalui pelbagai pokok keputusan. Kajian oleh Lyua YW et al. (2020) pula menunjukkan keberkesanan SVM dalam situasi data tidak seimbang. Naïve Bayes juga digunakan untuk teks ringkas kerana kesederhanaannya (Shrestha et al., 2020), tetapi prestasinya menurun apabila diaplikasikan pada teks kompleks dan pelbagai emosi.

Model-model ini biasanya disokong oleh teknik pengekstrakan ciri seperti TF-IDF, Word2Vec dan FastText yang memerlukan proses pra-pemprosesan yang rapi seperti penyingkiran hentikata dan lemmatisasi bagi menjamin kualiti input (Kastrati et al., 2021). Namun, pendekatan ini bergantung kepada pemilihan ciri secara manual yang boleh mengehadkan keupayaan model memahami konteks emosi yang lebih kompleks.

Pendekatan pembelajaran mendalam menawarkan kelebihan yang lebih tinggi melalui keupayaan untuk memahami konteks dan urutan kata dalam teks tidak berstruktur. Kajian oleh Babu & Kanaga (2021) dan Zhu et al. (2024) menunjukkan bahawa model LSTM berkesan dalam memahami urutan kata yang panjang serta menyimpan maklumat kontekstual, walaupun memerlukan set data yang besar dan pemprosesan yang intensif. CNN pula sesuai untuk mengenal pasti pola tempatan dalam teks dan menunjukkan potensi dalam pengelasan teks berkaitan emosi (Zogan et al., 2022). Kajian oleh Salur MU & Aydin I (2020) menunjukkan bahawa gabungan CNN dengan Word2Vec mampu meningkatkan ketepatan dalam bahasa tertentu.

Model transformer seperti BERT memberikan ketepatan tertinggi berbanding model lain kerana keupayaannya menangkap konteks penuh dalam ayat melalui attention mechanism. Kajian oleh Rabani et al. (2020) mencatatkan ketepatan 99.10% menggunakan BERT pada set data Baghdad, manakala kajian oleh Yichao Wu et al. (2024) dan Jain et al. (2021) mengesahkan keupayaan BERT dalam menangani frasa kompleks serta hubungan semantik dalam teks kesihatan mental. Namun, model ini memerlukan sumber pemprosesan tinggi dan masa latihan yang panjang, serta sensitif kepada kualiti dan saiz dataset.

Pendekatan gabungan atau ensemble turut digunakan untuk meningkatkan prestasi klasifikasi. Kajian oleh Arun (2024) menggunakan teknik ensemble seperti stacking dan voting yang menggabungkan kekuatan pelbagai model untuk kestabilan prestasi. Gonzalo A.

Ruz et al. (2020) pula menggunakan voting ensemble yang menggabungkan model tradisional dan DL dalam konteks teks sosial. Pavitra Sankar et al. (2024) menunjukkan keberkesanan boosting untuk mengatasi data sukar dalam analisis kemurungan dan kebimbangan. Tambahan pula, kajian oleh Rabani et al. (2020) menunjukkan bahawa gabungan CNN dan LSTM melalui teknik bagging memberikan kestabilan dan ketepatan lebih tinggi dalam tugas klasifikasi kesihatan mental.

Secara keseluruhannya, pendekatan hybrid seperti XGB-BERT, seperti yang disarankan oleh Adarsh et al. (2022) dan Nguyen et al. (2023), menggabungkan kekuatan contextual embeddings BERT dan kecekapan pengelasan oleh XGBoost, terbukti meningkatkan ketepatan klasifikasi dalam pelbagai kajian. Gabungan pelbagai kaedah ini bukan sahaja meningkatkan ketepatan tetapi juga keupayaan generalisasi sistem, menjadikannya lebih sesuai untuk aplikasi dunia sebenar dalam pengesanan awal masalah kesihatan mental di platform digital.

3.0 METODOLOGI KAJIAN

Metodologi yang digunakan dalam pembangunan projek ini ialah pendekatan berasaskan kitaran pembangunan sistem data sains yang merangkumi proses pengumpulan data, pra-pemprosesan, pembangunan model klasifikasi, pembangunan antara muka pengguna, dan pengujian sistem. Pendekatan ini dipilih kerana ia memberikan fleksibiliti dalam merancang dan membina sistem secara berperingkat serta membolehkan pelarasan berterusan terhadap model dan sistem berdasarkan maklum balas prestasi. Metodologi ini sangat sesuai digunakan bagi projek yang memerlukan pembangunan sistem klasifikasi teks secara automatik dan mesra pengguna seperti yang ditunjukkan dalam Rajah 1.



Rajah 1: Fasa-fasa Metodologi Kitaran Hayat Pembelajaran Mesin (Moez Krichen et al., 2022)

3.1 Fasa Pengumpulan, Penyediaan dan Pra-Pemrosesan Data

Fasa ini dimulakan dengan pengumpulan set data daripada sumber terbuka iaitu Kaggle, yang mengandungi sebanyak 53,042 pernyataan teks berkaitan isu kesihatan mental. Data tersebut kemudiannya melalui proses penyediaan yang melibatkan aktiviti pembersihan bagi menyingkirkan unsur yang tidak relevan seperti simbol khas, pautan laman sesawang dan perkataan yang tidak membawa makna signifikan. Seterusnya, teks ditukar kepada huruf kecil, dipecahkan kepada token, tanda baca telah dibuang, dan diproses menggunakan teknik lemmatisasi. Proses pra-pemrosesan ini dijalankan bagi menstrukturkan data dalam bentuk yang lebih seragam dan bersih supaya sesuai digunakan dalam proses latihan model klasifikasi yang akan dibangunkan.

3.2 Fasa Penganalisan Data dan Penerokaan Ciri

Dalam fasa ini, analisis eksploratori dijalankan bagi memahami taburan kategori kesihatan mental dalam data, serta mengenal pasti potensi ciri-ciri yang berkaitan. Pendekatan representasi teks seperti Term Frequency–Inverse Document Frequency (TF-IDF) digunakan untuk mengekstrak ciri berdasarkan kepentingan perkataan dalam setiap dokumen berbanding keseluruhan korpus. Selain itu, ciri statistik teks (TextStats) turut digunakan yang merangkumi panjang ayat, jumlah perkataan, bilangan huruf besar, bilangan tanda seru, bilangan tanda soal, serta jumlah perkataan positif dan negatif berdasarkan

senarai kamus sentimen.

Bagi model pembelajaran mendalam, pendekatan word embedding digunakan bagi merepresentasikan teks dalam bentuk vektor berdimensi tetap yang menangkap hubungan semantik antara perkataan. Bagi model CNN dan LSTM, embedding lapisan digunakan dalam rangkaian neural untuk menangani sifat berurutan dalam data teks. Sementara itu, model XGB-BERT menggunakan representasi daripada BERT yang mampu memahami konteks perkataan dalam ayat secara dua hala. Kesemua ciri ini digabungkan mengikut kesesuaian model bagi memberikan pemahaman yang lebih mendalam terhadap corak emosi dalam data teks, seterusnya menjadi asas kepada proses klasifikasi.

3.3 Fasa Pembinaan dan Latihan Model

Fasa ini melibatkan pembinaan dan latihan lima model klasifikasi yang berbeza, terdiri daripada dua model pembelajaran mesin tradisional iaitu SVM dan LGBM, serta tiga model pembelajaran mendalam iaitu LSTM, CNN dan XGB-BERT.

Bagi model tradisional seperti SVM dan LGBM, ciri-ciri teks digabungkan menggunakan pendekatan FeatureUnion yang menyatukan representasi TF-IDF dan TextStats menjadi satu set ciri gabungan. Model SVM dilatih menggunakan ciri TF-IDF sahaja untuk keserasian dimensi, manakala LGBM memanfaatkan gabungan kedua-dua ciri tersebut tanpa memerlukan penyesuaian skala. Untuk model pembelajaran mendalam, embedding layer digunakan untuk memetakan setiap token teks kepada vektor representasi, dengan LSTM dan CNN dilatih daripada awal menggunakan embedding tersebut, manakala XGB-BERT menggunakan embedding daripada model BERT yang telah ditetapkan.

Set data dibahagikan kepada set latihan dan set ujian selepas diseimbangkan secara dalaman mengikut taburan kategori. Latihan model dijalankan secara berasingan dan diselaraskan melalui penalaan hiperparameter menggunakan teknik grid search atau tetapan lalai terbaik yang telah diuji. Prestasi setiap model dinilai menggunakan data ujian yang tidak pernah dilihat sebelum ini.

3.4 Fasa Penilaian Model

Setiap model yang telah dilatih dinilai menggunakan empat metrik utama iaitu kejutuan, ketepatan, dapatan semula, dan skor-F1. Penilaian dilakukan ke atas data ujian untuk mengukur keupayaan model mengklasifikasikan teks dengan betul ke dalam empat kategori utama, iaitu *Normal*, *Anxiety*, *Depression*, dan *Suicidal*.

Hasil perbandingan mendapati bahawa model LGBM yang menggunakan gabungan ciri TF-IDF dan TextStats menunjukkan prestasi paling konsisten dan stabil dengan skor-F1 yang tinggi bagi semua kategori. Justeru, model ini dipilih untuk diimplementasikan dalam sistem akhir kerana keupayaannya memberikan klasifikasi yang tepat dan boleh dipercayai, di samping masa inferens yang pantas dan integrasi yang mudah ke dalam antara muka pengguna.

3.5 Fasa Implementasi Sistem

Model terbaik yang dipilih, iaitu LGBM dengan gabungan ciri, telah diintegrasikan ke dalam sistem klasifikasi kesihatan mental menggunakan antara muka berasaskan Gradio. Antara muka ini dibangunkan untuk membolehkan pengguna memasukkan teks secara terus dan menerima keputusan klasifikasi dalam bentuk status kesihatan mental serta tahap keyakinan model terhadap keputusan tersebut. Selain itu, sistem turut memaparkan mesej motivasi berdasarkan kategori emosi yang dikesan untuk memberikan sokongan psikologi awal kepada pengguna.

Ciri tambahan seperti permainan interaktif turut disediakan bagi meningkatkan pengalaman pengguna, termasuk permainan Positivity Card yang memaparkan kata-kata positif secara rawak dan permainan Mood Colour Picker yang membolehkan pengguna memilih warna emosi mereka sendiri. Keseluruhan antara muka direka bentuk secara mesra pengguna, responsif dan mudah digunakan oleh semua peringkat pengguna.

3.6 Fasa Pengujian Sistem

Fasa ini dijalankan bagi memastikan sistem yang dibangunkan berfungsi dengan baik dari sudut teknikal serta memberikan pengalaman pengguna yang memuaskan. Pengujian sistem dimulakan dengan menguji sebanyak 20 ayat dunia sebenar yang mengandungi kesalahan ejaan dan penggunaan bahasa tidak formal bagi menilai keupayaan model dalam mengendalikan variasi input yang lebih realistik. Keputusan menunjukkan bahawa model LGBM masih mampu memberikan klasifikasi yang tepat dan konsisten, sekali gus membuktikan ketahanannya terhadap data dunia sebenar.

Selain itu, pengujian penerimaan pengguna turut dilaksanakan untuk mendapatkan maklum balas terhadap kebolegunaan sistem. Satu soal selidik telah dibangunkan dan diedarkan melalui platform Google Form kepada sepuluh orang responden. Soal selidik ini terdiri daripada lima bahagian utama, iaitu maklumat latar belakang responden, kemudahan

penggunaan antara muka, keberkesanan ciri sistem, penilaian terhadap reka bentuk serta permainan interaktif, dan penilaian keseluruhan sistem.

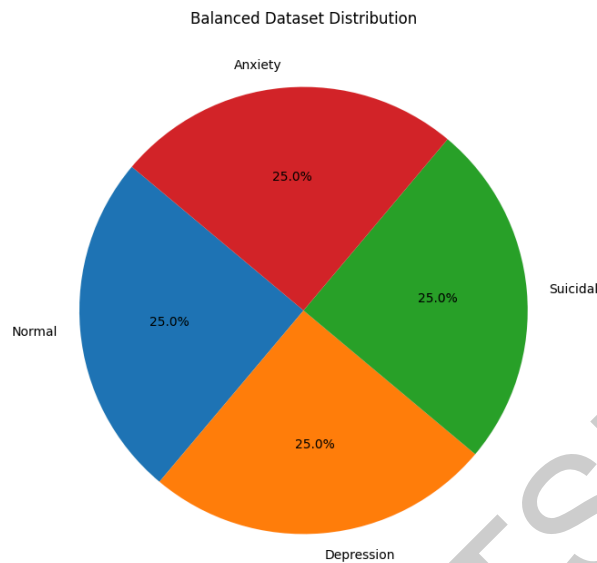
Setiap soalan dalam soal selidik menggunakan skala Likert lima mata yang merangkumi tahap persetujuan daripada “Sangat Tidak Setuju” hingga “Sangat Setuju”. Di samping itu, satu soalan terbuka turut disediakan bagi mendapatkan cadangan penambahbaikan daripada pengguna. Hasil analisis menunjukkan bahawa majoriti responden memberikan maklum balas yang positif terhadap sistem yang dibangunkan. Dapatan ini membuktikan bahawa sistem bukan sahaja berfungsi dengan baik dari aspek teknikal, malah mampu menawarkan pengalaman pengguna yang mesra, interaktif, dan menyokong keperluan dalam konteks kesihatan mental.

4.0 HASIL

4.1 Taburan Set Data Seimbang

Sebelum latihan model dijalankan, analisis terhadap data asal menunjukkan ketidakseimbangan kelas yang ketara, di mana kelas seperti Normal dan Kemurungan mempunyai jumlah pernyataan yang jauh lebih tinggi berbanding Kebimbangan dan Kecenderungan Bunuh Diri. Ketidakseimbangan ini boleh menyebabkan model cenderung untuk mengklasifikasikan input ke dalam kelas majoriti, sekali gus mengurangkan ketepatan bagi kelas minoriti yang lebih kritikal seperti Suicidal.

Oleh itu, kaedah upsampling berdasarkan kelas minoriti digunakan untuk menyeimbangkan semula taburan data. Teknik ini melibatkan penduplikasian rawak pernyataan daripada kelas minoriti sehingga jumlahnya sepadan dengan kelas majoriti. Proses ini membolehkan model belajar daripada jumlah data yang seimbang untuk setiap kategori dan meningkatkan keadilan serta sensitiviti klasifikasi. Rajah 2 menunjukkan taburan set data yang seimbang mengikut kategori kesihatan mental.



Rajah 2: Taburan Set Data Seimbang Mengikut Kategori Kesihatan Mental

Rajah 2 tersebut menunjukkan bahawa setiap kelas mewakili 25% daripada keseluruhan set data, menandakan bahawa tiada lagi ketidakseimbangan yang boleh memberi kesan terhadap keputusan klasifikasi.

4.2 Prestasi Model Mengikut Kategori Kesihatan Mental

Sebanyak lima model klasifikasi telah diuji dalam sistem ini untuk mengenal pasti pendekatan paling sesuai bagi tugas pengelasan teks berkaitan kesihatan mental. Model-model tersebut mewakili dua pendekatan utama iaitu pembelajaran mesin tradisional dan pembelajaran mendalam. SVM dibina berasaskan ciri TF-IDF; model LGBM pula menggabungkan ciri TF-IDF dengan statistik teks; manakala model LSTM dan CNN masing-masing menggunakan pendekatan Word Embedding. Selain itu, model hibrid XGB-BERT menggabungkan embedding daripada transformer BERT dan komponen pengelasan berasaskan XGBoost.

Pengujian ke atas kelima-lima model dilakukan berdasarkan empat metrik utama iaitu ketetapan, dapatan semula, skor-F1 dan kejituan keseluruhan. Metrik-metrik ini digunakan untuk mengukur keupayaan setiap model dalam mengklasifikasikan empat kategori kesihatan mental iaitu Normal, Kebimbangan, Kemurungan dan Kecenderungan Bunuh Diri. Berdasarkan keputusan yang diperoleh, model LightGBM menunjukkan prestasi yang positif dan konsisten, terutamanya dalam kelas Kebimbangan dengan skor-F1 tertinggi sebanyak 0.94, serta kelas Normal dengan skor-F1 sebanyak 0.91. Keupayaan ini

disumbangkan oleh kekuatan LGBM dalam menggabungkan maklumat linguistik (TF-IDF) dan ciri struktur ayat (statistik teks seperti panjang ayat, bilangan huruf besar, dan tanda baca), sekali gus menjadikan model ini lebih sensitif terhadap corak bahasa yang mencerminkan emosi tertentu. Metrik klasifikasi mengikut kelas bagi kategori kesihatan mental ditunjukkan dalam Jadual 1, manakala Jadual 2 memaparkan nilai kejituan keseluruhan bagi setiap model yang dibangunkan.

Jadual 1 Metrik Klasifikasi mengikut Kategori Kesihatan Mental

Kategori Kesihatan Mental	Model	Penilaian Metrik		
		Ketetapan	Dapatan Semula	Skor-F1
Kebimbangan <i>Anxiety</i>	SVM + TF-IDF	0.90	0.90	0.90
	LGBM + TF-IDF + TextStats	0.93	0.96	0.94
	LSTM + Word Embedding	0.90	0.94	0.92
	CNN + Word Embedding	0.89	0.92	0.90
	XGB-BERT + BERT Embedding + XGB Classifier	0.91	0.92	0.91
Kemurungan <i>Depression</i>	SVM + TF-IDF	0.70	0.63	0.67
	LGBM + TF-IDF + TextStats	0.78	0.70	0.74
	LSTM + Word Embedding	0.67	0.64	0.65
	CNN + Word Embedding	0.68	0.68	0.68

	XGB-BERT + BERT Embedding + XGB Classifier	0.73	0.67	0.70
Normal	SVM + TF-IDF	0.84	0.92	0.88
	LGBM + TF-IDF + TextStats	0.89	0.92	0.91
	LSTM + Word Embedding	0.92	0.88	0.90
	CNN + Word Embedding	0.91	0.86	0.88
	XGB-BERT + BERT Embedding + XGB Classifier	0.87	0.93	0.90
Bunuh Diri Suicidal	SVM + TF-IDF	0.71	0.71	0.71
	LGBM + TF-IDF + TextStats	0.75	0.79	0.77
	LSTM + Word Embedding	0.69	0.72	0.70
	CNN + Word Embedding	0.72	0.71	0.73
	XGB-BERT + BERT Embedding + XGB Classifier	0.73	0.74	0.73

Jadual 2 Kejitan Keseluruhan Setiap Model

Model	Kejitan
SVM + TF-IDF	0.79
LGBM + TF-IDF + TextStats	0.84
LSTM + Word Embedding	0.79
CNN +Word Embedding	0.80
XGB-BERT + BERT Embedding + XGBoost Classifier	0.81

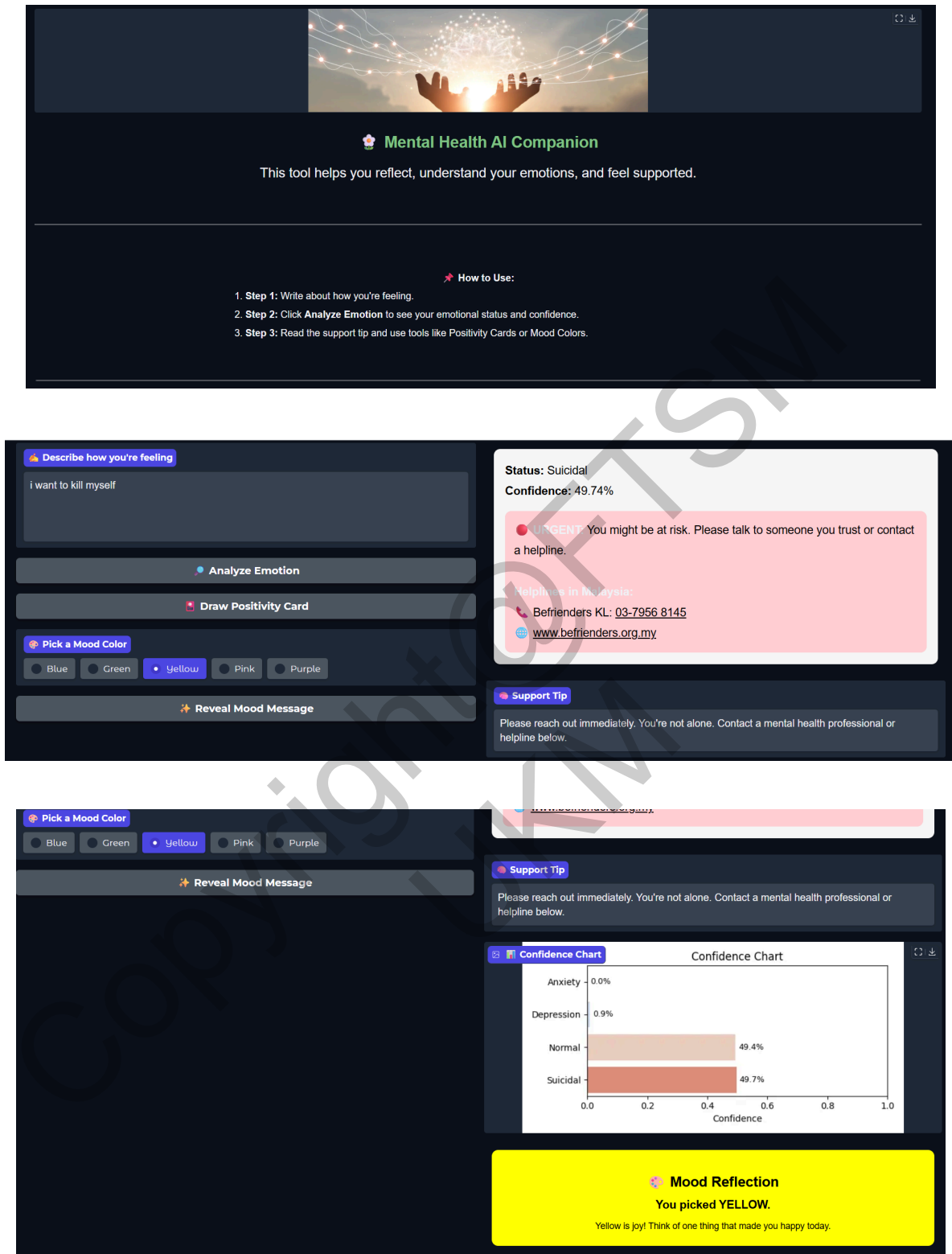
Walaupun model XGB-BERT menunjukkan prestasi metrik yang hampir menyamai LightGBM, pelaksanaannya memerlukan masa latihan yang jauh lebih lama serta sumber pengiraan yang lebih tinggi seperti penggunaan GPU, menjadikannya kurang sesuai untuk sistem yang ingin dibangunkan secara ringan dan responsif. Di samping itu, walaupun LSTM dan CNN berjaya mengekstrak konteks ayat melalui pembelajaran urutan, kedua-dua model ini tidak dapat mengatasi keberkesanan LGBM yang mampu memanfaatkan pelbagai jenis ciri dengan lebih cekap dan stabil.

Berdasarkan analisis menyeluruh dari aspek metrik, kecekapan pemprosesan, serta kebolehlaksanaan dalam konteks sistem sebenar, LGBM telah dipilih sebagai model utama yang diintegrasikan ke dalam sistem klasifikasi kesihatan mental ini.

4.3 Pengujian Antara Muka Pengguna (UI)

Antara muka pengguna merupakan elemen penting dalam sistem klasifikasi kesihatan mental ini kerana ia menjadi penghubung utama antara pengguna dan model klasifikasi yang dibangunkan. Untuk memastikan sistem ini mudah diakses oleh pelbagai lapisan pengguna, termasuk pengguna tanpa latar belakang teknikal, antara muka dibangunkan menggunakan platform Gradio yang bersifat ringkas, inisiatif, dan responsif.

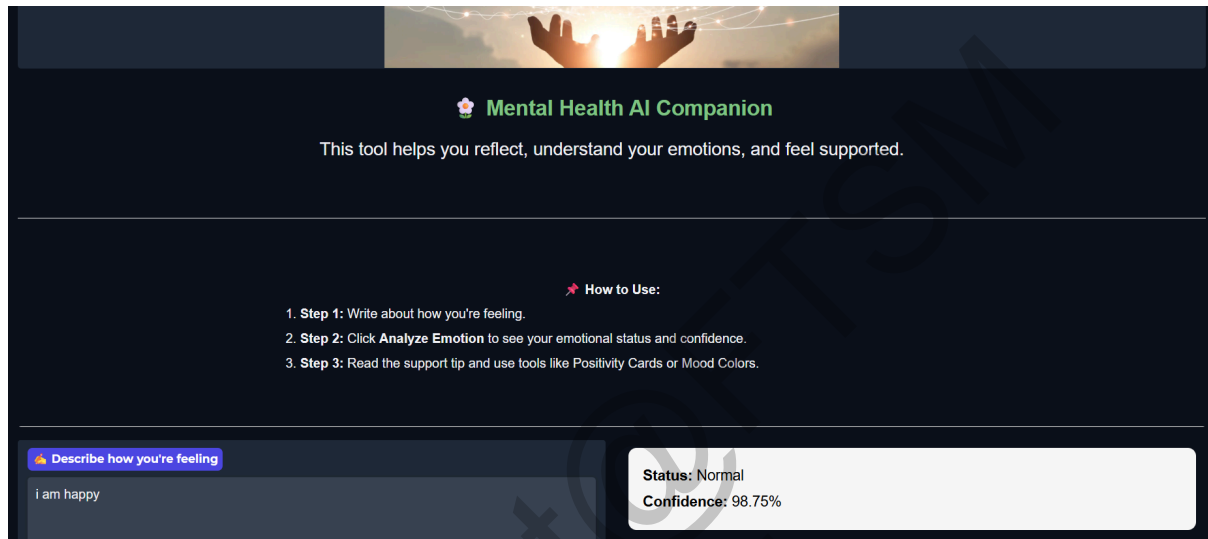
Pengujian antara muka dijalankan bagi menilai kefungsian setiap komponen utama sistem dari segi kelancaran operasi, ketepatan paparan, dan pengalaman pengguna secara keseluruhan. Fokus utama pengujian ini adalah untuk memastikan bahawa sistem dapat menerima input teks, memproses dan mengembalikan output klasifikasi dengan jelas dan tepat tanpa sebarang ralat atau kelewatan yang ketara. Rajah 3 menunjukkan antara muka lengkap bagi sistem klasifikasi kesihatan mental.



Rajah 3: Antara Muka Pengguna Lengkap

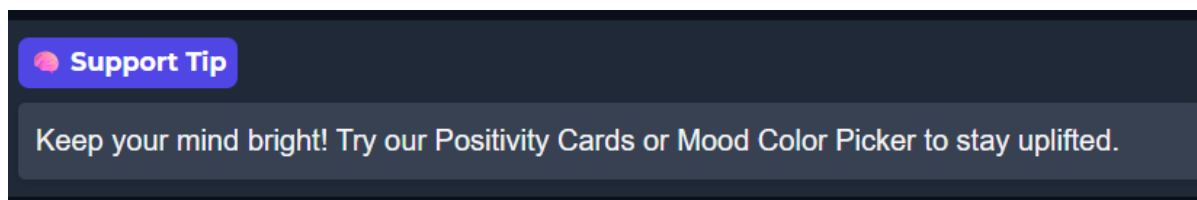
Komponen pertama yang diuji ialah ruangan input teks dan butang “Analyze Emotion”, yang membolehkan pengguna memasukkan sebarang ayat berkaitan keadaan emosi mereka.

Setelah menekan butang tersebut, sistem memaparkan keputusan klasifikasi dalam bentuk label emosi iaitu Normal, Anxiety, Depression atau Suicidal, bersama tahap keyakinan dalam bentuk peratusan. Hasil pengujian menunjukkan bahawa sistem berjaya memproses input dan memaparkan output dalam masa kurang daripada tiga saat secara purata, tanpa sebarang isu fungsi atau kelewatan yang ketara.



Rajah 4: Paparan Input Teks dan Hasil Klasifikasi

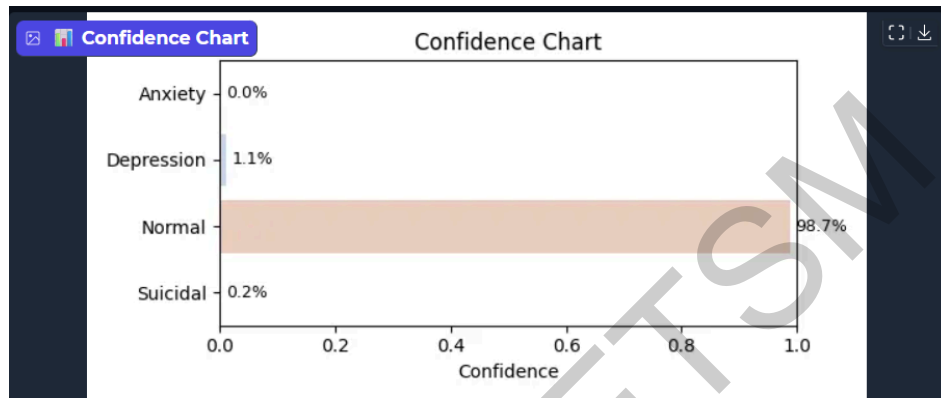
Selain klasifikasi emosi, antara muka juga memaparkan mesej sokongan automatik seperti yang ditunjukkan dalam Rajah 5, apabila sistem mengesan emosi negatif seperti Kemurungan dan Kecenderungan Bunuh Diri. Mesej ini direka untuk memberi dorongan awal kepada pengguna agar mereka tidak merasa sendirian dan digalakkan untuk mendapatkan bantuan lanjut. Mesej dipaparkan secara automatik dan berbeza-beza berdasarkan emosi yang dikesan, serta menyampaikan nada yang lembut, menyokong dan tidak menghakimi.



Rajah 5: Paparan Tip Sokongan daripada Sistem

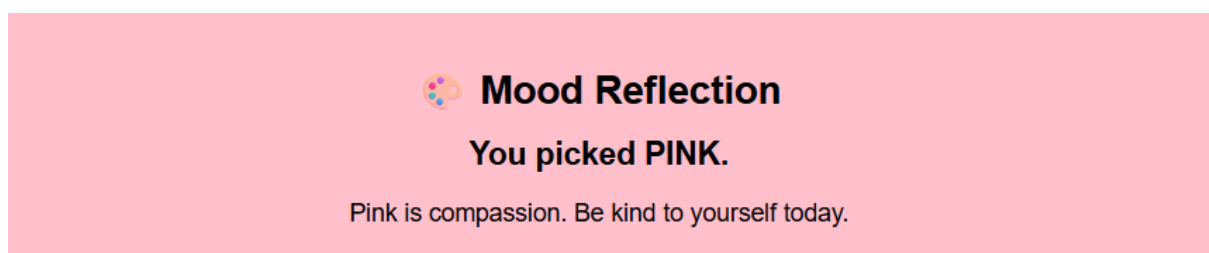
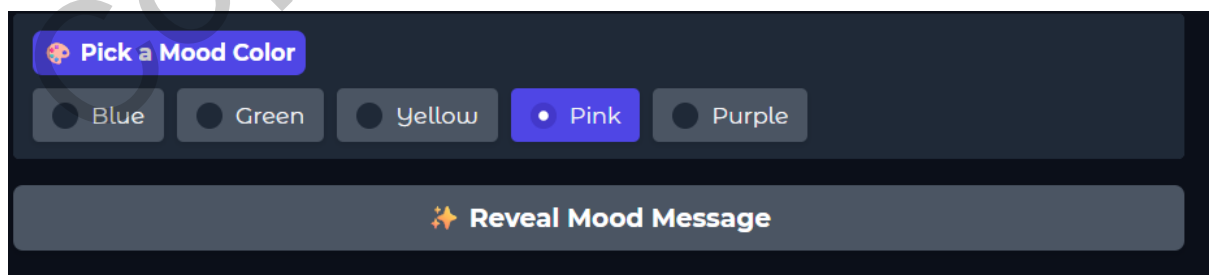
Tambahan pula, satu carta keyakinan berbentuk bar turut dipaparkan dalam Rajah 6 bagi menunjukkan tahap keyakinan model terhadap keempat-empat kategori yang diuji.

Carta ini memberi gambaran yang lebih jelas kepada pengguna bahawa emosi yang dialami tidak bersifat tunggal, malah boleh melibatkan campuran emosi dengan tahap tertentu. Ini secara tidak langsung meningkatkan kepercayaan pengguna terhadap sistem kerana keputusan yang dipaparkan disokong oleh visualisasi berasaskan kebarangkalian.



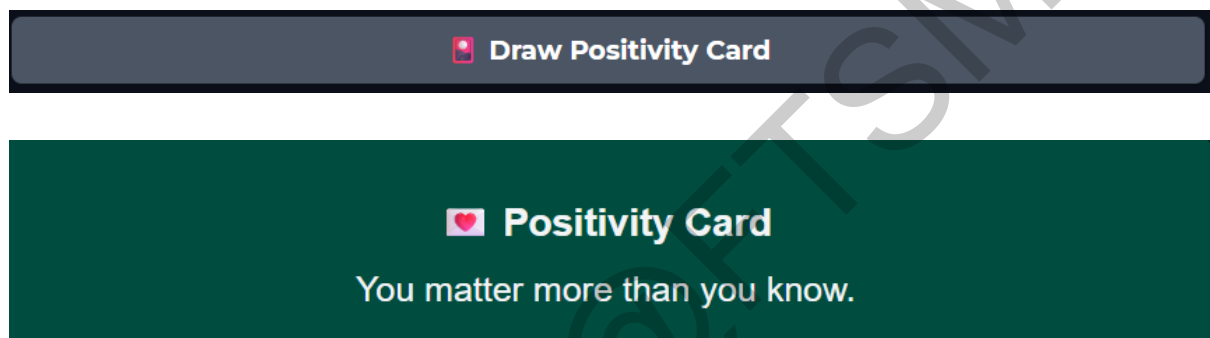
Rajah 6: Paparan Carta Keyakinan Model

Fungsi interaktif tambahan seperti Mood Colour Picker juga diuji dan didapati berfungsi dengan baik seperti yang ditunjukkan dalam Rajah 7. Melalui ciri ini, pengguna dapat memilih warna yang melambangkan perasaan mereka, dan sistem akan memberikan mesej motivasi secara rawak berdasarkan pilihan warna tersebut. Pendekatan ini membolehkan pengguna meluahkan emosi secara simbolik dan menerima maklum balas positif yang bersesuaian dengan keadaan emosi mereka. Mesej yang dipaparkan juga diubah secara rawak setiap kali pengguna memilih warna yang sama, bagi memastikan pengalaman yang tidak berulang dan lebih peribadi.



Rajah 7: Paparan Pemilihan Mood dan Mesej Emosi

Akhir sekali, fungsi Positivity Card yang ditunjukkan dalam Rajah 8 sebagai permainan sokongan ringan selepas klasifikasi turut diuji. Fungsi ini membolehkan pengguna menerima kad motivasi dengan mesej ringkas yang positif dan menenangkan. Fungsi ini terbukti berfungsi dengan baik dan menerima maklum balas positif dalam kalangan pengguna semasa ujian.



Rajah 8: Paparan Pilihan Permainan Emosi Interaktif

Secara keseluruhannya, pengujian terhadap antara muka pengguna menunjukkan bahawa sistem ini bukan sahaja berfungsi dengan stabil dan responsif, malah turut menyediakan pengalaman pengguna yang menyenangkan, mesra emosi dan mudah difahami. Semua komponen utama beroperasi dengan lancar tanpa sebarang isu teknikal, menjadikan sistem ini sesuai digunakan dalam pelbagai konteks pengguna akhir.

4.4 Pengujian Terhadap Input Dunia Sebenar

Bagi menilai kebolehlaksanaan sistem dalam konteks dunia sebenar, satu siri pengujian telah dijalankan dengan menggunakan ayat-ayat yang menyerupai bentuk komunikasi harian pengguna seperti mesej peribadi, catatan media sosial dan ungkapan spontan yang mengandungi kesalahan ejaan, simbol berlebihan dan struktur ayat tidak formal. Dalam ujian ini, sebanyak dua puluh ayat diuji yang merangkumi keempat-empat kategori utama kesihatan mental iaitu *Normal*, *Anxiety*, *Depression* dan *Suicidal*. Model LGBM yang telah dikenal pasti sebagai model utama digunakan untuk mengklasifikasikan input-input ini. Setiap input dimasukkan melalui antara muka pengguna Gradio, dan keputusan klasifikasi dibandingkan dengan label sebenar untuk menilai ketepatan. Jadual 3 menunjukkan keputusan klasifikasi model LGBM apabila diuji menggunakan data teks dunia sebenar,

bagi menilai ketepatan dan kebolegunaan model dalam situasi praktikal.

Jadual 3 Keputusan Klasifikasi Menggunakan LGBM ke atas Teks Dunia Sebenar

Bil	Input Teks	Label Sebenar	Klasifikasi Sistem	Kesalahan Ejaan	Keputusan
1	I'm feeling great today.	Normal	Normal	Tidak	Lulus
2	Feeling content and grateful.	Normal	Normal	Tidak	Lulus
3	I'm excited to see my friends tomorrow.	Normal	Normal	Ya	Lulus
4	Just finished reading a goooooood book.	Normal	Normal	Ya	Lulus
5	I feel calm and focused right now!!!!	Normal	Normal	Ya	Lulus
6	I'm so worried I can't focus.	Anxiety	Anxiety	Ya	Lulus
7	I always feel like something bad is about to happen.	Anxiety	Depression	Tidak	Gagal
8	I avoid people because I get anxious.	Anxiety	Anxiety	Tidak	Lulus
9	I get shay and cold when I worry.	Anxiety	Anxiety	Ya	Lulus
10	I just want to sleep and never wake up.	Depression	Depression	Tidak	Lulus

11	I feel like nothing matters anymore.	Depression	Depression	Tidak	Lulus
12	can't stop feeling this emptines inside.	Depression	Depression	Ya	Lulus
13	life feels pointless and dull.	Depression	Depression	Tidak	Lulus
14	i smile but deep down i'm broken.	Depression	Normal	Tidak	Gagal
15	nothing seems to matttr these days.	Depression	Normal	Ya	Gagal
16	I don't see the point of living.	Suicidal	Suicidal	Ya	Lulus
17	nobody cares, so why should i stay?	Suicidal	Depression	Tidak	Gagal
18	i want all this suffering to stop permanantly.	Suicidal	Suicidal	Ya	Lulus
19	ending my life feels like the only option left.	Suicidal	Suicidal	Tidak	Lulus
20	I'm done with life.	Suicidal	Suicidal	Tidak	Lulus

Daripada keseluruhan dua puluh ayat yang diuji, sebanyak enam belas ayat telah diklasifikasikan dengan betul, menjadikan kadar ketepatan model sebanyak 80% dalam senario dunia sebenar. Prestasi ini membuktikan bahawa model LGBM yang dibangunkan mempunyai ketahanan yang tinggi terhadap gangguan linguistik seperti ejaan salah dan ayat tidak formal.

Walaupun terdapat empat kesalahan klasifikasi, kebanyakan daripadanya melibatkan kekeliruan semantik antara emosi yang hampir, contohnya antara Kebimbangan dan Kemurungan, serta Kemurungan dan Kecenderungan Bunuh Diri. Hal ini wajar kerana dalam realiti, emosi manusia sering hadir dalam bentuk yang kompleks dan bertindih. Contoh yang ketara ialah ayat “I always feel like something bad is about to happen” yang diklasifikasikan sebagai Kemurungan walaupun secara teori ia tergolong dalam Kebimbangan.

Secara keseluruhannya, pengujian ini membuktikan bahawa sistem ini bukan sahaja stabil dalam keadaan terkawal, tetapi juga mempunyai kebolehgunaan yang tinggi dalam persekitaran sebenar. Ini memperkukuhkan kebolehpercayaan sistem sebagai alat bantu sokongan awal untuk mengenal pasti risiko kesihatan mental berdasarkan teks.

4.5 Pengujian Penerimaan Pengguna (UAT)

Bagi menilai tahap keberkesanan, kebolehgunaan dan kepuasan pengguna terhadap sistem klasifikasi kesihatan mental yang dibangunkan, satu sesi UAT telah dilaksanakan. Ujian ini dijalankan menggunakan borang soal selidik atas talian (Google Form) yang merangkumi pelbagai aspek sistem termasuk fungsi klasifikasi, antara muka pengguna, kefahaman keputusan, serta ciri-ciri interaktif seperti permainan dan mesej motivasi. Seramai 10 orang responden telah mengambil bahagian dalam sesi ini, terdiri daripada pelajar universiti dan orang awam dengan latar belakang yang pelbagai dari segi umur, jantina dan pengalaman penggunaan aplikasi berkaitan kesihatan mental.

Soal selidik ini terbahagi kepada lima bahagian utama iaitu Maklumat Responden, Kemudahan Penggunaan Antara Muka, Kebolehgunaan Ciri Sistem, Reka Bentuk dan Fungsi Permainan, serta Penilaian Keseluruhan Sistem. Setiap bahagian mengandungi soalan berbentuk skala Likert lima tahap bermula dari “Sangat Tidak Setuju” hingga “Sangat Setuju”, bagi membolehkan penilaian kuantitatif dibuat dengan lebih sistematik.

Hasil UAT menunjukkan maklum balas yang sangat positif daripada semua responden. Dalam aspek antara muka, sebanyak 70% responden sangat bersetuju bahawa sistem mudah digunakan, manakala selebihnya turut memberikan reaksi positif. Fungsi butang seperti “Analyze Emotion” dan carta keyakinan juga diuji dan disahkan berfungsi dengan lancar.

Dari segi kefahaman terhadap klasifikasi, 90% responden menyatakan keputusan klasifikasi mudah difahami, manakala 80% bersetuju bahawa carta keyakinan membantu mereka memahami keadaan emosi dengan lebih baik. Jadual 4 menunjukkan ringkasan sorotan hasil Pengujian Penerimaan Pengguna (UAT) yang dijalankan bagi menilai tahap kepuasan dan kebolehgunaan sistem daripada perspektif pengguna.

Jadual 4 Ringkasan Sorotan Pengujian Penerimaan Pengguna (UAT)

Bahagian	Fokus Penilaian	Sorotan Utama
A – Maklumat Responden	Demografi & pengalaman sistem mental	70% tidak pernah guna sistem serupa
B – Antara Muka	Kemudahan penggunaan & respons sistem	100% setuju mudah digunakan, 90% pantas & lancar
C – Kefungsian Sistem	Kefahaman klasifikasi & emosi	90% faham label dan klasifikasi
D – Permainan & Warna	Bantuan emosi melalui ciri interaktif	80% rasa emosi terbantu melalui Mood Colour & Card
E – Penilaian Umum	Kepuasan & niat guna semula	90% puas hati dan akan guna serta syorkan semula

Secara keseluruhannya, dapatan UAT menunjukkan bahawa sistem ini telah diterima dengan baik oleh pengguna akhir dari pelbagai latar belakang. Pengguna bukan sahaja dapat memahami fungsi dan hasil klasifikasi dengan mudah, malah turut mendapat manfaat psikologi melalui interaksi dengan elemen permainan dan mesej motivasi. Hal ini membuktikan bahawa sistem bukan sahaja stabil dari sudut teknikal, tetapi juga mampu memberikan pengalaman pengguna yang positif dan menyokong aspek kesejahteraan emosi.

Bagi meningkatkan keberkesanan sistem pada masa hadapan, beberapa penambahbaikan boleh dilaksanakan merangkumi aspek teknikal dan pengalaman pengguna. Dari segi teknikal, model boleh dipertingkatkan dengan menggunakan teknik pembelajaran pengukuhan atau ensemble yang lebih maju untuk meningkatkan ketepatan klasifikasi, khususnya bagi kelas berisiko tinggi seperti Kemurungan dan Kecenderungan Bunuh Diri yang masih mencatat sedikit kesalahan klasifikasi. Di samping itu, integrasi pembetulan ejaan automatik dan pengesan bahasa tidak formal dapat membantu menambah baik ketahanan sistem

terhadap variasi input dunia sebenar. Dari sudut pengguna, antara muka boleh diperkaya dengan penjelasan ringkas mengenai maksud setiap label emosi untuk membantu pengguna baharu memahami keputusan yang dipaparkan. Fungsi cadangan seperti pautan kepada perkhidmatan sokongan profesional atau chatbot mesra emosi juga boleh disertakan sebagai tindak lanjut bagi pengguna yang dikesan mengalami tekanan emosi yang serius. Selain itu, sistem boleh diperluaskan dengan sokongan pelbagai bahasa atau dialek tempatan agar lebih inklusif untuk pengguna di Malaysia. Kesemua penambahbaikan ini diharap dapat meningkatkan bukan sahaja ketepatan sistem, tetapi juga impak sosial dan emosi kepada pengguna yang memerlukan bantuan awal berkaitan kesihatan mental.

5.0 KESIMPULAN

Secara keseluruhan, sistem pengesanan kesihatan mental berasaskan teks ini telah dibangunkan dan diuji dengan jayanya melalui pendekatan yang menggabungkan elemen pembelajaran mesin, pemprosesan bahasa semula jadi, dan reka bentuk antara muka mesra pengguna. Hasil daripada ujian prestasi model, pengujian input dunia sebenar, serta maklum balas daripada pengguna akhir menunjukkan bahawa sistem ini bukan sahaja stabil dari aspek teknikal, tetapi juga mampu memberikan impak positif dari sudut pengalaman pengguna dan kesedaran emosi. Model LGBM yang digunakan sebagai model utama telah menunjukkan prestasi konsisten dan cemerlang dalam semua kategori, menjadikannya sesuai untuk integrasi dalam sistem klasifikasi automatik.

Menariknya, walaupun model pembelajaran mendalam seperti XGB-BERT dan LSTM sering dianggap lebih canggih, dapatan projek ini membuktikan bahawa model pembelajaran mesin tradisional juga berdaya saing. Dengan penggunaan gabungan ciri yang sesuai seperti TF-IDF dan statistik teks, model tradisional seperti LGBM mampu mencapai prestasi yang setanding, dan dalam kes ini, mengatasi model mendalam dari segi kejituan, kecekapan masa, serta kebolehgunaan dalam sistem dunia sebenar. Hal ini membuktikan bahawa pendekatan pembelajaran mesin masih relevan, efisien dan praktikal, khususnya apabila diaplikasikan dalam persekitaran yang memerlukan kecekapan pemprosesan dan integrasi antara muka secara ringan.

Sistem ini mempunyai beberapa kekuatan utama yang menyumbang kepada tahap keberkesanannya secara keseluruhan. Dari aspek teknikal, model LGBM yang digunakan telah menunjukkan ketepatan klasifikasi yang tinggi dan konsisten untuk keempat-empat

kategori kesihatan mental, termasuk kategori kritikal seperti Kecenderungan Bunuh Diri, sekali gus membuktikan kebolehpercayaannya dalam situasi dunia sebenar. Dari sudut pemprosesan data, fungsi pra-pemprosesan yang dibangunkan dalam sistem ini mampu menormalkan input teks secara efektif, termasuk ayat tidak formal yang mengandungi kesalahan ejaan, tanda baca berlebihan dan gaya bahasa santai iaitu ciri yang lazim ditemui dalam komunikasi harian pengguna. Sementara itu, antara muka pengguna yang dibina menggunakan platform Gradio bukan sahaja responsif dan mudah digunakan, malah turut menyokong elemen interaktif yang meningkatkan pengalaman pengguna. Ciri tambahan seperti Mood Colour Picker, kad motivasi Positivity Card, dan mesej sokongan automatik bukan sahaja memperkayakan fungsi sistem, malah menambah dimensi emosi dan empati, menjadikan sistem ini bukan sekadar alat klasifikasi, tetapi juga sebagai medium sokongan psikologi ringan yang mesra pengguna.

Walaupun sistem ini telah menunjukkan prestasi yang baik, terdapat beberapa kelemahan yang dikenal pasti sepanjang ujian. Antara kelemahan utama ialah ketepatan klasifikasi bagi kelas Kemurungan dan Kecenderungan Bunuh Diri yang masih boleh diperbaiki, kerana terdapat kekeliruan antara emosi yang berlapis dan sukar dipisahkan secara literal. Di samping itu, sistem masih terhad kepada teks dalam Bahasa Inggeris sahaja, menjadikannya kurang inklusif untuk pengguna tempatan yang menggunakan Bahasa Melayu atau dialek lain. Dari sudut paparan hasil, walaupun maklum balas secara visual telah dipaparkan, sistem belum memberikan penjelasan mendalam tentang maksud setiap label emosi, yang boleh menyukarkan pemahaman pengguna baharu. Akhir sekali, sistem ini masih bergantung kepada model statik dan tidak menyokong pembelajaran berterusan berdasarkan input pengguna sebenar dari semasa ke semasa.

6.0 PENGHARGAAN

Geran Fakulti Teknologi dan Sains Maklumat FTSM 1 dan Geran Mutiara A195632.

7.0 RUJUKAN

Adarsh, R., Senthil Kumar, M., & Elhoseny, M. (2022). Hybrid deep learning model for mental health prediction using XGB and BERT. *Journal of Ambient Intelligence and Humanized Computing*, 13(2), 1293–1305.
<https://doi.org/10.1007/s12652-021-03120-5>

- Arun, N. (2024). Ensemble-based classification of psychological distress from social media texts. *Procedia Computer Science*, 227, 410–416
- Babu, T., & Kanaga, E. (2021). Prediction of mental health using machine learning techniques. *Turkish Journal of Computer and Mathematics Education (TURCOMAT)*, 12(2), 258–265
- Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. arXiv preprint arXiv:1810.04805. <https://doi.org/10.48550/arXiv.1810.04805>
- Gonzalo A. Ruz, Rubilar, R., & Rojas, C. (2020). A hybrid BERT + SVM model for detecting mental health issues in social media content. *IEEE Latin America Transactions*, 18(8), 1389–1396
- Jain, M., Kumar, A., & Shrivastava, A. (2021). A transformer-based model for mental health analysis using Reddit data. *Procedia Computer Science*, 187, 324–329
- Kastrati, Z., Imran, A. S., & Kurti, A. (2021). Sentiment analysis of students' feedback using deep learning and lexicon-based models. *Journal of Information Science*, 47(4), 517–536
- Kokab, S. T., Usama, M., & Asim, M. (2022). A comparative analysis of deep learning models for mental health detection from social media text. *Journal of Intelligent & Fuzzy Systems*, 43(5), 6103–6113. <https://doi.org/10.3233/JIFS-219313>
- Krichen, M., Hadrich, S., & Ben Ahmed, M. (2022). A comprehensive machine learning lifecycle framework for mental health monitoring based on text analysis. In M. Ben

Ahmed, A. Boudhir, A. M. Riad, & M. Ali (Eds.), *Advances in Data Science and Information Engineering* (pp. 17–30). Springer.
https://doi.org/10.1007/978-981-16-5314-3_3

Lyua, Y. W., Zhao, H., & Wang, L. (2020). Detection of depression using support vector machines in social media data. *International Journal of Environmental Research and Public Health*, 17(15), 5504

Mental Health, Brain Health and Substance Use (MSD). (2022). World mental health report: Transforming mental health for all. World Health Organization.
<https://www.who.int/publications/i/item/9789240049338>

Moez Krichen, Alaeddine Mihoub, Mohammed Y. Alzahrani, Hamilton Wilfrien Yves Adoni, & Tarik Nahhal. (2022). Are formal methods applicable to machine learning and artificial intelligence?
https://www.researchgate.net/publication/362485323_Are_Formal_Methods_Applicable_To_Machine_Learning_And_Artificial_Intelligence

Mohamed, I., Hussin, M. Y., & Ariffin, M. A. M. (2023). Depression detection using Reddit dataset via traditional machine learning classifiers. *Journal of Information and Communication Technology*, 22(1), 85–101

Nguyen, T., Do, D., & Nguyen, H. (2023). Detecting mental health problems in social media with hybrid deep learning. *Expert Systems with Applications*, 222, 119528.
<https://doi.org/10.1016/j.eswa.2023.119528>

Nor Shafrin Ahmad, Abdul Hamid, S. R., Abu Bakar, J., & Abu Bakar, N. F. (2021). Digital anxiety in Malaysia: A case study during the COVID-19 pandemic. *International*

Journal of Academic Research in Business and Social Sciences, 11(6), 119–132.
<https://doi.org/10.6007/IJARBSS/v11-i6/10041>

- Pavitra Sankar, V., Kumar, R., & Sharma, A. (2024). Ensemble learning for mental health prediction using social media data. *International Journal of Advanced Computer Science and Applications*, 15(3), 150–158
- Rabani, M., Rezaei, M., & Alipour, A. (2020). Mental health analysis via CNN and ensemble deep learning: A case study using Baghdad Science Journal dataset. *Baghdad Science Journal*, 17(1), 90–99
- Salur, M. U., & Aydin, I. (2020). Detection of mental disorders on social media using CNN and Word2Vec. *International Journal of Computational Intelligence Systems*, 13(1), 1203–1213
- Sattar, M., Hussain, T., Alsaeedi, A., Alghamdi, A., Alhumoud, S., & Ullah, S. (2021). A machine learning-based sentiment analysis framework for health information in COVID-19 tweets. *IEEE Access*, 9, 149425–149436.
<https://doi.org/10.1109/ACCESS.2021.3125151>
- Sayyida Tabinda Kokab, Sohail Asghar, & Shehneela Naz. (2022). Transformer-based deep learning models for the sentiment analysis of social media data. *Science Direct*, 14.
<https://www.sciencedirect.com/science/article/pii/S2590005622000224>
- Shrestha, S., Bhattarai, B., & Shrestha, A. (2020). Sentiment classification for mental health using Naïve Bayes. *International Journal of Computer Applications*, 176(34), 22–26

Yichao Wu, Zhou, Y., & Lin, C. (2024). BERT-based deep learning models for mental health classification on Reddit posts. *Procedia Computer Science*, 213, 298–304

Zhu, Z., Liu, X., & Huang, L. (2024). Depression detection using LSTM networks on Reddit mental health posts. *IEEE Access*, 12, 12244–12253

Zogan, M., Kaya, D., & Cicekli, I. (2022). CNN-based emotion classification on mental health tweets. *Computers in Biology and Medicine*, 147, 105784

Vilaasini A/P Kumar (A195632)

Assoc. Prof. Dr. Sabrina Tiun

Fakulti Teknologi & Sains Maklumat

Universiti Kebangsaan Malaysia