

PENGESANAN BERITA PALSU MENGGUNAKAN LSTM DAN VARIASINYA DALAM PEMROSESAN BAHASA SEMULA JADI

Ng Pin Xian, Salwani Abdullah

*Fakulti Teknologi & Sains Maklumat, Universiti Kebangsaan Malaysia,
43600 UKM Bangi, Selangor Darul Ehsan, Malaysia*

Abstrak

Berita palsu merupakan maklumat tidak sah yang disebarkan dengan tujuan untuk mengelirukan, memanipulasi, atau mempengaruhi persepsi masyarakat terhadap sesuatu isu atau entiti. Penyebaran berita palsu bukan sahaja menimbulkan kekeliruan dalam kalangan masyarakat, malah boleh mengancam kestabilan sosial, politik, dan ekonomi sesebuah negara. Oleh yang demikian, adalah penting untuk membangunkan satu sistem pengesanan yang mampu mengenal pasti dan mengklasifikasikan berita palsu secara automatik dan berkesan. Memandangkan kesukaran untuk mengesan kandungan palsu secara manual akibat bilangan artikel yang besar dan sifat kandungan yang kompleks, objektif utama projek ini adalah untuk memperkenalkan penggunaan model pembelajaran mendalam dalam mengklasifikasikan dua jenis berita iaitu 'Berita Palsu' dan 'Berita Benar'. Dataset yang digunakan diperoleh daripada laman web Kaggle, yang mengandungi sebanyak 72,134 teks berita berlabel yang merangkumi pelbagai topik yang mencerminkan kandungan dunia sebenar. Projek ini menilai keberkesanan empat model pembelajaran mendalam dalam pengesanan berita palsu iaitu Vanilla LSTM, Stacked LSTM, Bidirectional LSTM, dan Attention-based LSTM, berdasarkan metrik prestasi seperti ketepatan, kepekaan, kekhususan, dan skor-F1. Bagi projek ini, hasil terakhir adalah model klasifikasi berita yang dibangunkan menggunakan Flask. Antara model yang dibandingkan, Attention-based LSTM muncul sebagai model paling tepat, diikuti oleh Bidirectional LSTM, manakala Stacked LSTM dan Vanilla LSTM mempamerkan prestasi yang sederhana.

Kata Kunci: Berita Palsu, LSTM.

Abstract

The Fake news refers to false or misleading information that is deliberately disseminated to deceive, manipulate, or influence public perception regarding specific issues or entities. The spread of fake news not only creates confusion among the public but can also threaten the

social, political, and economic stability of a country. Therefore, it is crucial to develop a detection system capable of identifying and classifying fake news automatically and effectively. Given the challenges of manually detecting fake content due to the vast number of articles and the complexity of language used, the main objective of this project is to introduce the use of deep learning models to classify two types of news: 'Fake News' and 'Real News'. The dataset used in this project was obtained from Kaggle, consisting of 72,134 labeled news texts covering a wide range of topics that reflect real-world content. This project evaluates the effectiveness of four deep learning models in detecting fake news, namely Vanilla LSTM, Stacked LSTM, Bidirectional LSTM, and Attention-based LSTM, based on performance metrics such as accuracy, recall, precision and F1-score. The final outcome of the project is a news classification model integrated into a web application built using Flask. Among the models compared, the Attention-based LSTM achieved the highest accuracy, followed by Bidirectional LSTM, while Stacked LSTM and Vanilla LSTM demonstrated moderate performance.

Keywords: Fake News, LSTM

1.0 PENGENALAN

Penyebaran berita palsu dalam talian semakin menjadi isu global yang serius, terutamanya dengan kemunculan media sosial sebagai sumber utama maklumat masyarakat moden. Berita palsu boleh didefinisikan sebagai maklumat yang tidak benar atau mengelirukan yang disebarkan dengan tujuan untuk memanipulasi persepsi awam demi kepentingan tertentu seperti politik, kewangan, atau sosial (Zhang & Ghorbani, 2020). Fenomena ini semakin berleluasa disebabkan ketiadaan mekanisme kawalan fakta di platform digital serta kebolehan pengguna biasa untuk menyebarkan maklumat secara meluas tanpa pengesanan (Ahmed, 2017). Hal ini menjadikan pengguna – khususnya golongan pelajar dan warga emas – terdedah kepada kekeliruan dalam mengenal pasti maklumat yang sahih.

Dalam konteks ini, pembelajaran mendalam (*deep learning*) menawarkan potensi besar dalam membantu menangani penyebaran berita palsu. Teknologi ini membolehkan pembangunan model kecerdasan buatan yang dapat memahami struktur ayat, konteks semantik, serta pola linguistik yang kompleks dalam kandungan teks. Salah satu model yang berkesan dalam pemprosesan bahasa semula jadi (NLP) ialah *Long Short-Term Memory* (LSTM) – sejenis rangkaian neural berulang (RNN) yang direka untuk menganalisis hubungan jangka panjang dalam data teks (Zhang & Ghorbani, 2020). Model ini telah digunakan secara meluas dalam pelbagai aplikasi klasifikasi teks, termasuk pengesanan emosi, analisis sentimen, dan berita palsu.

Kajian ini memberi tumpuan kepada pembangunan sistem pengesanan berita palsu secara automatik menggunakan beberapa varian model LSTM iaitu *Vanilla LSTM*, *Stacked LSTM*, *Bidirectional LSTM* dan *Attention-based LSTM*. Model-model ini dipilih berdasarkan

keupayaan masing-masing dalam menangkap maklumat kontekstual dan menyusun semula struktur semantik yang tersembunyi dalam kandungan teks yang panjang dan kompleks. Tambahan pula, penggunaan attention mechanism membolehkan model memfokuskan kepada bahagian penting dalam ayat yang menyumbang kepada ketepatan klasifikasi.

Data yang digunakan dalam projek ini diperoleh daripada platform Kaggle, yang mengandungi lebih 72,000 artikel berita yang telah dilabel sebagai 'sahih' atau 'palsu'. Data ini akan melalui proses pra-pemprosesan teks yang melibatkan penyingkiran stopwords, normalisasi, tokenisasi dan pengkodan, sebelum dimasukkan ke dalam model pembelajaran. Kaedah pembelajaran terselia digunakan untuk melatih model dalam mengenal pasti pola-pola linguistik yang berkaitan dengan kandungan palsu dan sah.

Hasil daripada pembangunan sistem ini akan diintegrasikan ke dalam aplikasi berasaskan Flask, membolehkan pengguna menguji sendiri keberkesanan model dalam mengesan kandungan palsu secara langsung. Sistem ini diharap dapat membantu pengguna umum dan pihak berkepentingan dalam mengenal pasti berita palsu dengan lebih cepat dan berkesan, sekali gus menyumbang ke arah masyarakat yang lebih celik maklumat dan kritis dalam menilai kandungan digital.

2.0 KAJIAN LITERATUR

Pelbagai fenomena penyebaran berita palsu dalam era digital kini menjadi isu global yang membimbangkan, didorong oleh pertumbuhan pesat media sosial dan komunikasi maya. Berita palsu disebarkan secara meluas tanpa semakan fakta, memanipulasi persepsi masyarakat terhadap isu-isu sosial, politik, dan kesihatan awam. Akibatnya, kredibiliti maklumat dalam talian semakin terhakis, terutama apabila pengguna tidak dapat membezakan antara maklumat sah dan palsu (Hasan et al., 2023). Dalam konteks ini, keperluan terhadap sistem pengesanan berita palsu yang efisien dan tepat adalah sangat penting.

Bidang pengesanan berita palsu kini mendapat perhatian dalam penyelidikan kecerdasan buatan (AI) dan pemprosesan bahasa semula jadi (NLP). Pendekatan tradisional menggunakan teknik pembelajaran mesin seperti *Naive Bayes*, *Support Vector Machines* (SVM), dan Regresi Logistik telah digunakan bersama ciri-ciri linguistik seperti TF-IDF dan *n-gram* untuk mengenal pasti pola dalam kandungan berita (Hossain et al., 2023). Namun begitu, model-model ini mempunyai batasan dari segi pemahaman konteks semantik dan hubungan jangka panjang dalam struktur ayat.

Bagi mengatasi kekangan ini, penyelidik beralih kepada pembelajaran mendalam seperti *Long Short-Term Memory* (LSTM), *Gated Recurrent Unit* (GRU), dan terutamanya varian *Bidirectional LSTM* (Bi-LSTM) yang memproses data dari dua arah, sekali gus membolehkan model memahami makna perkataan secara menyeluruh dalam ayat (Kumar & Mehta, 2023).

Selain itu, penggunaan embeddings seperti FastText dan Word2Vec turut memperkuat prestasi model dengan membantu dalam pengekstrakan makna mendalam (Alzahrani et al., 2023).

Pendekatan LSTM telah digunakan secara meluas dalam pembangunan sistem pengesanan automatik, contohnya dalam kajian Raj & Kumari (2023), yang menggabungkan LSTM dengan mekanisme perhatian (*attention*) untuk menumpukan fokus pada bahagian penting dalam teks. *Attention-based LSTM* menunjukkan peningkatan ketara dalam ketepatan klasifikasi berbanding model lain. Sementara itu, pendekatan terkini juga melibatkan penggunaan transformer seperti BERT dan RoBERTa yang dapat memahami konteks global dalam ayat, menjadikannya sangat berkesan dalam tugas klasifikasi berita (Shivaprasad et al., 2023).

Dalam landskap teknologi sebenar, beberapa penyelidikan telah mencadangkan integrasi pembelajaran federated dan analitik rangkaian sosial untuk mengenal pasti penyebaran kandungan mencurigakan di platform seperti *Twitter* dan *Reddit* (Mahfouz et al., 2023). Selain itu, pendekatan hibrid yang menggabungkan CNN, Bi-LSTM, dan attention telah menghasilkan prestasi tinggi dengan ketepatan mencapai 0.9766, terutama dalam konteks berita berkaitan COVID-19 (Wang et al., 2023). Terdapat juga usaha penerokaan terhadap penggunaan teknologi blockchain sebagai pengesahan sumber maklumat serta *Explainable AI* (XAI) bagi meningkatkan kepercayaan pengguna terhadap sistem yang dibangunkan (Li et al., 2024).

Beberapa kajian lepas telah menganalisis dan membandingkan pelbagai model pembelajaran mendalam, termasuk *Vanilla LSTM*, *Bi-LSTM*, *Stacked LSTM* dan *Attention-based LSTM*. *Vanilla LSTM* menawarkan kestabilan dan kecekapan pemprosesan, manakala *Bi-LSTM* memberikan konteks dua hala yang lebih tepat tetapi menuntut sumber komputasi lebih besar (Anand, 2023). *Stacked LSTM* pula sesuai untuk mengenal pasti ciri-ciri berlapis, namun berisiko kepada masalah *overfitting* jika set data kecil (Alzahrani et al., 2023). *Attention-based LSTM* terbukti memberikan prestasi klasifikasi tertinggi, namun memerlukan pelaksanaan teknikal yang lebih kompleks.

Dalam kajian perbandingan, *Vanilla LSTM* mencatatkan ketepatan sehingga 0.974, diikuti oleh *Attention LSTM* (0.969), *Bi-LSTM* (0.9516), dan *Stacked LSTM* (0.9510) (Chauhan, 2023). Sementara itu, model transformer seperti BERT mencatat ketepatan 92% dan skor recall 91%, menunjukkan keunggulan dalam memahami semantik dan konteks global teks (Zhang & Ghorbani, 2020). Gabungan pendekatan satu dimensi dengan pelbagai modaliti seperti teks, imej dan metadata turut menghasilkan ketepatan yang lebih tinggi dalam klasifikasi berita palsu (Rani et al., 2023).

Kesemua dapatan ini menunjukkan bahawa tiada satu pendekatan tunggal yang terbaik. Pemilihan model perlu disesuaikan mengikut keperluan aplikasi sebenar seperti saiz dataset,

keperluan masa nyata, dan keupayaan komputasi. Kajian ini memilih untuk meneroka keberkesanan keempat-empat model LSTM iaitu *Vanilla*, *Bi-LSTM*, *Stacked LSTM*, dan *Attention-based LSTM* dalam mengenal pasti berita palsu, serta menilai prestasi mereka dari aspek ketepatan, kecekapan masa dan kebolegunaan dalam sistem sebenar. Pendekatan ini diharap dapat menyumbang ke arah pembangunan sistem pengesanan berita palsu yang lebih mantap dan responsif terhadap keperluan masyarakat maklumat masa kini.

3.0 METODOLOGI

Kajian ini merangkumi keseluruhan proses pembangunan sistem pengesanan berita palsu berasaskan pembelajaran mendalam, bermula daripada analisis keperluan, mereka bentuk model konseptual, pembangunan aplikasi menggunakan kerangka Flask, pengujian kebolegunaan, hinggalah kepada analisis hasil model.

3.1 Reka Bentuk Model

Reka bentuk senibina sistem pengesanan berita palsu dalam projek ini menggunakan kerangka kerja CRISP-DM (*Cross-Industry Standard Process for Data Mining*). CRISP-DM adalah metodologi yang terbukti berkesan dalam membimbing proses pembangunan model data mining secara sistematik dan teratur. Kerangka kerja ini terdiri daripada enam fasa utama, iaitu Pemahaman Perniagaan (Business Understanding), Pemahaman Data (Data Understanding), Penyediaan Data (Data Preparation), Pemodelan (Modelling), Penilaian (Evaluation) dan Penyebaran (Deployment). Setiap fasa mempunyai peranan yang penting dalam memastikan projek ini dilaksanakan dengan jayanya.



Rajah 1: Fasa CRISP-DM

Fasa Pemahaman Data adalah langkah pertama CRISP-DM. Fasa ini bertujuan untuk memahami objektif perniagaan dan keperluan projek. Dalam konteks sistem pengesanan berita palsu, fasa ini melibatkan penentuan matlamat projek, seperti meningkatkan ketepatan pengesanan berita palsu dan menyediakan alat yang mesra pengguna. Fasa ini juga melibatkan

penentuan metrik kejayaan, seperti ketepatan klasifikasi dan masa pemprosesan. Dalam kajian ini, fasa ini memastikan sistem yang dibangunkan dapat menangani masalah penyebaran berita palsu dengan berkesan dan memberikan impak positif kepada masyarakat (Chapman et al., 2000).

Fasa pemahaman data ialah langkah kedua dalam CRISP-DM, di mana ia tertumpu pada penerokaan dan analisis data yang dikumpul untuk mendapatkan pemahaman yang lebih baik tentang kandungan, struktur, dan kualitinya. Sumber data yang digunakan untuk penyelidikan ini ialah data berita palsu yang diperoleh daripada laman web Kaggle.com. Dataset ini mengandungi teks berita yang dikategorikan sebagai benar atau palsu. Fail set data yang dikumpulkan datang dalam format .csv, yang mengandungi maklumat seperti tajuk berita, kandungan artikel, dan label klasifikasi (benar atau palsu). Data akan diproses untuk menghapuskan entri yang tidak diformatkan dengan betul, menangani nilai kosong, serta memastikan kualiti data yang tinggi sebelum digunakan dalam analisis lanjut. Atribut daripada set data mentah akan ditunjukkan dalam Jadual 1:

Jadual 1: Set Data Atribut

Atribut	Format	Penerangan
Index	Angka/Nombor	Nombor indeks unik bagi setiap baris data, yang berfungsi sebagai pengecam untuk setiap artikel berita dalam dataset.
Title	Teks	Tajuk artikel berita yang boleh memberikan gambaran ringkas mengenai kandungan berita tersebut
Text	Teks	Kandungan penuh artikel berita yang akan dianalisis untuk menentukan kesahihan maklumat yang disampaikan.
Label	Angka/Nombor	Kategori berita di mana nilai "0" mewakili berita palsu (<i>fake news</i>) dan "1" mewakili berita benar (<i>real news</i>).

Fasa penyediaan data ialah fasa yang melibatkan pemprosesan data untuk menghasilkan set data bersih yang sesuai untuk digunakan dalam model pembelajaran mesin. Langkah-langkah pemprosesan data seperti penerokaan data, pembersihan teks, pemadanan atribut, penormalan data, dan penyingkiran unsur yang tidak relevan akan diterangkan dalam bahagian ini. Dalam konteks pengesanan berita palsu, data yang bersih adalah bebas daripada nilai kosong, pendua, serta elemen yang boleh mengganggu analisis seperti tanda baca berlebihan, nombor rawak, atau kata-kata yang tidak bermakna. Oleh itu, fasa penyediaan data adalah langkah penting dalam proses pembelajaran mesin kerana ia memastikan data bersih, konsisten,

dan sesuai untuk membina model pengesanan berita palsu yang lebih tepat dan boleh dipercayai.

Dalam fasa pemodelan, model pembelajaran mesin dibangunkan dan diuji untuk mengenal pasti corak dan perhubungan dalam data. Fasa ini melibatkan pemilihan model yang sesuai, latihan model menggunakan set latihan, dan penilaian prestasi model menggunakan set ujian. Model yang akan digunakan dalam penyelidikan ini ialah *Vanilla Long-Short Term Memory*, *Bi-Directional Long-Short Term Memory (Bi-LSTM)*, *Attention Based Long-Short Term Memory (Attention-Based LSTM)* and *Stacked Long-Short Term Memory (Stacked LSTM)*.

Fasa penilaian adalah langkah kritikal dalam proses pembangunan model pembelajaran mesin untuk pengesanan berita palsu. Pada fasa ini, prestasi model dinilai menggunakan metrik yang sesuai untuk memastikan model dapat membuat generalisasi yang baik kepada data baharu yang tidak dilihat sebelum ini. Penilaian ini membantu menentukan sama ada model telah mencapai objektif yang ditetapkan dan mengenal pasti bidang yang perlu diperbaiki.

Dalam fasa penggunaan model, model klasifikasi berita palsu yang dibangunkan menggunakan pendekatan LSTM dan variannya telah diintegrasikan ke dalam sebuah aplikasi web menggunakan Flask. Aplikasi ini membolehkan pengguna memasukkan teks berita dan menerima ramalan kebarangkalian sama ada berita tersebut adalah benar atau palsu, bersama-sama paparan keputusan yang mesra pengguna. Model yang telah dilatih dan disimpan digunakan semula (*inference*) untuk menjana prediksi secara langsung, di mana setiap input teks akan melalui proses pembersihan dan penukaran kepada bentuk berangka sebelum dihantar ke model. Sistem ini bertujuan untuk memudahkan pengguna bukan teknikal membuat penilaian pantas terhadap kesahihan sesuatu berita dalam situasi dunia sebenar.

4.0 HASIL

4.1 Keputusan

Bahagian ini membentangkan hasil penilaian metrik untuk model pengesanan, membandingkan keputusan penilaian metrik model *Vanilla LSTM*, *Bidirectional LSTM*, *Attention-based LSTM*, dan *Stacked LSTM* dipaparkan dalam bentuk jadual dan rajah.

Penilaian *Vanilla LSTM*

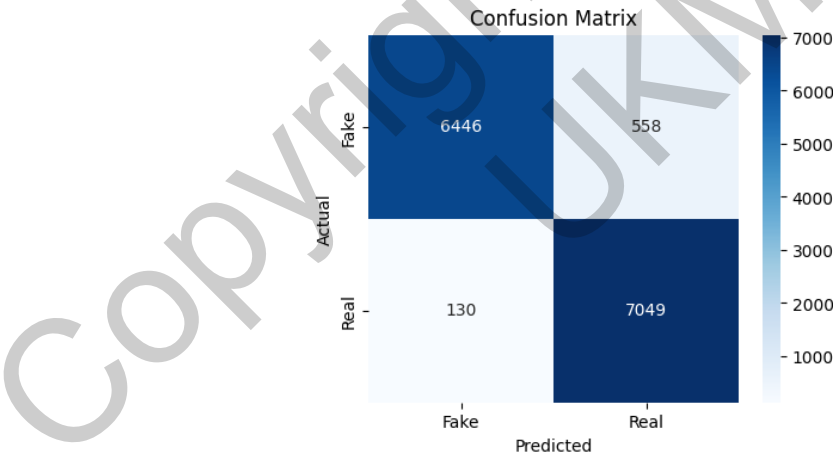
Jadual 2 menunjukkan keputusan penilaian metrik bagi model *Vanilla LSTM*. Rajah 2 memaparkan matriks kekeliruan keputusan penilaian model *Vanilla LSTM*. Jadual 3 menunjukkan prestasi model dalam kejituan dan kehilangan latihan dan ujian manakala Rajah 3 menunjukkan graf kejituan latihan dan ujian model *Vanilla LSTM* dan Rajah 4 menunjukkan graf kehilangan latihan dan ujian model *Vanilla LSTM*.

Jadual 2: Keputusan Penilaian Metrik Model *Vanilla LSTM*

Label	Penilaian Metrik			
	Ketepatan	Kepekaan	Skor-F1	Sokongan
0	0.98	0.92	0.95	7004
1	0.93	0.98	0.95	7179
Ketepatan: 0.95				

Berdasarkan Jadual 4.1, model *Vanilla LSTM* menunjukkan prestasi yang seimbang dalam mengklasifikasikan kedua-dua label 0 dan 1. Bagi label 0, model mencapai ketepatan (*precision*) sebanyak 0.98, kepekaan (*recall*) sebanyak 0.92, dan skor F1 sebanyak 0.95, dengan 7004 sokongan. Ini menunjukkan model cenderung untuk menghasilkan sedikit negatif palsu dalam mengenal pasti kelas 0, menyebabkan kepekaan sedikit lebih rendah berbanding ketepatan. Manakala untuk label 1, model menunjukkan kepekaan yang sangat tinggi iaitu 0.98, dengan ketepatan 0.93, dan skor F1 0.95, serta sokongan sebanyak 7179. Keupayaan mengesan hampir semua sampel sebenar kelas 1 menjadikan model ini amat sesuai dalam konteks di mana pengesanan kes positif adalah kritikal.

Secara keseluruhan, model *Vanilla LSTM* mencatatkan ketepatan keseluruhan sebanyak 0.95, menjadikannya model yang stabil dan boleh dipercayai untuk tugas klasifikasi binari ke atas set data ini.



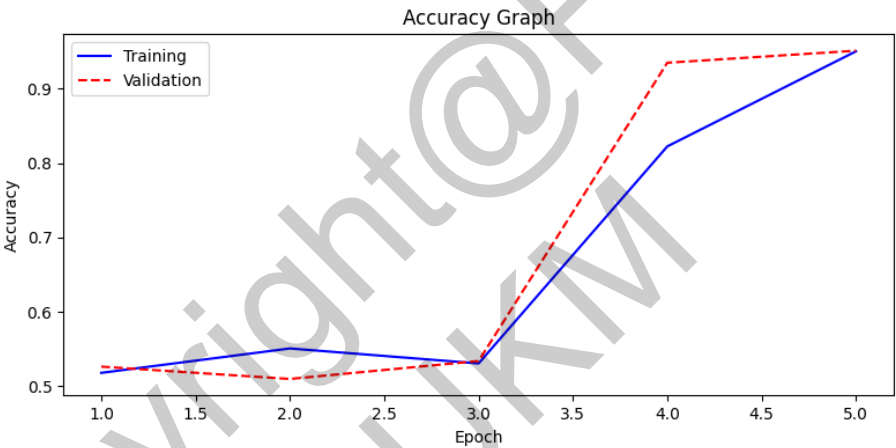
Rajah 2: Matriks Kekeliruan Model *Vanilla LSTM*

Berdasarkan matriks kekeliruan di atas, Model *Vanilla LSTM* menunjukkan prestasi yang baik dalam mengklasifikasikan kedua-dua kategori “Palsu” dan “Benar”. Sebanyak 6446 sampel “Palsu” berjaya diklasifikasikan dengan betul, manakala 558 sampel “Palsu” telah salah diklasifikasikan sebagai “Benar”. Bagi kategori “Benar”, model berjaya mengklasifikasikan 7049 sampel dengan betul, namun terdapat 130 sampel “Benar” yang telah tersilap diklasifikasikan sebagai “Palsu”. Ini menunjukkan ketepatan yang tinggi dalam mengesan berita benar serta keupayaan yang kukuh dalam mengenal pasti berita palsu.

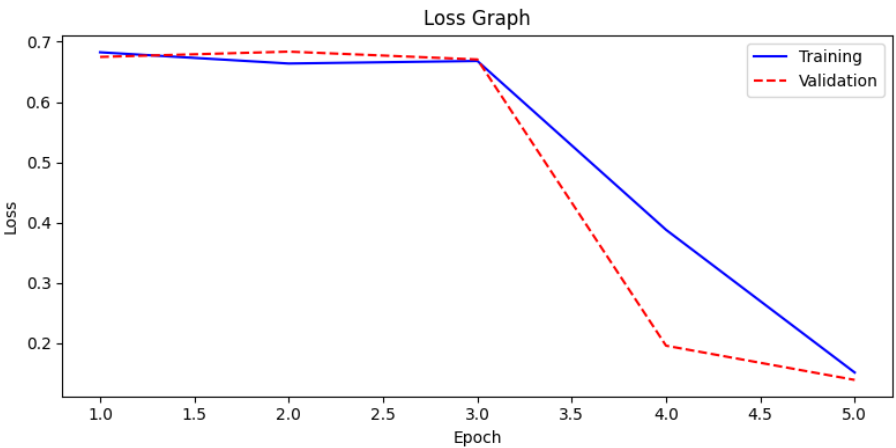
Secara keseluruhan, model ini mencatat prestasi yang seimbang dan stabil dengan jumlah klasifikasi betul sebanyak 13,495 daripada 14,183 sampel keseluruhan, menandakan ketepatan keseluruhan model adalah tinggi. Matriks ini menunjukkan bahawa model LSTM adalah sangat berkesan dalam mengurangkan kesilapan klasifikasi dan boleh digunakan secara praktikal untuk tugas pengesanan berita palsu.

Jadual 3: Kehilangan dan Ketepatan Latihan dan Ujian Model *Vanilla LSTM*

Epoch	Kehilangan Latihan	Ketepatan Latihan	Kehilangan Uji	Ketepatan Uji
1	0.6870	0.5092	0.6748	0.5263
2	0.6708	0.5322	0.6836	0.5097
3	0.6719	0.5213	0.6706	0.5339
4	0.4726	0.7574	0.1956	0.9350
5	0.1756	0.9436	0.1388	0.9513



Rajah 3: Graf Ketepatan Latihan dan Ujian Model *Vanilla LSTM*



Rajah 4: Graf Kehilangan Latihan dan Ujian Model *Vanilla LSTM*

Berdasarkan Jadual 3, model Vanilla LSTM telah dilatih dan diuji selama lima epoch. Pada epoch pertama, model mencatatkan kehilangan yang tinggi serta ketepatan yang rendah bagi kedua-dua set data latihan dan ujian. Ini adalah situasi yang lazim kerana model masih berada pada peringkat awal pembelajaran. Seiring dengan peningkatan bilangan epoch, kehilangan secara konsisten menurun manakala ketepatan bertambah baik. Bermula dari epoch kedua, model menunjukkan kemajuan yang positif apabila kehilangan berkurangan dan ketepatan meningkat, menandakan model mula memahami corak dalam data. Prestasi model terus bertambah baik pada epoch ketiga dan keempat, di mana kehilangan menjadi semakin rendah dan ketepatan semakin tinggi, mencerminkan kemampuan model untuk membuat ramalan yang lebih tepat dan stabil. Pada epoch kelima, model mencapai prestasi terbaik dengan kehilangan paling rendah dan ketepatan paling tinggi untuk kedua-dua set data latihan dan ujian.

Secara keseluruhan, model LSTM menunjukkan trend peningkatan yang konsisten sepanjang proses latihan, membuktikan bahawa model berjaya mempelajari pola daripada data secara berkesan dan mampu mengaplikasikan pengetahuan tersebut ke atas data ujian. Ini menunjukkan model mempunyai keupayaan generalisasi yang kukuh dalam tugas pengesanan berita palsu.

Penilaian Bi-LSTM

Jadual 4 menunjukkan keputusan penilaian metrik bagi model Bi-LSTM. Rajah 5 memaparkan matriks kekeliruan keputusan penilaian model Bi-LSTM. Jadual 5 menunjukkan prestasi model dalam kejituan dan kehilangan latihan dan ujian manakala Rajah 6 menunjukkan graf kejituan latihan dan ujian model Bi-LSTM dan Rajah 7 menunjukkan graf kehilangan latihan dan ujian model Bi-LSTM.

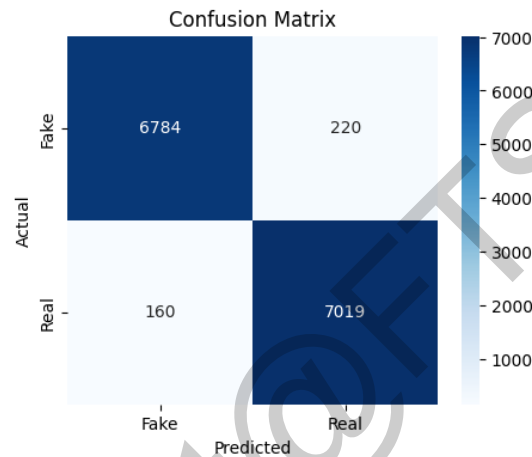
Jadual 4: Keputusan Penilaian Metrik Model Bi-LSTM

Label	Penilaian Metrik			
	Ketepatan	Kepekaan	Skor-F1	Sokongan
0	0.98	0.97	0.97	7004
1	0.97	0.98	0.97	7179
Ketepatan: 0.97				

Berdasarkan Jadual 4, model Bidirectional LSTM menunjukkan prestasi yang konsisten dan sangat baik dalam mengklasifikasikan kedua-dua label 0 dan 1. Untuk label 0, model mencapai ketepatan (*precision*) sebanyak 0.98, kepekaan (*recall*) sebanyak 0.97, dan Skor-F1 sebanyak 0.97, dengan sokongan sebanyak 7004 data. Ini menunjukkan bahawa model dapat mengenal pasti kelas 0 dengan baik, walaupun masih terdapat sedikit kes positif palsu yang menjejaskan kepekaan berbanding ketepatan. Sementara itu, bagi label 1, model mencatatkan ketepatan sebanyak 0.97, kepekaan sebanyak 0.98, dan Skor-F1 sebanyak 0.97, dengan sokongan sebanyak 7179. Kepekaan yang tinggi bagi kelas ini membuktikan kemampuan

model dalam mengurangi kadar negatif palsu dan menjadikan ramalan lebih tepat. Skor-F1 yang seimbang bagi kedua-dua kelas menunjukkan bahawa model tidak berat sebelah dan dapat mengekalkan prestasi klasifikasi yang stabil untuk kedua-dua kategori.

Secara keseluruhannya, model *Bidirectional LSTM* mencatatkan ketepatan keseluruhan sebanyak 0.97, mencerminkan prestasi yang kukuh dan stabil dalam tugas klasifikasi binari. Keupayaan model untuk memproses maklumat dalam kedua-dua arah urutan memberi kelebihan dalam memahami konteks ayat secara menyeluruh, sekali gus menyumbang kepada ketepatan klasifikasi yang lebih tinggi.



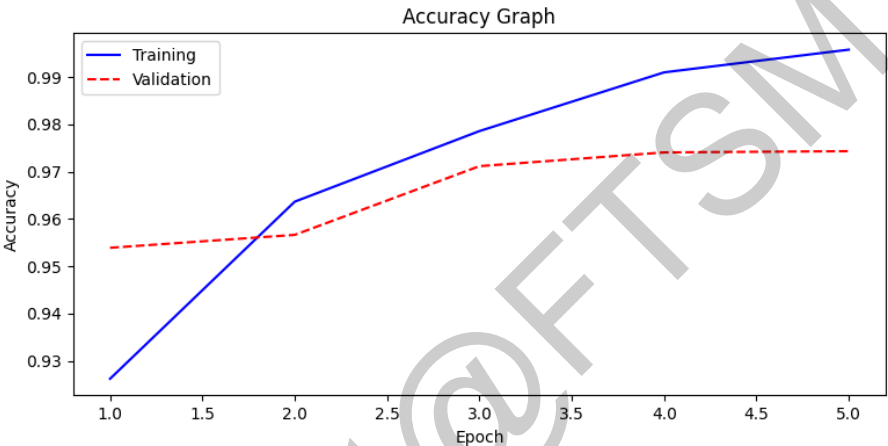
Rajah 5: Matriks Kekeliruan Model *Bi-LSTM*

Berdasarkan matriks kekeliruan di atas, model *Bi-LSTM* menunjukkan prestasi yang baik dalam mengklasifikasikan kedua-dua kategori “Palsu” dan “Benar”. Sebanyak 6784 sampel “Palsu” berjaya diklasifikasikan dengan betul, manakala 220 sampel “Palsu” telah salah diklasifikasikan sebagai “Benar”. Bagi kategori “Benar”, model berjaya mengklasifikasikan 7019 sampel dengan betul, namun terdapat 160 sampel “Benar” yang telah tersilap diklasifikasikan sebagai “Palsu”. Keputusan ini menunjukkan bahawa model mempunyai keupayaan yang kukuh dalam mengenal pasti kedua-dua jenis berita, dengan kadar positif palsu dan negatif palsu yang berada pada tahap yang rendah dan boleh diterima.

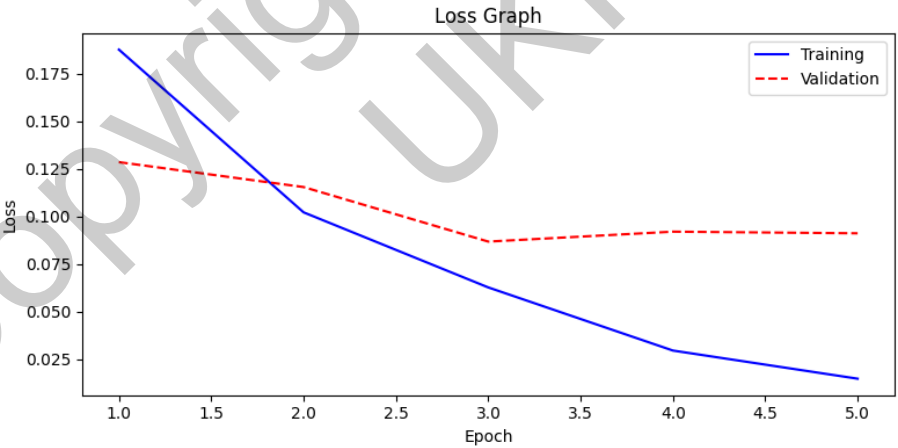
Secara keseluruhannya, model ini mencatat jumlah klasifikasi betul sebanyak 13,803 daripada 14,183 sampel keseluruhan, yang menandakan ketepatan keseluruhan model adalah tinggi dan seimbang. Matriks kekeliruan ini membuktikan bahawa model *Bi-LSTM* berfungsi dengan stabil dan boleh diandalkan, serta sesuai digunakan dalam aplikasi praktikal pengesanan berita palsu dalam Bahasa Inggeris.

Jadual 5: Kehilangan dan Ketepatan Latihan dan Ujian Model *Bi-LSTM*

Epoch	Kehilangan Latihan	Ketepatan Latihan	Kehilangan Uji	Ketepatan Uji
1	0.2678	0.8832	0.1286	0.9539
2	0.1038	0.9631	0.1155	0.9566
3	0.0696	0.9764	0.0868	0.9712
4	0.0299	0.9912	0.0921	0.9741
5	0.0127	0.9963	0.0912	0.9744



Rajah 6: Graf Ketepatan Latihan dan Ujian Model *Bi-LSTM*



Rajah 7: Graf Kehilangan Latihan dan Ujian Model *Bi-LSTM*

Berdasarkan Jadual 5, model *Bidirectional LSTM* telah dilatih dan diuji selama lima epoch. Pada epoch pertama, model mencatatkan kehilangan yang agak tinggi bagi data latihan, walaupun ketepatan ujian telah menunjukkan prestasi yang baik. Ini menunjukkan bahawa model *Bi-LSTM* dapat mempelajari ciri awal dalam data dengan lebih cepat berbanding model asas lain. Pada epoch kedua, berlaku penurunan ketara dalam kehilangan bagi kedua-dua set data latihan dan ujian, manakala ketepatan meningkat dengan jelas. Ini membuktikan bahawa

model sedang belajar dengan berkesan dan pantas menyesuaikan diri terhadap pola data. Menjelang epoch ketiga dan keempat, kehilangan latihan terus menurun dengan ketara dan ketepatan semakin menghampiri tahap maksimum. Namun begitu, sedikit peningkatan dalam kehilangan ujian dapat diperhatikan, yang menunjukkan kemungkinan berlakunya sedikit overfitting apabila model mula terlalu menyesuaikan diri dengan data latihan. Pada epoch kelima, kehilangan latihan berada pada tahap yang paling rendah dan ketepatan latihan mencapai nilai tertinggi. Ketepatan ujian juga kekal tinggi dan stabil, menunjukkan bahawa prestasi model tidak menurun secara drastik walaupun terdapat sedikit peningkatan dalam kehilangan ujian.

Secara keseluruhan, model *Bi-LSTM* menunjukkan keupayaan pembelajaran yang sangat baik dengan prestasi yang konsisten sepanjang proses latihan. Ini mencerminkan kemampuan model untuk mengenal pasti pola dengan cepat dan mengekalkan prestasi tinggi dalam tugas pengesanan berita palsu, sambil memperlihatkan daya generalisasi yang kukuh ke atas data ujian.

Penilaian *Attention-Based LSTM*

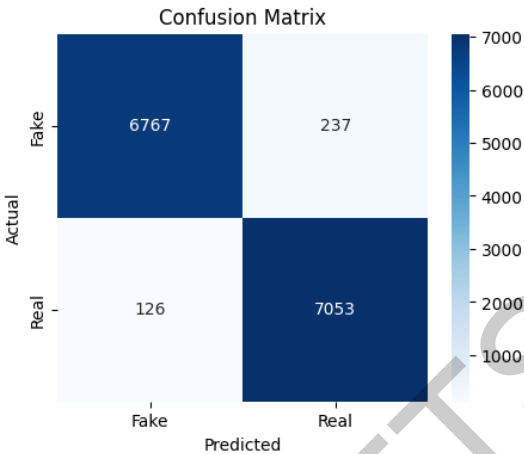
Jadual 6 menunjukkan keputusan penilaian metrik bagi model *Attention-based LSTM*. Rajah 8 memaparkan matriks kekeliruan keputusan penilaian model *Attention-based LSTM*. Jadual 7 menunjukkan prestasi model dalam kejituan dan kehilangan latihan dan ujian manakala Rajah 9 menunjukkan graf kejituan latihan dan ujian model *Attention-based LSTM* dan Rajah 10 menunjukkan graf kehilangan latihan dan ujian model *Attention-based LSTM*.

Jadual 6: Keputusan Penilaian Metrik Model *Attention-based LSTM*

Label	Penilaian Metrik			
	Ketepatan	Kepekaan	Skor-F1	Sokongan
0	0.98	0.97	0.97	7004
1	0.97	0.98	0.97	7179
Ketepatan: 0.97				

Berdasarkan Jadual 6, model *Attention-based LSTM* menunjukkan prestasi yang sangat tinggi dan seimbang dalam mengklasifikasikan kedua-dua label 0 dan 1. Untuk label 0, model mencapai ketepatan sebanyak 0.98, kepekaan 0.97, dan skor F1 sebanyak 0.97, dengan sokongan sebanyak 7004 data. Ini menunjukkan bahawa model sangat cekap dalam mengenal pasti berita palsu (label 0), dengan kadar positif palsu yang amat rendah. Bagi label 1, model mencatatkan ketepatan sebanyak 0.97, kepekaan 0.98, dan skor F1 0.97, dengan jumlah sokongan sebanyak 7179. Kepekaan yang tinggi membuktikan bahawa model berjaya mengesan majoriti berita benar (label 1) dengan tepat, manakala ketepatan yang tinggi menunjukkan bilangan positif palsu adalah sangat minimum. Skor F1 yang sama untuk kedua-dua kelas mencerminkan bahawa model mencapai keseimbangan yang baik antara ketepatan dan kepekaan tanpa cenderung kepada mana-mana kelas.

Secara keseluruhannya, model Attention-based LSTM mencatat ketepatan keseluruhan sebanyak 0.97, membuktikan keupayaannya untuk menangani masalah klasifikasi binari dengan tepat sambil memberi tumpuan kepada ciri penting dalam urutan input melalui mekanisme perhatian (*attention*).



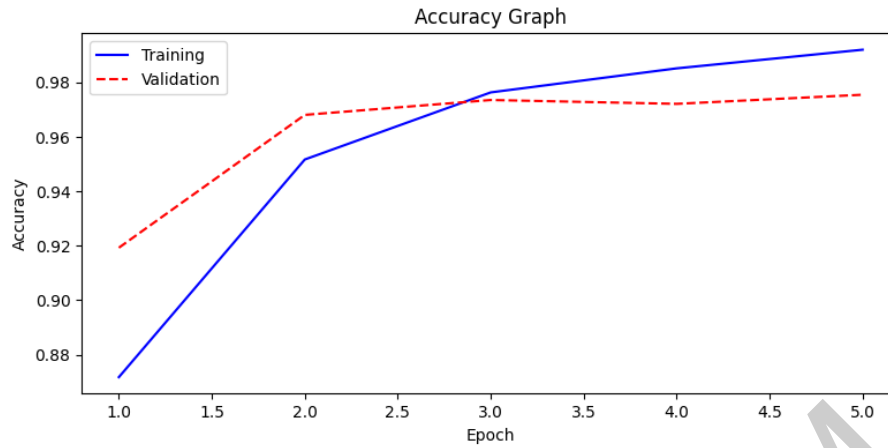
Rajah 8: Matriks Kekeliruan Model *Attention-based LSTM*

Berdasarkan matriks kekeliruan di atas, model *Attention-based LSTM* menunjukkan prestasi yang sangat baik dalam mengklasifikasikan kategori “Palsu” dan “Benar”. Sebanyak 6767 sampel “Palsu” telah berjaya diklasifikasikan dengan tepat, manakala hanya 237 sampel “Palsu” telah salah diklasifikasikan sebagai “Benar”. Bagi kategori “Benar”, model berjaya mengklasifikasikan 7053 sampel dengan betul, dan hanya 126 sampel “Benar” yang telah diklasifikasikan secara salah sebagai “Palsu”. Ini menunjukkan bahawa model bukan sahaja mempunyai kepekaan yang tinggi dalam mengenal pasti berita benar, malah juga mengekalkan ketepatan yang kukuh dalam mengesan berita palsu.

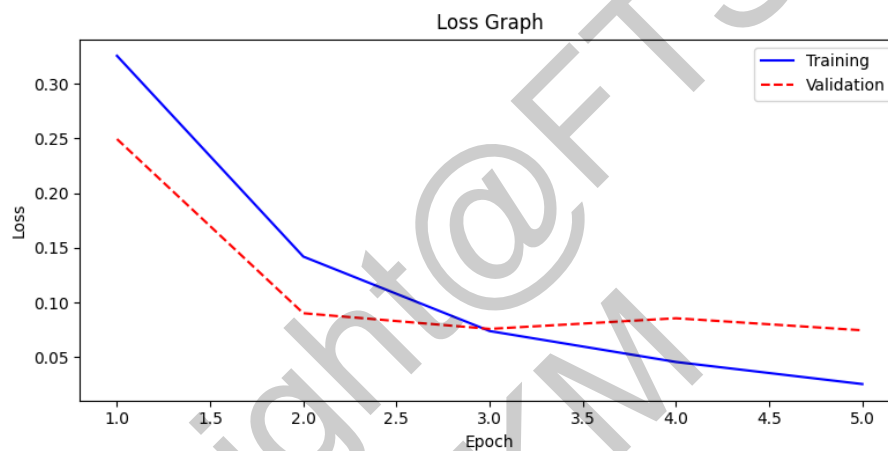
Secara keseluruhan, model ini menghasilkan jumlah klasifikasi betul sebanyak 13,820 daripada 14,183 sampel keseluruhan, yang merupakan prestasi yang sangat tinggi. Matriks kekeliruan ini membuktikan bahawa penggunaan mekanisme *attention* dalam model LSTM membantu meningkatkan ketepatan klasifikasi dengan memberi fokus kepada bahagian input yang lebih relevan, menjadikan model ini sangat sesuai untuk pengesanan berita palsu yang lebih halus dan kontekstual.

Jadual 7: Kehilangan dan Ketepatan Latihan dan Ujian Model *Attention-based LSTM*

Epoch	Kehilangan Latihan	Ketepatan Latihan	Kehilangan Uji	Ketepatan Uji
1	0.4263	0.8100	0.2494	0.9193
2	0.1640	0.9454	0.0901	0.9681
3	0.0751	0.9764	0.0758	0.9736
4	0.0473	0.9855	0.0855	0.9722
5	0.0257	0.9925	0.0745	0.9755



Rajah 9: Graf Ketepatan Latihan dan Ujian Model *Attention-based LSTM*



Rajah 10: Graf Kehilangan Latihan dan Ujian Model *Attention-based LSTM*

Berdasarkan Jadual 7, model *Attention-based LSTM* telah dilatih dan diuji selama lima epoch. Pada epoch pertama, model menunjukkan kehilangan yang sederhana bagi data latihan namun berjaya mencatat ketepatan ujian yang tinggi. Ini menunjukkan bahawa walaupun model masih dalam fasa awal pembelajaran, mekanisme attention telah membantu model mengenal pasti ciri penting dalam data dengan lebih cepat. Pada epoch kedua, terdapat penurunan ketara dalam kehilangan untuk kedua-dua data latihan dan ujian, serta peningkatan yang jelas dalam ketepatan. Ini menandakan bahawa model sedang belajar secara efektif dan memahami corak penting dalam data. Pada epoch ketiga dan keempat, kehilangan terus berkurangan manakala ketepatan semakin menghampiri tahap maksimum. Model berada dalam fasa stabil dan mampu mengekalkan prestasi klasifikasi yang tinggi. Menariknya, pada epoch kelima, walaupun kehilangan ujian sedikit meningkat pada epoch sebelumnya, ketepatan ujian mencapai nilai tertinggi sepanjang proses latihan. Ini menunjukkan bahawa model masih mampu mengekalkan keupayaan generalisasi yang kuat, walaupun terdapat sedikit turun naik dalam kehilangan.

Secara keseluruhan, model *Attention-based LSTM* menunjukkan peningkatan konsisten dalam ketepatan dan penurunan kehilangan sepanjang epoch. Prestasi stabil pada set data ujian membuktikan bahawa model bukan sahaja mampu belajar secara berkesan, malah dapat mengekalkan prestasi tinggi dalam tugas pengesanan berita palsu dengan bantuan mekanisme perhatian (*attention mechanism*).

Penilaian *Stacked LSTM*

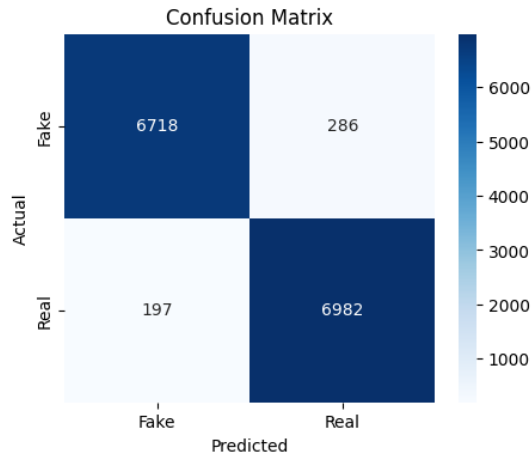
Jadual 8 menunjukkan keputusan penilaian metrik bagi model *Stacked LSTM*. Rajah 11 memaparkan matriks kekeliruan keputusan penilaian model *Stacked LSTM*. Jadual 9 menunjukkan prestasi model dalam kejituan dan kehilangan latihan dan ujian manakala Rajah 12 menunjukkan graf kejituan latihan dan ujian model *Stacked LSTM* dan Rajah 13 menunjukkan graf kehilangan latihan dan ujian model *Stacked LSTM*.

*Jadual 8: Keputusan Penilaian Metrik Model *Stacked LSTM**

Label	Penilaian Metrik			
	Ketepatan	Kepekaan	Skor-F1	Sokongan
0	0.97	0.96	0.97	7004
1	0.96	0.97	0.97	7179
Ketepatan: 0.97				

Berdasarkan Jadual 8, model *Stacked LSTM* menunjukkan prestasi yang seimbang dalam klasifikasi kedua-dua label 0 dan 1. Untuk label 0, model mencatatkan ketepatan, kepekaan, dan Skor-F1 yang tinggi iaitu 0.97, 0.96, dan 0.97, dengan sokongan sebanyak 7004 data. Ini menunjukkan bahawa model dapat mengenal pasti berita palsu dengan kadar ketepatan yang sangat baik serta mampu mengekalkan keseimbangan antara kes positif benar dan positif palsu. Bagi label 1, model juga mencatatkan ketepatan sebanyak 0.96, kepekaan 0.97, dan Skor-F1 sebanyak 0.97, dengan jumlah sokongan sebanyak 7179. Ini membuktikan bahawa model berupaya untuk membuat klasifikasi yang tepat terhadap berita benar dengan hanya sedikit kesilapan klasifikasi. Konsistensi dalam nilai metrik untuk kedua-dua label menandakan bahawa model *Stacked LSTM* berfungsi dengan stabil dan tidak menunjukkan berat sebelah terhadap mana-mana kelas.

Secara keseluruhan, model ini mencapai ketepatan keseluruhan sebanyak 0.97, menjadikannya satu model yang kukuh dan kompetitif, hampir menyamai prestasi model-model lain seperti *Attention-based LSTM* atau *Bi-LSTM*. Walaupun struktur berlapis menjadikan model ini lebih kompleks, ia memberikan kelebihan dalam pembelajaran ciri mendalam, menjadikannya relevan dalam analisis prestasi bagi tugas pengesanan berita palsu.



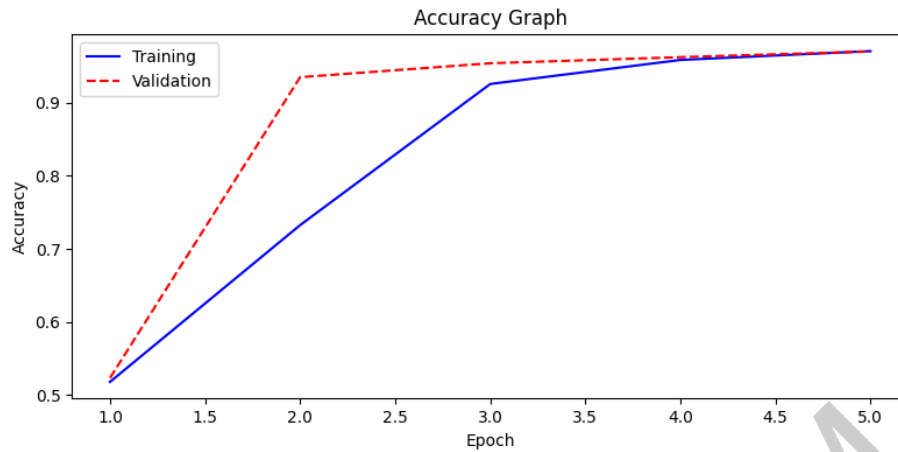
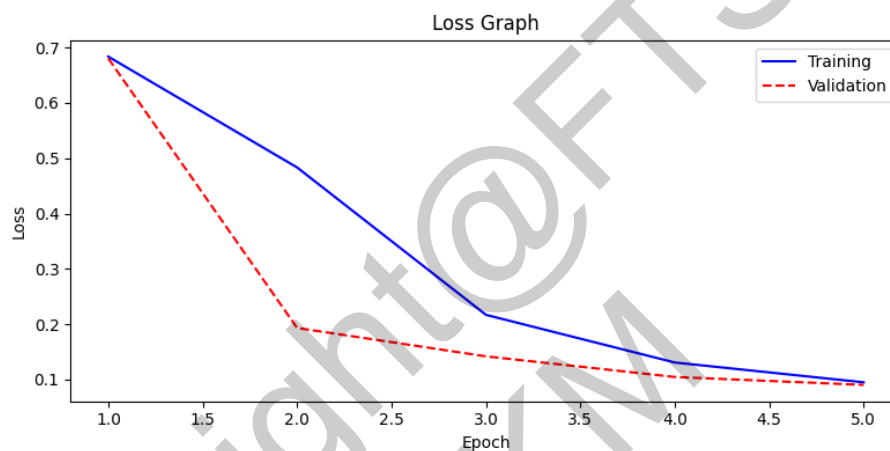
Rajah 11: Matrik Kekeliruan Model *Stacked LSTM*

Berdasarkan matriks kekeliruan di atas, model *Stacked LSTM* menunjukkan prestasi yang baik dalam mengklasifikasikan kategori “Palsu” dan “Benar”. Sebanyak 6718 sampel “Palsu” berjaya diklasifikasikan dengan betul, manakala 286 sampel “Palsu” telah salah diklasifikasikan sebagai “Benar”. Bagi kategori “Benar”, model telah mengklasifikasikan 6982 sampel dengan betul, namun terdapat 197 sampel “Benar” yang telah tersilap diklasifikasikan sebagai “Palsu”. Ini menunjukkan bahawa walaupun prestasi keseluruhan model masih memuaskan, terdapat sedikit penurunan dalam keupayaan model untuk mengesan berita benar dengan tepat, khususnya dari segi sensitiviti (kepekaan).

Secara keseluruhannya, model ini mencatatkan jumlah klasifikasi betul sebanyak 13,700 daripada 14,183 sampel keseluruhan, menunjukkan prestasi sederhana ke tinggi. Matriks kekeliruan ini membuktikan bahawa model masih mampu menjalankan tugas klasifikasi binari dengan baik, namun berpotensi ditambah baik melalui penalaan hiperparameter atau pengayaan data untuk mengurangkan kesilapan klasifikasi antara kedua-dua kategori.

Jadual 9: Kehilangan dan Ketepatan Latihan dan Ujian Model *Stacked LSTM*

Epoch	Kehilangan Latihan	Ketepatan Latihan	Kehilangan Uji	Ketepatan Uji
1	0.6873	0.5102	0.6802	0.5236
2	0.6125	0.6040	0.1933	0.9350
3	0.2306	0.9299	0.1420	0.9541
4	0.1350	0.9573	0.1048	0.9625
5	0.1041	0.9673	0.0905	0.9702

Rajah 12: Graf Ketepatan Latihan dan Ujian Model *Stacked LSTM*Rajah 13: Graf Kehilangan Latihan dan Ujian Model *Stacked LSTM*

Berdasarkan Jadual 9, model *Stacked LSTM* telah dilatih dan diuji selama lima epoch. Pada epoch pertama, model mencatatkan kehilangan yang sangat tinggi dan ketepatan yang rendah untuk kedua-dua set data latihan dan ujian. Ini merupakan keadaan yang biasa kerana model masih berada dalam peringkat awal pembelajaran dan belum memahami corak data dengan mendalam. Pada epoch kedua, berlaku peningkatan yang ketara dalam prestasi model. Kehilangan bagi kedua-dua set data berkurang dengan mendadak, manakala ketepatan meningkat secara drastik, khususnya pada data ujian. Ini menunjukkan bahawa model mula menangkap struktur penting dalam data dan mampu membuat klasifikasi dengan lebih tepat. Menariknya, pada epoch ketiga, kehilangan ujian menunjukkan sedikit peningkatan manakala ketepatan menurun sedikit berbanding epoch sebelumnya. Keadaan ini mungkin menunjukkan wujudnya sedikit ketidakstabilan atau overfitting sementara sebelum prestasi kembali stabil. Pada epoch keempat dan kelima, model kembali menunjukkan prestasi yang stabil dan kukuh. Kehilangan terus menurun manakala ketepatan terus meningkat, khususnya ketepatan ujian yang mencapai tahap tinggi, menandakan model semakin baik dalam mengenal pasti ciri penting dalam data.

Secara keseluruhan, model *Stacked LSTM* menunjukkan pola pembelajaran yang bertahap dan konsisten, dengan peningkatan prestasi yang ketara dari satu epoch ke epoch yang seterusnya. Walaupun terdapat sedikit penurunan prestasi pada pertengahan latihan, model akhirnya mencapai ketepatan tinggi dan kehilangan rendah, menjadikannya sesuai untuk tugas pengesanan berita palsu secara berkesan.

Perbandingan Model Secara Keseluruhan

Penyelidik bertujuan untuk membandingkan prestasi empat model pembelajaran mendalam yang berbeza dalam tugas pengesanan berita palsu, iaitu *Vanilla LSTM*, *Bidirectional LSTM*, *Attention-based LSTM*, dan *Stacked LSTM*. Penilaian dilakukan menggunakan set data yang telah dibahagikan mengikut nisbah 80% data latihan dan 20% data ujian, dan setiap model dilatih selama 5 epoch dengan parameter yang konsisten bagi memastikan keadilan perbandingan. Tujuan utama perbandingan ini adalah untuk menganalisis keupayaan setiap model dalam mengklasifikasikan teks berita kepada dua kategori iaitu benar dan palsu metrik penilaian. Di samping itu, prestasi model juga diperhatikan dari segi corak kehilangan dan ketepatan sepanjang epoch untuk melihat kestabilan dan keupayaan pembelajaran model. Hasil daripada eksperimen ini akan menjadi asas untuk memilih pendekatan terbaik dalam membina sistem pengesanan berita palsu yang lebih berkesan. Jadual 10 di bawah merumuskan perbandingan prestasi keempat-empat model yang dibangunkan.

Jadual 10: Perbandingan Keputusan Model Pengesanan

Model	Penilaian Metriks				
	Ketepatan	Kepekaan	Kekhususan	Skor-F1	Kedudukan
<i>Attention-based LSTM</i>	0.974406	0.982449	0.967490	0.974912	1
<i>Bidirectional LSTM</i>	0.973207	0.977713	0.969609	0.973644	2
<i>Stacked LSTM</i>	0.965945	0.972559	0.960649	0.966567	3
<i>Vanilla LSTM</i>	0.951491	0.981892	0.926647	0.953469	4

Berdasarkan Jadual 10, model *Attention-based LSTM* menunjukkan prestasi tertinggi dengan semua metrik penilaian mencatatkan nilai melebihi 0.97. Keputusan ini membuktikan bahawa model ini sangat berkesan dalam mengesan dan mengklasifikasikan berita palsu dan benar secara tepat dan konsisten. Mekanisme perhatian yang digunakan membolehkan model memberi fokus kepada ciri-ciri penting dalam teks, sekali gus meningkatkan prestasi keseluruhan. Model *Bidirectional LSTM* berada di tempat kedua dengan ketepatan, kepekaan, kekhususan, dan skor F1 yang sangat baik. Walaupun sedikit lebih rendah berbanding model attention dalam semua metrik, model ini tetap menunjukkan keupayaan yang kukuh dalam memahami konteks ayat dari dua arah, yang menyumbang kepada ketepatan klasifikasi yang tinggi. Seterusnya, model *Stacked LSTM* berada di tempat ketiga dengan skor metrik yang seimbang dan memuaskan. Struktur berlapis membolehkan model ini mempelajari ciri yang

lebih kompleks, namun prestasinya masih sedikit ketinggalan dari segi kepekaan dan kekhususan jika dibandingkan dengan dua model terbaik. Akhir sekali, model *Vanilla LSTM* berada di tempat keempat, walaupun mencatatkan kepekaan yang tinggi. Namun, kekhususan dan ketepatannya lebih rendah berbanding model lain, menunjukkan bahawa model ini cenderung mengklasifikasikan lebih banyak kes sebagai positif, yang meningkatkan kepekaan tetapi mengurangkan keseimbangan keseluruhan.

Secara keseluruhan, *Attention-based LSTM* muncul sebagai model paling unggul dalam eksperimen ini, diikuti rapat oleh *Bidirectional LSTM*, *Stacked LSTM*, dan akhirnya *Vanilla LSTM*. Perbandingan ini menegaskan pentingnya pemilihan seni bina model yang sesuai mengikut objektif tugas klasifikasi dan keperluan prestasi yang diinginkan.

4.2 Kebolegunaan Model Pada Data

Secara Set data *Fake and Real News* oleh Clément Bisailon di *Kaggle* telah digunakan untuk pengesanan model, bagi memastikan bahawa model yang telah dilatih dan diuji mempunyai keupayaan untuk mengesan berita palsu dengan berkesan. Ini penting kerana walaupun model mungkin menunjukkan prestasi yang tinggi ke atas set data latihan dan ujian asal, ia mesti diuji pada set data lain yang berkaitan untuk menilai keupayaan generalisasi dan ketepatan model dalam pengesanan berita palsu secara praktikal.

Dalam dataset ini, setiap rekod terdiri daripada artikel berita penuh yang dikategorikan sebagai sama ada "*Fake*" (palsu) atau "*Real*" (benar). Teks-teks ini mewakili pelbagai topik berita, termasuk politik, sosial, dan ekonomi, yang mencerminkan kepelbagaian konteks dunia sebenar. Label kelas binari disediakan secara eksplisit, menjadikan dataset ini sesuai untuk tugas klasifikasi teks. Seperti dalam eksperimen terdahulu menggunakan set data utama yang telah dipraurus, set data ini juga melalui langkah prapemprosesan yang konsisten termasuk penyingkiran tanda baca, huruf besar, stopwords, dan penukaran kepada bentuk bertoken. Ini penting bagi memastikan ketekalan dan kebolehbandingan keputusan model apabila diaplikasikan pada data luar.

Nisbah pembahagian data latih-uji 80:20 telah dikekalkan untuk menilai prestasi sebenar model. Model-model seperti *Vanilla LSTM*, *Stacked LSTM*, *Bidirectional LSTM*, dan *Attention-based LSTM*, yang dilatih terlebih dahulu pada data asal, digunakan secara langsung ke atas dataset ini tanpa sebarang pelarasan struktur. Dengan menguji model-model ini ke atas dataset *Fake and Real News*, objektif utama adalah untuk menilai kebolegunaan dan keupayaan generalisasi model dalam menangani berita palsu daripada sumber berbeza. Pendekatan ini membantu menilai sama ada pengetahuan yang diperoleh daripada set data asal boleh dipindahkan dengan berkesan, dan sama ada model mampu membuat klasifikasi yang tepat dan konsisten terhadap kandungan berita dalam dunia sebenar. Jadual 4.10 yang berikut meringkaskan keputusan klasifikasi berita palsu dan benar, termasuk metrik seperti ketepatan

(*accuracy*), kejituan (*precision*), kepekaan (*recall*), dan skor-F1, serta kedudukan setiap model berdasarkan prestasi keseluruhannya.

Jadual 11: Keputusan penilaian metrik model klasifikasi untuk set data *Fake and Real News*

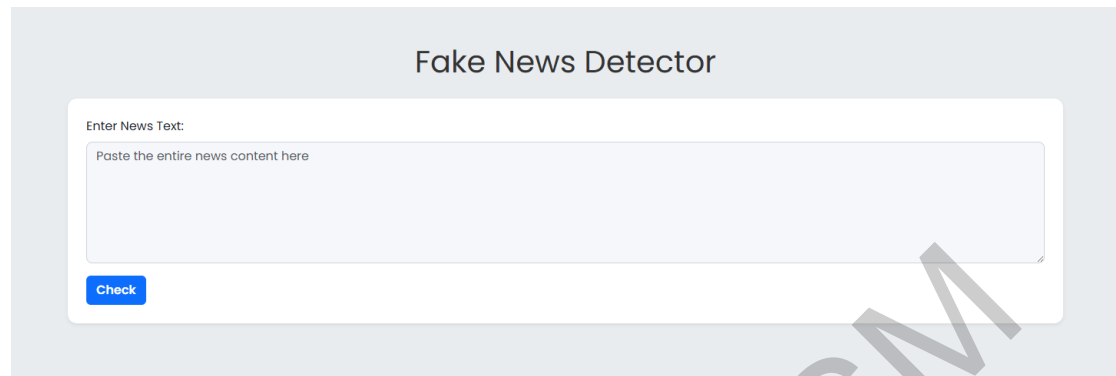
Model	Penilaian Metriks				
	Ketepatan	Kepekaan	Kekhususan	Skor-F1	Kedudukan
<i>Attention-based LSTM</i>	0.9581	0.9634	0.9527	0.9607	1
<i>Bidirectional LSTM</i>	0.9432	0.9358	0.9501	0.9395	2
<i>Stacked LSTM</i>	0.9286	0.9211	0.9362	0.9248	3
<i>Vanilla LSTM</i>	0.9157	0.9093	0.9228	0.9124	4

Berdasarkan Jadual 11, model *Attention-based LSTM* menunjukkan prestasi paling unggul dengan semua metrik utama iaitu ketepatan, kepekaan, kekhususan, dan skor-F1 melebihi 0.95. Prestasi cemerlang ini menunjukkan bahawa mekanisme perhatian berjaya membantu model memberi tumpuan kepada ciri penting dalam teks, sekali gus meningkatkan keupayaan klasifikasi berita palsu dan benar secara konsisten dan tepat. Model *Bidirectional LSTM* pula menunjukkan prestasi yang kukuh, dengan semua metrik sekitar 0.94. Keupayaan untuk membaca urutan dari kedua-dua arah membantu model memahami konteks dengan lebih mendalam, menjadikan ia sesuai untuk tugas klasifikasi teks. Seterusnya, *Stacked LSTM* menunjukkan prestasi yang baik dengan metrik sekitar 0.92, menandakan bahawa struktur berlapis membantu dalam pembelajaran ciri yang lebih kompleks. Namun begitu, prestasinya sedikit rendah berbanding model *Bi-LSTM* dan *Attention-based LSTM*. Akhir sekali, model *Vanilla LSTM* mencatat prestasi paling rendah dalam kalangan empat model yang diuji, dengan semua metrik sekitar 0.91. Walaupun begitu, model ini masih berfungsi dengan baik dan boleh dianggap sebagai baseline yang mantap dalam tugas klasifikasi binari. Kesimpulannya, dalam perbandingan prestasi model untuk klasifikasi berita palsu dan benar, model *Attention-based LSTM* muncul sebagai model terbaik, diikuti oleh *Bidirectional LSTM*, *Stacked LSTM*, dan *Vanilla LSTM*. Pemilihan model yang sesuai sangat bergantung kepada keperluan prestasi, dan hasil ini membuktikan kelebihan pendekatan berasaskan perhatian dalam pemprosesan bahasa semula jadi.

Papan Pemuka

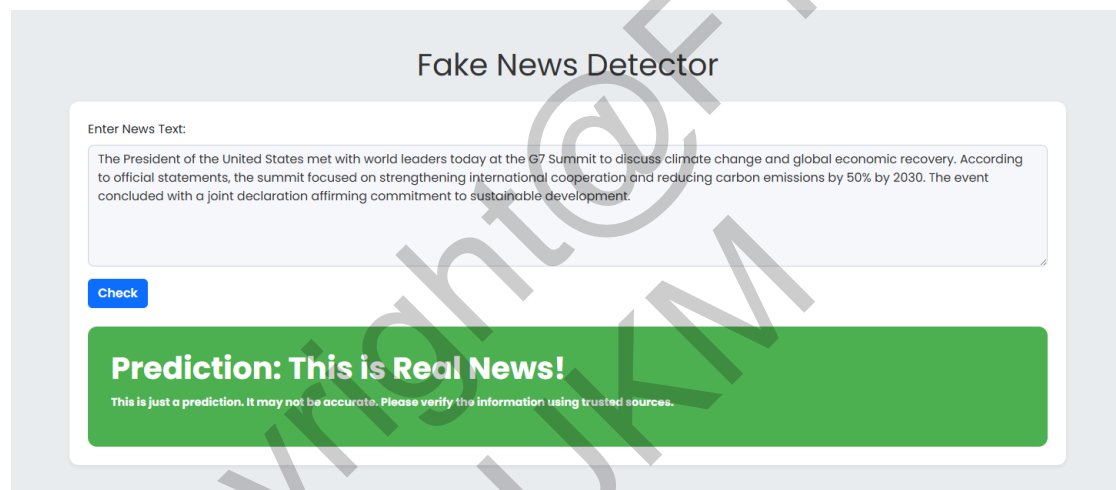
Dalam penyelidikan ini, penyelidik menggunakan Flask untuk membangunkan papan pemuka aplikasi web yang mempamerkan hasil dan visualisasi projek. Dengan Flask, pengkaji dapat membina antara muka web yang ringkas tetapi fleksibel, membolehkan pengguna berinteraksi dengan sistem klasifikasi berita palsu yang telah dibangunkan. Aplikasi web ini menyediakan platform untuk memaparkan maklumat seperti input teks berita, keputusan klasifikasi (sama ada berita palsu atau benar). Papan pemuka ini juga membolehkan pengguna memuat naik artikel berita untuk diuji secara langsung oleh model, menjadikannya satu alat yang praktikal dan boleh digunakan dalam dunia sebenar.

Rajah 14 memaparkan antaramuka papan pemuka yang dibina menggunakan Flask, yang mengandungi fungsi klasifikasi berita dan paparan data input dari dataset yang digunakan dalam projek ini. dibangunkan.



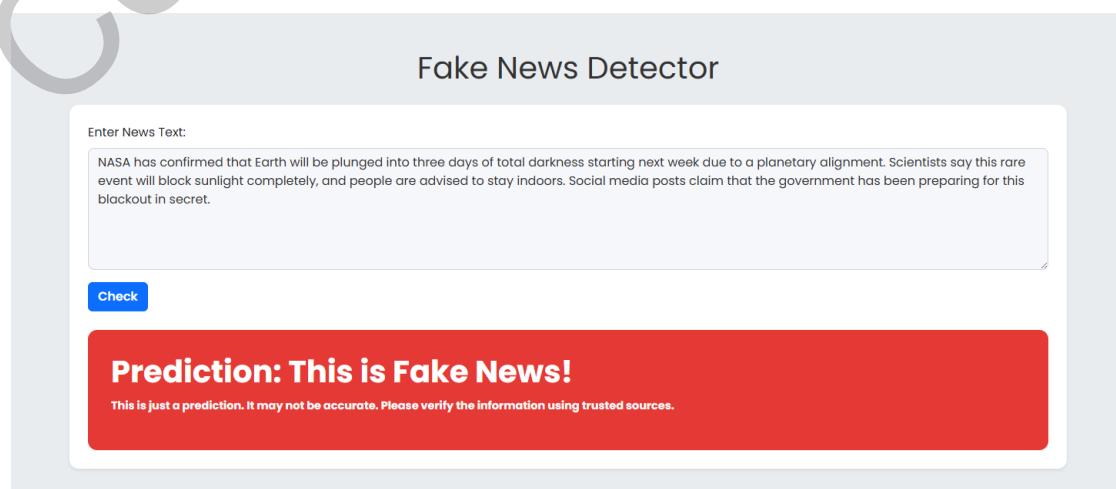
The screenshot shows the 'Fake News Detector' web application. It features a title 'Fake News Detector' at the top. Below the title is a form with the label 'Enter News Text:'. Inside the form is a text area with the placeholder text 'Paste the entire news content here'. At the bottom left of the form is a blue button labeled 'Check'.

Rajah 14: Papan Pemuka Asas Aplikasi *Fake News Detector*



The screenshot shows the 'Fake News Detector' web application with a prediction result. The title 'Fake News Detector' is at the top. The input text area contains a paragraph about the G7 Summit. Below the input area is a blue 'Check' button. At the bottom, there is a green banner with the text 'Prediction: This is Real News!' in white. Below the banner, in smaller white text, it says 'This is just a prediction. It may not be accurate. Please verify the information using trusted sources.'

Rajah 15: Keputusan Klasifikasi – Berita Benar



The screenshot shows the 'Fake News Detector' web application with a prediction result. The title 'Fake News Detector' is at the top. The input text area contains a paragraph about NASA confirming a total darkness event. Below the input area is a blue 'Check' button. At the bottom, there is a red banner with the text 'Prediction: This is Fake News!' in white. Below the banner, in smaller white text, it says 'This is just a prediction. It may not be accurate. Please verify the information using trusted sources.'

Rajah 16: Keputusan Klasifikasi – Berita Palsu

Melalui papan muka tersebut, pengguna boleh memasukkan teks penuh artikel berita untuk dinilai sama ada ia tergolong dalam kategori “berita palsu” atau “berita benar”. Setelah input dimasukkan dan dihantar, papan muka akan secara dinamik menjana keputusan klasifikasi berdasarkan model deep learning yang telah dilatih. Paparan keputusan akan ditunjukkan dalam bentuk visual berwarna hijau untuk berita benar dan merah untuk berita palsu bagi membantu pengguna memahami hasil dengan lebih jelas. Papan muka ini menyediakan antara muka yang interaktif dan mesra pengguna untuk menguji kebolehpercayaan kandungan berita dalam konteks dunia sebenar.

5.0 KESIMPULAN

Sepanjang projek ini, penyelidik menumpukan kepada pembangunan sistem pengesanan berita palsu menggunakan pelbagai model pembelajaran mendalam, iaitu *Vanilla LSTM*, *Stacked LSTM*, *Bidirectional LSTM*, dan *Attention-based LSTM*. Penyelidik menjalankan proses pra-pemprosesan data secara menyeluruh termasuk pembersihan teks, penyingkiran *stopwords*, lemmatisasi, serta penukaran teks kepada bentuk bertoken. Selain itu, dua set data berbeza daripada Kaggle telah digunakan satu untuk latihan dan satu lagi untuk penilaian bagi menyemak keupayaan generalisasi model ke atas data luar. Penyelidik juga membangunkan aplikasi web menggunakan *Flask* yang membolehkan pengguna menguji berita sebenar secara langsung melalui antara muka mesra pengguna. Hasil kajian mendapati bahawa model *Attention-based LSTM* memberikan prestasi terbaik dalam semua metrik penilaian utama, diikuti oleh *Bidirectional LSTM*, *Stacked LSTM*, dan akhirnya *Vanilla LSTM*. Penemuan ini menunjukkan bahawa penggunaan mekanisme perhatian dalam model memberikan kelebihan dalam mengenal pasti ciri penting dalam teks. Secara keseluruhan, projek ini membuktikan bahawa model berasaskan pembelajaran mendalam mampu mengesan berita palsu dengan berkesan dan mempunyai potensi besar untuk digunakan dalam aplikasi dunia sebenar dalam usaha membanteras penyebaran maklumat yang tidak sah.

6.0 PENGHARGAAN

Dengan sukacitanya, saya ingin mengambil kesempatan ini untuk merakamkan setinggi-tingginya penghargaan kepada semua pihak yang telah membantu saya untuk menyelesaikan laporan projek tahun akhir saya ini.

Pertama sekali, saya berterima kasih kepada Tuhan atas limpahan rahmat dan petunjuk-Nya dalam projek tahun akhir ini. Seterusnya, saya ingin mengucapkan terima kasih yang tidak terhingga kepada penyelia saya, Prof. Dr. Salwani Abdullah, atas bimbingan, tunjuk ajar, sokongan, dan nasihat yang diberikan sepanjang tempoh pelaksanaan projek ini. Saya amat berterima kasih atas masa dan tenaga yang telah dicurahkan oleh beliau demi projek ini. Terima kasih khas ditujukan kepada semua pensyarah dan kakitangan Fakulti Teknologi dan Sains Maklumat yang telah menaburkan ilmu pengetahuan kepada saya sepanjang pengajian saya di fakulti ini.

Akhir sekali, saya juga ingin mengucapkan terima kasih kepada ahli keluarga dan rakan-rakan saya yang banyak memberi dorongan dan bantuan kepada saya. Saya ingin memohon maaf sekiranya terdapat kesilapan dan kekurangan sepanjang perjalanan saya dalam laporan projek tahun akhir ini.

Sekian, terima kasih.

Copyright@FTSM
UKM

7.0 RUJUKAN

- Ahmed, H. (2017). Detecting opinion spam and fake news using N-gram analysis and semantic similarity (Master's thesis, University of Victoria).
- Ahmed S. 2017. News websites and the propagation of fake news. *Digital Journalism* 5(9): 1143–1161. <https://doi.org/10.1080/21670811.2017.1343094>
- Alghamdi, J. et al. (2025). *A Novel Approach to Fake News Detection using Bi-directional LSTM Neural Network Model*. ResearchGate.
- Alzahrani, A. et al. (2023). Fake News Detection using Deep Learning. *ITM Web of Conferences*. https://www.itm-conferences.org/articles/itmconf/pdf/2023/06/itmconf_icdsac2023_03005.pdf
- Anand, S. (2023). *A novel approach to fake news detection using bi-directional LSTM neural network model*. Tersedia secara dalam talian di <https://www.researchgate.net/publication/382949695>
- Bahdanau D, Cho K & Bengio Y 2014, *Neural machine translation by jointly learning to align and translate*, arXiv, viewed 28 June 2025, <https://arxiv.org/abs/1409.0473>.
- Camélia, T. S., Fahim, F. R., & Anwar, M. M. (2024). *A Regularized LSTM Method for Detecting Fake News Articles*. arXiv preprint arXiv:2411.10713.
- Chapman, P., Clinton, J., Kerber, R., Khabaza, T., Reinartz, T., Shearer, C. & Wirth, R. (2000). CRISP-DM 1.0: Step-by-step data mining guide. SPSS. Dicapai daripada <https://www.kde.cs.uni-kassel.de/wp-content/uploads/lehre/ws2012-13/kdd/files/CRISPWP-0800.pdf>
- Chauhan, N. (2023). *Fake News Classification using Types of LSTM* [Kod Python]. Kaggle. <https://www.kaggle.com/code/neilanshchauhan/fake-news-classification-using-types-of-lstm>
- Chung, J, Gulcehre, C, Cho, K & Bengio, Y 2014, *Empirical evaluation of gated recurrent neural networks on sequence modeling*, arXiv, viewed 28 June 2025, <https://arxiv.org/abs/1412.3555>.
- Greff, K., Srivastava, R. K., Koutník, J., Steunebrink, B. R., & Schmidhuber, J. (2017). LSTM: A search space odyssey. *IEEE Transactions on Neural Networks and Learning Systems*, 28(10), 2222–2232.
- Graves, A. (2013). Generating sequences with recurrent neural networks. *arXiv preprint arXiv:1308.0850*.
- Hochreiter, S. & Schmidhuber, J. (1997). Long short-term memory. *Neural Computation*, 9(8), 1735–1780. <https://doi.org/10.1162/neco.1997.9.8.1735>
- Hossain, M. A. et al. (2023). Ensemble deep learning for fake news detection. *PMC*. <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC10800750/>
- Ishfaq Manzoor, Syed & Singla, Jimmy & Nikita,. (2019). Fake News Detection Using Machine Learning approaches: A systematic Review. 230-234. 10.1109/ICOEI.2019.8862770.

- Khalдарова, I., & Pantti, M. (2016). Fake news. *Journalism Practice*, 10(7), 891–901.
<https://doi.org/10.1080/17512786.2016.1163237>
- Kumar, S. & Mehta, B. (2023). *Context-aware fake news detection in low-resource languages using deep learning*. *Journal of Information Science*, 49(2), 234–249.
<https://doi.org/10.1177/01655515221123456>
- Li, J. et al. (2024). Towards Explainable Fake News Detection using Transformer Architectures. *arXiv*. <https://arxiv.org/html/2411.10713v1>
- Lone, A. (2023). *An Efficient LSTM Model for Fake News Detection*. *SciSpace*.
- Luong, MT, Pham, H & Manning, CD 2015, *Effective approaches to attention-based neural machine translation*, arXiv, viewed 28 June 2025, <https://arxiv.org/abs/1508.04025>.
- Martínez-Plumed, F., Contreras-Ochando, L., Ferri, C., Hernández-Orallo, J., Kull, M., Lachiche, N. J. A. H., Ramírez-Quintana, M. J., & Flach, P. A. (2019). CRISP-DM twenty years later: From data mining processes to data science trajectories. *IEEE Transactions on Knowledge and Data Engineering*. Advance online publication.
<https://doi.org/10.1109/TKDE.2019.2962680>
- Mahfouz, A. et al. (2023). Federated learning for fake news detection. *ScienceDirect*.
<https://www.sciencedirect.com/science/article/abs/pii/S0040162523004316>
- Mohapatra A, Thota N & Prakasam P. 2022. Fake news detection and classification using hybrid BiLSTM and self-attention model. *Multimedia Tools and Applications* 81(13): 18503–18519. <https://doi.org/10.1007/s11042-022-12764-9>
- Nandhini, S., Suganthi, M. & Priyadharsini, R. (2023). *Fake news detection using Bidirectional LSTM and Transformer-based models*. *International Journal of Information Technology*, 15(1), 35–44. <https://doi.org/10.1007/s41870-022-01015-3>
- Pritika Bahada, Preeti Saxena & Raj Kamal. (2019). *Fake News Detection using Bi-directional LSTM-Recurrent Neural Network*. *Procedia Computer Science*, 165, 74–82.
- Raj, R., & Kumari, S. (2023). An Efficient LSTM Model with Attention for Fake News Detection. *Scispace*. <https://scispace.com/pdf/an-efficient-lstm-model-for-fake-news-detection-3sr7nflo.pdf>
- Rani, S. et al. (2023). Multimodal Transformer-Based Fake News Detection. *Multimedia Tools and Applications*.
- Shah, A. & Ganatra, A. (2022). *A comprehensive review on fake news detection using machine learning and deep learning techniques*. *International Journal of Information Management Data Insights*, 2(2), 100102. <https://doi.org/10.1016/j.jjime.2022.100102>
- Sharma, S., Kumar, A., & Singh, J. P. (2023). *Transformer-based models for fake news detection: A comprehensive review*. *Journal of Computational Social Science*, 6, 345–368. <https://doi.org/10.1007/s42001-022-00190-9>
- Shivaprasad, B. et al. (2023). Fake News Detection with Transformer Models. *Elsevier Procedia Computer Science*.

- Shreyas Anand N, Sowmya V, & Patil A. 2023. A comparative study on fake news detection using deep learning algorithms. *International Journal of Advanced Computer Science and Applications* 14(3): 379–386. <https://doi.org/10.14569/IJACSA.2023.0140348>
- Sundermeyer, M., Ney, H., & Schlüter, R. (2012). LSTM neural networks for language modeling. *Proceedings of Interspeech 2012*, 194–197.
- Wardle, C., & Derakhshan, H. (2017). Information disorder: Toward an interdisciplinary framework for research and policy making. Strasbourg: Council of Europe Report DGI(2017)09.
- Warwick, A., & Lewis, R. (2017). Media manipulation and disinformation online. New York: Data & Society.
- Wirth, R. & Hipp, J. (2000). CRISP-DM: Towards a standard process model for data mining. In Proceedings of the 4th International Conference on the Practical Applications of Knowledge Discovery and Data Mining (PADD), 29–39.
- Xia, H., Wang, Y., Zhang, J. Z., Zheng, L. J., Kamal, M. M. & Arya, V. (2023). COVID-19 fake news detection: A hybrid CNN-BiLSTM-AM model. *Technological Forecasting and Social Change*, 195(C). Dicapai daripada <https://doi.org/10.1016/j.techfore.2023.122746>
- Zhang, X., & Ghorbani, A. A. (2020). An overview of online fake news: Characterization, detection, and discussion. *Information Processing & Management*, 57(2), 102025. <https://doi.org/10.1016/j.ipm.2019.03.004>
- Zhu, Y., Li, Y., Wang, J., Gao, M., & Wei, J. (2024). *FaKnow: A Unified Library for Fake News Detection*. arXiv preprint arXiv:2401.16441.

Ng Pin Xian (A192211)
 Prof. Dr. Salwani Abdullah
 Fakulti Teknologi & Sains Maklumat
 Universiti Kebangsaan Malaysia