

AUTOMASI PENGESANAN AKTIVITI PENCEROBOHAN SIBER MENGUNAKAN PEMBELAJARAN MESIN DAN PEMBELAJARAN MENDALAM

¹Nur 'Eiza Athira Eikhmerizal Bazura, ²Rohizah Abd Rahman

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

Abstrak

Pencerobohan siber merujuk kepada sebarang usaha yang berniat jahat untuk menceroboh, mencuri, mengganggu atau merosakkan sistem rangkaian, data dan infrastruktur digital. Dalam era digital ini, peningkatan aktiviti siber yang bersifat merosakkan telah mewujudkan keperluan yang mendesak terhadap sistem pengesanan serangan yang lebih cekap dan pantas. Objektif kajian bertujuan membangunkan sebuah model pengesanan pencerobohan siber dengan menggunakan pendekatan pembelajaran mesin dan pembelajaran mendalam. Beberapa algoritma telah diuji, termasuk Logistic Regression, Random Forest, Support Vector Machine (SVM), Multilayer Perceptron (MLP) dan Autoencoder bagi menilai keberkesanan dalam mengesan corak serangan berdasarkan ciri tingkah laku log masuk seperti bilangan cubaan log masuk, log masuk yang gagal skor reputasi IP. Metodologi yang akan digunakan adalah berasaskan model Cross Industry Standard Process for Data Mining (CRISP-DM). Model ini merangkumi enam fasa utama iaitu pemahaman bisnes, pemahaman data, penyediaan data, pembangunan model, penilaian dan penyebaran. Kajian ini akan dilarikan menggunakan bahasa pengaturcaraan Python. Terdapat 4 langkah utama dalam projek ini (1) melakukan pembersihan dan eksplorasi data analisis (2) pemilihan model pembelajaran mesin dan pembelajaran mendalam (3) membangunkan dan latihan model pengesanan dan (4) membangunkan sebuah antara muka menggunakan StreamLit bagi analisis ringkasan projek. Hasil daripada kajian ini menunjukkan bahawa model SVM mencatatkan purata metrik penilaian tertinggi berbanding model-model lain. Oleh itu, model SVM telah dipilih untuk digunakan dalam proses pengesanan masa nyata melalui antara muka pengguna Streamlit bagi tujuan pengujian set data. Melalui integrasi model ini, sistem berjaya menjalankan proses pengesanan aktiviti pencerobohan siber ini. Diharap hasil kajian ini mampu menyokong pembangunan sistem pengesanan pencerobohan siber yang lebih responsif, tepat, dan praktikal. Juga model yang dibangunkan diharap boleh digunapakai untuk algoritma pembelajaran mendalam terkini yang dipertingkatkan pada masa akan datang.

Kata kunci: Pencerobohan Siber, Pembelajaran Mesin, Pembelajaran Mendalam

Abstract

Cyber intrusion refer to any malicious attempts to intrude, steal, disrupt, or damage network systems, data, and digital infrastructure. In today's digital era, the rise of destructive cyber activities has created an urgent need for more efficient and rapid detection systems for attacks. This study aims to develop a cyber attack detection model using machine learning and deep learning approaches. Several algorithms were tested, including Logistic Regression, Random Forest, Support Vector Machine (SVM), Multilayer Perceptron (MLP), and Autoencoder to evaluate their effectiveness in detecting attack patterns based on login behavior features such as the number of login attempts, failed logins, and IP reputation score. The methodology adopted in this project follows the Cross Industry Standard Process for Data Mining (CRISP-DM) model, which consists of six main phases: business understanding, data understanding, data preparation, model development, evaluation, and deployment. This project will be implemented using the Python programming language. There are four main steps involved in the project: (1) performing data cleaning and exploratory data analysis, (2) selecting appropriate machine learning and deep learning models, (3) developing and training the detection model, and (4) building a user interface using Streamlit to present a summary analysis of the project. The findings of this study show that the SVM model recorded the highest average evaluation metric compared to other models. Therefore, the SVM model was selected to be used in the real-time detection process through the Streamlit user interface for dataset testing purposes. Through the integration of this model, the system successfully performed the detection of cyber intrusion activities. In conclusion, the outcomes of this project are expected to support the development of a cyber attack detection system that is more responsive, accurate, and practical for real-world applications.

Keywords: Cyber Intrusion, Machine Learning , Deep Learning

1.0 PENGENALAN

Penggunaan internet telah meningkat dengan cepat, terutamanya dalam dekad yang lepas (Ghanem et al. 2022). Dalam bidang keselamatan komputer, pengkomputeran selari telah terbukti sebagai kaedah yang berkesan untuk mempercepatkan proses memecahkan kata laluan menggunakan teknik *brute force* serta serangan kamus (*dictionary attack*) (Alkhwaja et al. 2023). Serangan *brute force* ialah kaedah di mana penyerang mencuba semua kemungkinan dengan mengulangi pelbagai kombinasi kata laluan atau kunci penyulitan sehingga menemui yang betul. Ia merupakan kaedah yang biasa digunakan untuk memecahkan kata laluan. Disebabkan oleh peningkatan kebergantungan terhadap pendigitalan, pelbagai insiden

keselamatan seperti akses tanpa kebenaran sering berlaku. Sistem pengesanan pencerobohan (*Intrusion detection system*, IDS), tembok api (*firewall*) dan perisian antivirus hanyalah sebahagian daripada langkah keselamatan yang tersedia (Hamza et al., 2024). Penggunaan teknologi maklumat yang meluas serta kemunculan dan perkembangan ruang siber telah menyumbang kepada kemajuan dan kemakmuran ekonomi dan masyarakat. Namun begitu, ia juga telah membawa risiko keselamatan dan cabaran baru (Zhang et al. 2022).

Seterusnya, tempoh pengiraan yang diperlukan untuk menemui kata laluan melalui teknik *brute force* bergantung kepada pelbagai faktor, termasuk panjang dan kerumitan set aksara kata laluan tersebut, serta kerumitan pengiraan algoritma penyulitan yang digunakan. Selain itu, sekiranya proses pengiraan dijalankan pada satu pemproses sahaja, ia mungkin memerlukan lebih banyak masa untuk memecahkan kata laluan. Namun, dengan mengagihkan beban kerja ke beberapa pemproses atau peranti, proses tersebut dapat dipercepatkan. Ini bukan sahaja menjadikan pemecahan kata laluan lebih efisien, malah turut meningkatkan peluang untuk berjaya. Apabila sesuatu kata laluan dianggap sangat selamat, kaedah *brute force* menjadi sangat diperlukan. Walau bagaimanapun, terdapat juga kaedah alternatif untuk memecahkan kata laluan yang dianggap lebih berkesan, seperti serangan kamus (*dictionary attack*) (Alkhwaja et al. 2023). Dalam serangan kamus (*dictionary attack*), penggadam akan menggunakan fail kamus yang mengandungi ratusan atau mungkin jutaan kata laluan popular dan mencubanya satu per satu sehingga menemui yang betul. Sekiranya serangan ini berjaya, penggadam berkemungkinan besar dapat mengakses sistem serta semua data yang terkandung di dalamnya (Hamza & surayh Al-Janabi 2024).

Untuk menangani isu ini, penggunaan pembelajaran mesin (*machine learning*) dan pembelajaran mendalam (*deep learning*) digunakan untuk menganalisis corak tingkah laku pengguna secara menyeluruh bagi mengesan potensi serangan. Sebagai contoh, apabila data pengguna dikumpulkan termasuk jumlah percubaan log masuk serta bilangan kegagalan log masuk. Pembelajaran mesin dan pembelajaran mendalam boleh dilatih untuk mengenal pasti pola-pola mencurigakan. Corak seperti jumlah percubaan log masuk yang tinggi dalam masa singkat atau kadar kegagalan yang luar biasa boleh menjadi petunjuk kepada cubaan serangan seperti *brute force*. Melalui proses pembelajaran berasaskan data ini, sistem dapat membezakan tingkah laku pengguna yang biasa dan yang berpotensi berbahaya, lalu memberikan amaran awal

atau mengambil tindakan bagi mengelakkan pencerobohan selanjutnya. Pendekatan ini bukan sahaja meningkatkan kecekapan dalam pengesanan serangan, malah juga membolehkan sistem keselamatan bertindak secara proaktif dan pintar.

Laporan teknikal ini dibahagikan kepada kajian literatur, metodologi kajian, hasil kajian, kesimpulan, penghargaan serta rujukan. Kajian literatur merupakan koleksi analisis kajian terdahulu yang berkaitan dengan topik pengesanan aktiviti pencerobohan siber. Metodologi kajian menjelaskan tentang kaedah dan pendekatan yang digunakan dalam menjalankan kajian. Hasil kajian memaparkan keputusan yang diperoleh. Seterusnya, kesimpulan menyimpulkan keseluruhan kajian dan memberi ringkasan tentang hasil kajian serta implikasinya. Penghargaan adalah pernyataan terima kasih kepada individu, kumpulan atau pihak yang telah memberikan sokongan, bantuan atau sumbangan dalam menjalankan kajian. Rujukan merupakan senarai sumber atau bahan yang digunakan dalam projek ini.

2.0 KAJIAN LITERATUR

Perbandingan jenis algoritma oleh kajian-kajian terdahulu membantu kita mengenalpasti batasan dan ketepatan bagi setiap kajian. Menurut kepada kajian Hamza pada tahun 2024, mereka memfokuskan kepada pengesanan serangan brute force terhadap protokol SSH dan FTP dengan menggunakan teknik pembelajaran mesin. Hasil kajiannya mendapati Random Forest (RF) menunjukkan ketepatan tertinggi iaitu 99.905 berbanding Logistic Regression (LR) dan Naive Bayes (NB) yang masing-masing mencapai 89.17% dan 48.61%. Seterusnya, kajian oleh A.Sharma dan rakan-rakan pada tahun 2024, menggabungkan model ramalan dan pengesanan anomali. Decision Tree (DT) mencapai ketepatan tertinggi iaitu 99.96% diikuti K-Nearest Neighbour (KNN) sebanyak 99.80%, Logistic Regression (LR) sebanyak 95.55% serta Naive Bayes (NB) sebanyak 87.51%. Seterusnya, menurut kajian Ashiku dan Dagli, mereka menggunakan Convolutional Neural Network (CNN) untuk membangunkan sistem IDS berdasarkan dataset UNSW-NB15 dan model mencapai ketepatan sebanyak 94.4%. Kajian Otoom dan rakan-rakan pada tahun 2023, mengesan serangan brute force dalam rangkaian IoT menggunakan Deep Neural Network (DNN) dan mencapai ketepatan melebihi 99.6% dengan kadar positif palsu yang sangat rendah, sekitar 0.02%. Akhirnya, kajian oleh Almahadeen dan

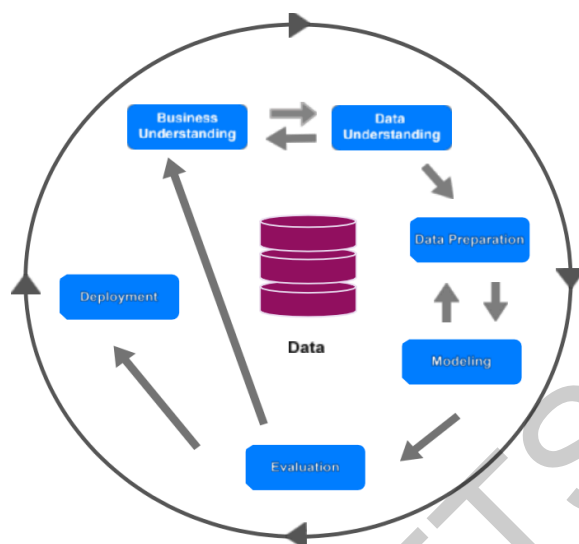
rakan-rakannya pada tahun 2024, menggabungkan Autoencoder dan Multilayer Perceptron (MLP) untuk pengesanan dalam sektor kewangan dan mencapai ketepatan sebanyak 99%.

Jadual 1.0 Perbandingan algoritma pembelajaran mesin dan pembelajaran mendalam bagi kajian literatur

Artikel	Algoritma	Rumusan	Ketepatan (%)
Hamza & Surayh Al-Janabi 2024	RF	Fokus pada serangan SSH dan FTP	99.90
	LR		89.17
	NB		48.61
A.Sharma et al. 2024	DT	Gabungkan model ramalan dan anomali serta mengkaji reka bentuk ciri dan pengurangan false positive.	99.96
	KNN		99.80
	LR		95.55
	NB		87.51
Ashiku & Dagli 2021	CNN	Penerokaan IDS berdasarkan CNN	94.40
Otoom et al. 2023	DNN	Pengesanan brute force dalam IoT serta kadar false positive yang rendah	99.70
Almahadeen et al. 2024	Autoencoder + MLP	Fokus pada keselamatan kewangan dengan menggabungkan model untuk meningkatkan keupayaan pengesanan	99.00

3.0 METODOLOGI

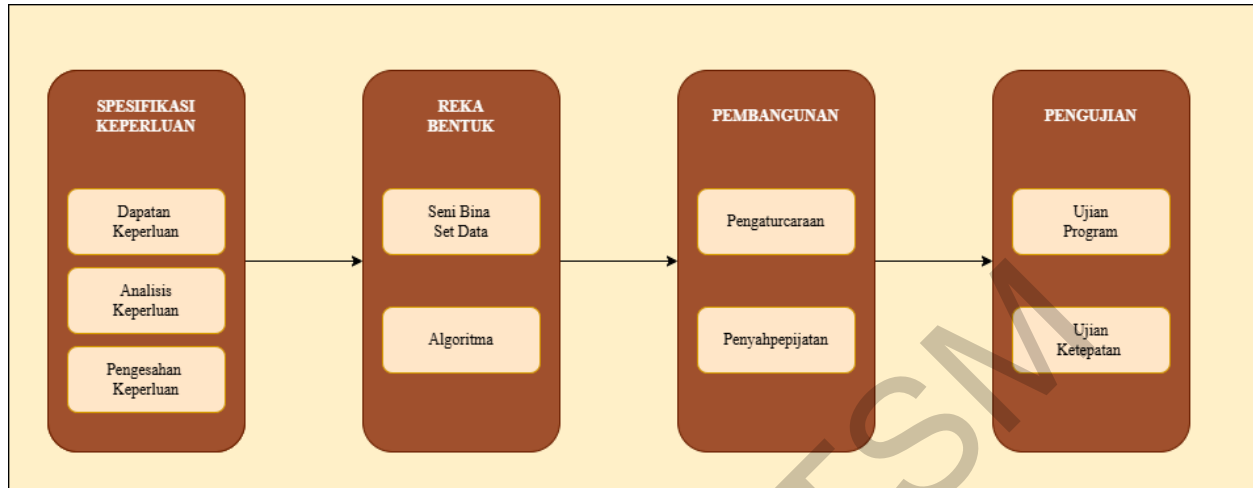
Metodologi yang digunakan bagi projek ini adalah selaras dengan enam fasa model *Cross Industry Standard Process for Data Mining* (CRISP-DM) yang merangkumi kitaran hidup pembangunan projek. Enam fasa tersebut adalah pemahaman bisnes, pemahaman data, penyediaan data, pemodelan, penilaian dan penyebaran. Rajah 1.0 merupakan gambaran visual terhadap fasa kitaran hidup metodologi yang terlibat dalam projek ini.



Rajah 1.0 Seni bina model CRISP-DM

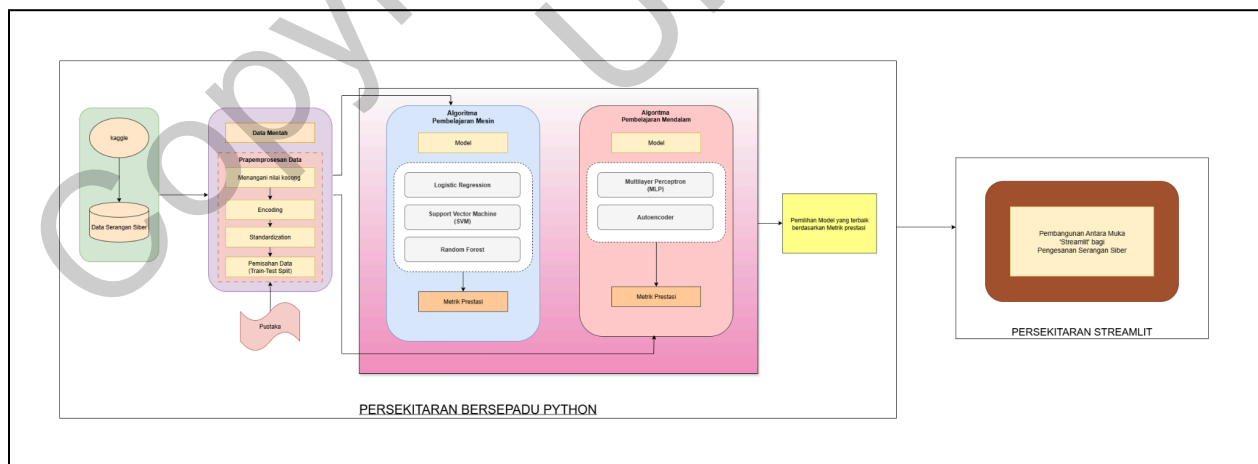
Sumber: Saepulrohman et al. 2025

Projek ini bermula dengan fasa pemahaman bisnes dimana fasa ini mengenal pasti objektif keselamatan siber dan keperluan projek, dengan fokus kepada pengesanan aktiviti pencerobohan siber berdasarkan tingkah laku pengguna. Bagi fasa pemahaman data, melibatkan analisis set data dari Kaggle yang mengandungi 11 atribut dan 9,537 rekod termasuk semakan isu seperti nilai hilang dan ketidakseimbangan kelas. Seterusnya, fasa penyediaan data di mana proses prapemprosesan dilakukan seperti *One-Hot Encoding* dan normalisasi nilai berangka, serta pengasingan data kepada ciri (X) dan label (Y). Fasa pemodelan pula, dimana beberapa algoritma digunakan termasuk Logistic Regression, Support Vector Machine, Random Forest, Multilayer Perceptron dan Autoencoder untuk melatih model pengesanan serangan. Selepas itu, model dinilai menggunakan metrik penilaian seperti *accuracy*, *precision*, *recall* dan *F1-Score* serta *ROC-AUC*. Akir sekali, fasa penyebaran dimana model akhir diintegrasikan dalam antara muka pengguna menggunakan Streamlit untuk analisis dan pemantauan keputusan secara interaktif dan masa nyata. Rajah 2.0 menunjukkan ringkasan bagi struktur pembangunan projek yang dibincangkan.



Rajah 2.0 Struktur Pembangunan Projek

Struktur pembangunan projek ini merangkumi beberapa fasa utama bagi memastikan kelancaran pelaksanaan. Ia bermula dengan penentuan spesifikasi keperluan berdasarkan analisis masalah. Seterusnya, fasa reka bentuk dijalankan bagi merancang sistem pengesanan pencerobohan siber. Fasa pembangunan dibahagikan kepada dua persekitaran iaitu Python dan Streamlit, bertujuan untuk membina dan menyahpejijat atur cara. Akhir sekali, fasa pengujian dilaksanakan bagi menilai prestasi model. Rajah 3.0 memberikan ringkasan bagi reka bentuk seni bina projek ini.



Rajah 3.0 Reka bentuk seni bina

Projek ini menggunakan set data pencerobohan siber yang diperoleh dari Kaggle dan dimuat naik oleh Samudrala pada 10 februari 2025. Set data tersebut mengandungi 9.537 baris

dan 11 ciri berkaitan dengan aktiviti sesi rangkaian pengguna. Dalam fasa prapemprosesan, ciri tidak relevan seperti *session_id* dibuang, ciri kategori ditukar kepada nilai berangka menggunakan *LabelEncoder* dan data dinormalisasi. Pemilihan ciri penting dilakukan menggunakan *ANOVA F-test* bagi mengenal pasti ciri yang paling signifikan terhadap sasaran *attack_detected*. Hasilnya, tiga ciri paling relevan dikenalpasti iaitu *login_attempts*, *failed_logins*, dan *ip_reputation_score*. Pemilihan ini membantu meningkatkan prestasi model dan mengurangkan kerumitan pengiraan.

Sebagai pengesahan, *Random Forest Classifier* digunakan untuk mengukur kepentingan ciri berdasarkan purata pengurangan impuriti. Hasilnya menunjukkan *ip_reputation_score* sebagai ciri paling dominan, diikuti oleh *failed_logins* dan *login_attempts*. Penemuan ini konsisten dengan hasil ANOVA, membuktikan bahawa ketiga-tiga ciri ini sangat relevan dalam pengesanan pencerobohan siber. Data bersih kemudiannya dibahagikan kepada set latihan dan set ujian. Algoritma pembelajaran mesin dan mendalam digunakan untuk melatih menilai prestasi model menggunakan metrik penilaian. Model dengan prestasi terbaik akan dipilih dan diintegrasikan ke dalam antara muka pengguna menggunakan Streamlit. Platform ini membolehkan pengguna memasukkan input seperti bilangan log masuk dan log masuk gagal serta skor reputasi IP, untuk diuji oleh model dalam mengesan pencerobohan siber secara masa nyata.

4.0 HASIL

4.1 Pembangunan Model Pembelajaran Mesin

Pada bahagian ini merangkumi pembangunan model pembelajaran mesin dan pembelajaran mendalam bagi mengesan aktiviti pencerobohan siber. Model ini dilatih menggunakan set data yang telah dipraproses, dan dinilai menggunakan metrik penilaian. Algoritma pembelajaran mesin yang digunakan adalah Logistic Regression, Support Vector Machine dan Random Forest, yang diimport daripada pustaka scikit-learn. Setiap model dilatih dengan data *X_train_scaled* dan *y_train*, serta diuji dengan *X_test_scaled*. Fungsi *predict()* digunakan untuk klasifikasi, manakala *predict_proba()* digunakan untuk kebarangkalian bagi pengiraan *ROC-AUC*.

Bagi SVM, parameter *probability=True* diaktifkan untuk membolehkan pengiraan kebarangkalian menggunakan pendekatan *Platt Scalling*. Hanya kebarangkalian kelas 1 iaitu serangan digunakan untuk menilai prestasi model. Hasil daripada semua metrik penilaian disimpan dalam struktur *ml_results* bagi membolehkan perbandingan sistematik antara model-model dilakukan, bagi mengenal pasti model terbaik.

4.2 Pembangunan Model Pembelajaran Mendalam

Selepas pembangunan model tradisional, model pembelajaran mendalam dibangunkan menggunakan Multilayer Perceptron (MLP) dan Autoencoder bagi mengesan aktiviti pencerobohan serangan siber.

Model MLP dibina menggunakan pustaka TensorFlow dan Keras dalam bentuk sequential, merangkumi lapisan *Dense (fully connected)*, *Dropout* bagi mencegah *overfitting*, serta *BatchNormalization* untuk menstabilkan pembelajaran. Pengoptimum Adam digunakan bersama *EarlyStopping* dan *ReduceLROnPlateau* bagi mengawal kadar pembelajaran dan menghentikan latihan apabila prestasi tidak lagi bertambah baik. Manakala Autoencoder digunakan untuk mengenal pasti corak data normal dan mengesan anomali berdasarkan ralat pembinaan semula (*reconstruction error*).

Model ini dilatih hanya dengan data normal dan threshold yang ditetapkan berdasarkan peratusan ke-95 ralat tertinggi data normal iaitu 1.2706. Mana-mana data dengan ralat melebihi nilai ini dianggap sebagai serangan, manakala selebihnya diklasifikasikan sebagai normal. Kaedah ini membantu meminimumkan *false positive* dan *false negative*, menjadikan sistem lebih tepat dalam mengesan aktiviti pencerobohan serangan siber yang sebenar. Kedua-dua model dinilai menggunakan metrik prestasi bagi memilih model yang terbaik untuk integrasi dalam antara muka pengguna.

4.3 Perbandingan Model Pengesanan Aktiviti Serangan Siber

Jadual 2.0 menunjukkan ringkasan prestasi bagi kesemua model pengesanan serangan rangkaian yang telah dibangunkan.

Jadual 2.0 Ringkasan prestasi model pengesanan serangan rangkaian

Model	Accuracy	Precision	Recall	F1-Score	ROC-AUC
Logistic Regression (LR)	0.7180	0.7043	0.6366	0.6687	0.7758
Support Vector Machine (SVM)	0.8643	1.0000	0.6964	0.8210	0.8453
Random Forest (RF)	0.7794	0.7500	0.7597	0.7548	0.8518
Multilayer Perceptron (MLP)	0.8632	0.9983	0.6940	0.8194	0.8563
Autoencoder	0.5639	0.5827	0.0868	0.1510	0.5717

Kajian membandingkan prestasi lima model yang berbeza dalam tugas pengesanan aktiviti pencerobohan siber berdasarkan lima metrik penilaian iaitu *Accuracy*, *Precision*, *Recall*, *F1-Score*, dan *ROC-AUC*. Hasil penilaian menunjukkan bahawa model Support Vector Machine (SVM) dan Multilayer Perceptron (MLP) mencatatkan prestasi terbaik secara konsisten dalam pelbagai aspek.

Model SVM merekodkan ketepatan tertinggi iaitu 86.43% dan *precision* yang sempurna menunjukkan kebolehnya mengenal pasti data serangan dengan sangat tepat tanpa menghasilkan amaran palsu (*false positive*). Nilai *recall* sebanyak 0.6964, *F1-score* sebanyak 0.8210 serta *ROC-AUC*, 0.8453. Hal ini mengukuhkan lagi prestasinya yang seimbang antara kepekaan dan ketepatan. Model MLP turut menunjukkan prestasi kompetitif dengan ketepatan sebanyak 86.32% dan membuktikan keupayaannya dalam menangkap pola kompleks dalam data. Sementara itu, model RF menunjukkan prestasi yang sederhana dengan ketepatannya sebanyak 77.94% dan metrik lain yang seimbang. Model LR mencatat prestasi lebih rendah dengan ketepatan sebanyak 71.80% dan *ROC-AUC*, 0.7758, manakala model Autoencoder merekodkan prestasi paling lemah dengan ketepatany sebanyak 56.39% sahaja.

Secara keseluruhannya, model SVM telah dipilih sebagai model terbaik untuk projek ini kerana ia menunjukkan prestasi yang stabil dan konsisten dengan merekod ketepatan yang paling tinggi serta mencapai *precision* yang sempurna tanpa mengorbankan keupayaan mengenal pasti serangan. Kebolehnya menghasilkan keputusan tepat tanpa amaran palsu menjadikannya

sangat sesuai untuk diintegrasikan ke dalam sistem pengesanan aktiviti pencerobohan siber berasaskan antaramuka pengguna Streamlit.

4.4 Pembangunan Streamlit

Aplikasi Streamlit dibina sebagai antara muka pengguna bagi menguji model pengesanan aktiviti pencerobohan siber. Pengguna memasukkan tiga input utama iaitu bilangan percubaan log masuk, log masuk gagal dan skor reputasi IP. Data ini dinormalisasikan menggunakan scaler yang telah dilatih, kemudian dianalisis oleh model SVM untuk menjana ramalan. Keputusan ditunjukkan dalam bentuk teks dan visual interkatif bagi membantu pengguna memahami sama ada input tersebut menunjukkan potensi serangan siber. Rajah 4.0 menunjukkan fail komponen penting dimaut naik.

```

77 # ----- Load Model & Scaler -----
78 try:
79     #model = tf.keras.models.load_model("mlp_model.keras", compile=False)
80     model = joblib.load("svm_model.pkl")
81     with open("scaler.pkl", "rb") as f:
82         scaler: StandardScaler = pickle.load(f)
83 except Exception as e:
84     st.error(f"❌ Failed to load model or scaler: {e}")
85     st.stop()

```

Rajah 4.0 Pengekodan bagi memuat naik fail komponen penting

Seterusnya, Rajah 5.0 dan Rajah 6.0 masing-masing menunjukkan kod bagi bahagian input bagi sistem ini dan hasilnya pada halaman antara muka pengguna.

```

93 # ----- Input Section -----
94 st.markdown("### 🚨 Enter login activity to detect potential intrusion:")
95 col1, col2 = st.columns(2)
96 with col1:
97     login_attempts = st.number_input("Login Attempts", min_value=0, step=1)
98 with col2:
99     failed_login = st.number_input("Failed Login Attempts", min_value=0, step=1)
100
101 ip_score = st.slider("IP Reputation Score (0 = Bad, 1 = Good)", 0.0, 1.0, step=0.01)

```

Rajah 5.0 Pengekodan bagi bahagian input sistem.

Rajah 6.0 Antara muka pengguna bagi bahagian input pengguna

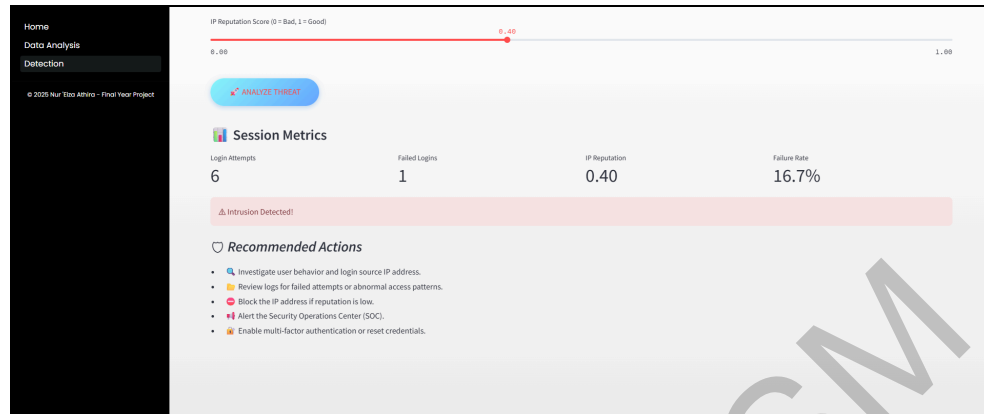
Antaramuka sistem menggunakan komponen Streamlit bagi membolehkan pengguna memasukkan tiga input utama, iaitu, cubaan log masuk, log masuk gagal dan skor reputasi IP. Input disusun secara kemas, manakala skor reputasi IP dipaparkan dalam bentuk gelangsar dengan julat 0.0 hingga 1.0. Pengguna memasukkan data secara manual, dan apabila butang "Analyze Threat" ditekan, sistem menjalankan proses ramalan menggunakan model yang telah dilatih. Keputusan dipaparkan dalam bentuk teks dan visual. Rajah 7.0 dan 8.0 masing-masing menunjukkan kod analisis dan ramalan pengesanan pencerobohan serta keputusan pengesanan berdasarkan input.

```

110 # ----- Prediction & Output -----
111 if analyze:
112     # Validate Inputs
113     input_data = np.array([[login_attempts, failed_login, ip_score]])
114     # Scale Inputs
115     input_scaled = scaler.transform(input_data)
116
117     prediction = model.predict(input_scaled)
118     #probability = float(prediction[0][0])
119     probability = model.predict_proba(input_scaled)[0][1]
120     # Determine if model predicts attack
121     model_says_attack = probability > 0.5

```

Rajah 7.0 Pengekodan analisis dan ramalan pengesanan pencerobohan



Rajah 8.0 Keputusan pengesanan serangan berdasarkan input pengguna

Dalam sistem ini, data input terlebih dahulu dinormalisasikan menggunakan *StandardScaler* yang telah dilatih, bagi memastikan konsistensi data dalam proses ramalan. Data yang telah diskalakan kemudiannya dihantar ke model Support Vector Machine (SVM) yang telah dipilih sebagai model terbaik. Model ini menghasilkan nilai *threshold* menggunakan fungsi *predict_proba()*, dan jika kebarangkalian melebihi *threshold* 0.5, aktiviti tersebut diklasifikasikan sebagai serangan siber. *Threshold*, 0.5 digunakan kerana ia adalah nilai lalai (*default threshold*) yang secara umum mewakili garis pemisah antara dua kelas dalam masalah klasifikasi binari. Contohnya, jika model mengembalikan nilai > 0.5 , sistem akan secara automatik memaparkan amaran “Intrusion Detected!”. Sebagai langkah susulan, sistem turut mencadangkan tindakan mitigasi seperti menyemak log aktiviti, menyekat IP yang mencurigakan, menghubungi Pusat Operasi Keselamatan (*Security Operations Center*, SOC), serta melaksanakan pengesanan berbilang faktor (*Multi-Factor Authentication*, MFA) atau penetapan semula kata laluan sebagai langkah mitigasi.

5.0 KESIMPULAN

Secara keseluruhannya, laporan ini telah menghuraikan dengan terperinci proses pembangunan sistem pengesanan aktiviti pencerobohan siber bermula dari penyediaan dan pra-pemprosesan data, pembahagian set data, pembinaan model pembelajaran mesin dan pembelajaran mendalam, sehinggalah kepada pelaksanaan sistem dalam bentuk antara muka pengguna (*User Interface*) menggunakan *Streamlit*. Pendekatan ini membolehkan pengguna akhir berinteraksi secara langsung dengan sistem melalui antara muka yang mesra pengguna tanpa memerlukan pengetahuan teknikal yang mendalam.

Pendekatan yang digunakan dalam sistem ini, membolehkan sistem membuat keputusan yang lebih tepat dan dapat dijelaskan. Kaedah ini juga membantu dalam mengurangkan kadar amaran palsu, *false positive* dan *false negative*. Selain itu, pembangunan antara muka menggunakan Streamlit telah membuktikan bahawa teknologi sumber terbuka mampu menyediakan platform interaktif dan mesra pengguna, membolehkan pengguna akhir menjalankan penilaian ancaman secara masa nyata dengan lebih mudah dan efisien. Reka bentuk ini juga memfokuskan kepada kebolegunaan (*usability*) dan kebolehpercayaan (*reliability*), yang merupakan elemen penting dalam sistem keselamatan siber moden.

Sistem yang dibangunkan telah mencapai objektif projek dari segi ketepatan, kebolehcapaian, dan kesesuaian untuk digunakan dalam situasi sebenar. Hasil daripada pelaksanaan ini membuktikan bahawa model pembelajaran mesin boleh digabungkan secara efektif dengan elemen visualisasi dan analisis sokongan untuk membentuk sebuah sistem pengesanan ancaman siber yang praktikal dan berkesan.

6.0 PENGHARGAAN

Alhamdulillah, bersyukur ke hadrat Allah S.W.T, Yang Maha Pengasih lagi Maha Penyayang kerana dengan izin dan berkat-Nya telah memberikan saya kesihatan yang baik, masa yang cukup dan kematangan fikiran untuk saya menyiapkan laporan projek tahun akhir ini. Saya dengan rendah hati menadah tangan tanda kesyukuran, terharu kerana telah menyiapkan projek tahun akhir saya bagi memenuhi sebahagian daripada syarat memperolehi Ijazah Sarjana Muda Sains Komputer dengan Kepujian dalam tempoh masa yang ditetapkan. Saya sekali lagi berasa sangat bersyukur kerana telah diberikan kekuatan untuk saya mengatasi segala masalah dan kekangan yang timbul sepanjang projek ini dijalankan. Saya sedar bahawa dugaan yang hadir ini telah membentuk peribadi saya untuk sentiasa peka dan disiplin dalam melakukan sesuatu kerja.

Seterusnya, saya ingin merakamkan setinggi-tinggi penghargaan kepada penyelia saya, Ts. Rohizah Binti Abd. Rahman, atas bimbingan dan dorongan sepanjang saya menyiapkan usulan projek ini tanpa lelah dan jemu. Beliau telah banyak memberikan tunjuk ajar serta bantuan, teguran dan nasihat yang sangat berguna kepada saya. Kepakaran dan pandangan beliau yang tidak ternilai telah memainkan peranan penting dalam membentuk hala tuju dan hasil usulan projek saya. Saya amat berterima kasih atas masa dan usaha yang beliau telah korbankan untuk saya dan kerja saya. Tidak lupa juga kepada Fakulti Teknologi dan Sains Maklumat dan Universiti Kebangsaan Malaysia atas kemudahan dan sumber yang disediakan kerana telah banyak membantu saya dalam menyiapkan projek ini serta pensyarah-pensyarah khususnya kepada mereka yang pernah menabur ilmu pengetahuan kepada saya sepanjang pengajian.

Akhir sekali, dengan kesempatan ini saya menghulurkan setinggi-tinggi ucapan penghargaan kepada beberapa pihak yang lain sama ada terlibat secara langsung mahupun tidak langsung dalam menjayakan penyiapan projek ini. Tidak lupa juga kepada semua ahli keluarga saya khususnya ayahanda dan bonda tercinta, Encik Eikhmerizal Bazura bin Baharin dan Puan Suriwati Azilah binti Abdullah serta sahabat seperjuangan saya, Intan Humaira, Iliana Hanin, Nur Ain Batrisyia dan Puteri Alissa yang telah memberi sokongan dan semangat serta sentiasa berdoa terhadap kejayaan saya di universiti. Saya juga ingin memohon maaf sekiranya terdapat kesalahan sepanjang pelaksanaan projek akhir tahun ini. Sekian, terima kasih.

7.0 RUJUKAN

- A. Sharma, H. Babbar and A. K. Vats, "Empowering Security: Machine Learning Solutions for Detecting Brute Force Attacks," 2024 4th Asian Conference on Innovation in Technology (ASIANCON), Pimari Chinchwad, India, 2024, pp. 1-5, doi: 10.1109/ASIANCON62057.2024.10838096.
- Alkhwaja, I., Albugami, M., Alkhwaja, A., Alghamdi, M., Abahussain, H., Alfawaz, F., ... & Min-Allah, N. (2023). Password cracking with brute force algorithm and dictionary attack using parallel programming. *Applied Sciences*, 13(10), 5979.
- Almahadeen, L., Mahadin, G. A., Santosh, K., Aarif, M., Deb, P., Syamala, M., & Bala, B. K. (2024). Enhancing Threat Detection in Financial Cyber Security Through Auto Encoder-MLP Hybrid Models. *International Journal of Advanced Computer Science & Applications*, 15(4).
- Ashiku, L., & Dagli, C. (2021). Network intrusion detection system using deep learning. *Procedia Computer Science*, 185, 239-247.
- Ghanem, W. A. H., Ghaleb, S. A. A., Jantan, A., Nasser, A. B., Saleh, S. A. M., Ngah, A., ... & Abiodun, O. I. (2022). Cyber intrusion detection system based on a multiobjective binary bat algorithm for feature selection and enhanced bat algorithm for parameter optimization in neural networks. *IEEE Access*, 10, 76318-76339.
- Hamza, A. A., & surayh Al-Janabi, R. J. (2024). Detecting brute force attacks using machine learning. In *BIO Web of Conferences* (Vol. 97, p. 00045). EDP Sciences.
- Otoom, A.F.; Eleisah, W.; Abdallah, E.E. Deep Learning for Accurate Detection of Brute Force attacks on IoT Networks. *Procedia Comput. Sci.* 2023, 220, 291–298
- Saepulrohman, A., Chairunnas, A., Denih, A., & Yasibang, N. D. S. (2025). Optimization of Stock Price Prediction Using Long Short-Term Memory (LSTM) Algorithm and Cross-Industry Standard Process Approach for Data Mining (CRISP-DM). *International Journal of Electronics and Communications Systems*, 5(1), 19-30.

Zhang, C., Jia, D., Wang, L., Wang, W., Liu, F., & Yang, A. (2022). Comparative research on network intrusion detection methods based on machine learning. *Computers & Security*, 121, 102861.

Nur 'Eiza Athira binti Eikhmerizal Bazura (A195719)
Ts. Rohizah binti Abd Rahman
Fakulti Teknologi & Sains Maklumat,
Universiti Kebangsaan Malaysia