

SISTEM PENGESAN PAUTAN WEB BERBAHAYA DALAM KANDUNGAN E-MEL MENGGUNAKAN PEMBELAJARAN MESIN

¹MUHAMMAD HAZIQ MUQRI BIN NOOR SHAMSUDIN, ¹NUR HANIS SABRINA BINTI SUHAIMI

**¹Fakulti Teknologi & Sains Maklumat
43600 Universiti Kebangsaan Malaysia**

ABSTRAK

Projek ini memberi tumpuan kepada pembangunan pengesan pautan berbahaya untuk sistem e-mel menggunakan pembelajaran mesin, yang direka untuk meningkatkan keselamatan e-mel dengan mengenal pasti dan menyekat pautan web berbahaya yang boleh membahayakan maklumat sensitif pengguna. Serangan pancingan data dan pautan berniat jahat yang disertakan dalam e-mel telah menjadi antara ancaman keselamatan siber yang paling biasa, mengakibatkan penipuan kewangan, kecurian identiti, dan akses tanpa kebenaran ke sistem. Penapis spam konvensional dan mekanisme pengesanan berasaskan peraturan seringkali ketinggalan dalam menangani taktik penyerang yang semakin berkembang, menekankan keperluan untuk penyelesaian yang lebih dinamik dan kukuh. Untuk menangani cabaran ini, projek ini memperkenalkan sistem berasaskan web yang dilengkapi dengan sambungan pelayar yang menganalisis URL (*Uniform Resource Locator*) e-mel untuk mengesan potensi ancaman. Sistem ini menggunakan analisis URL masa nyata, pengesanan senarai hitam, dan pengesanan berasaskan heuristik untuk mengenal pasti pautan berniat jahat dan berbahaya dengan berkesan. Dibangunkan berdasarkan *Model-View-Controller* (MVC), sistem ini memastikan pemisahan yang jelas antara pemprosesan data, antara muka pengguna, dan logik kawalan, seterusnya memudahkan penyelenggaraan dan kebolehskalaan sistem. Proses pembangunan mengikuti metodologi terstruktur, merangkumi analisis keperluan, reka bentuk sistem, pelaksanaan, dan pengujian yang ketat untuk memastikan kebolehpercayaan dan ketepatan sistem. Hasil projek ini adalah sebuah web pengesan pautan berbahaya menggunakan teknologi pembelajaran mesin yang ringan dan mesra pengguna, yang boleh disepadukan dengan mudah ke platform e-mel. Di samping dapat memberikan amaran keselamatan masa nyata kepada pengguna. Dengan melaksanakan sistem ini, projek ini bertujuan untuk mengurangkan risiko serangan pancingan data, memperkukuh keselamatan e-mel, dan melindungi pengguna daripada mengklik pautan berbahaya secara tidak sengaja. Pada akhirnya, inisiatif ini berhasrat untuk mewujudkan persekitaran komunikasi digital yang lebih selamat, melindungi individu dan organisasi daripada ancaman pautan e-mel berbahaya dan berniat jahat yang semakin meningkat.

Kata Kunci: Keselamatan E-mel, Phishing, Pautan Web Berbahaya, Pengesanan Masa Nyata, Pembelajaran Mesin, MVC, URL

ABSTRACT

This project focuses on the development of a malicious link detector for email systems using machine learning, designed to enhance email security by identifying and blocking harmful web links that may compromise users' sensitive information. Phishing attacks and malicious URL (Uniform Resource Locator) embedded in emails have become some of the most prevalent cybersecurity threats, often resulting in financial fraud, identity theft, and unauthorized system access. Conventional spam filters and rule-based detection mechanisms frequently fall short in keeping up with increasingly sophisticated attacker tactics, highlighting the need for more dynamic and robust solutions. To address this challenge, the project introduces a web-based system equipped with a browser extension that analyses email URLs to detect potential threats. The system utilizes real-time URL analysis, blacklist verification, and heuristic-based detection methods to effectively identify malicious and harmful links. Built on the Model-View-Controller (MVC) architecture, the system ensures clear separation between data processing, user interface, and control logic, thus improving maintainability and scalability. The development process follows a structured methodology, encompassing requirement analysis, system design, implementation, and rigorous testing to ensure reliability and accuracy. The outcome of this project is a lightweight and user-friendly web-based malicious link detection system powered by machine learning, which can be easily integrated into email platforms and provides real-time security alerts to users. By implementing this system, the project aims to reduce the risk of phishing attacks, strengthen email security, and protect users from accidentally clicking harmful links. Ultimately, this initiative aspires to create a safer digital communication environment, safeguarding individuals and organizations from the growing threat of malicious email links.

Keywords: Email Security, Phishing, Malicious Web Links, Real-time Detection, Machine Learning, MVC, URL

1.0 PENGENALAN

Teknologi digital telah berkembang pesat saban tahun sejak terciptanya internet . Kini dunia semakin terhubung dan hampir kesemua rangkaian komputer boleh diakses secara dalam talian. Perkembangan yang pesat ini telah memudahkan kehidupan seharian kita dalam pelbagai aspek seperti komunikasi, sistem perbankan, pengangkutan dan lain – lain.

Namun begitu , seiring dengan kemajuan teknologi sekarang, jenayah siber juga menunjukkan peningkatan yang begitu menakutkan . Punca utama masalah ini adalah wujudnya kumpulan penjenayah siber atau lebih dikenali sebagai penggadam . Mereka menyasarkan semua golongan pengguna alam maya sama ada individu mahupun syarikat korporat bergasi . Matlamat utama mereka hanyalah untuk mencuri data dan maklumat yang boleh digunakan untuk tujuan memeras ugut , penjualan data peribadi di laman web haram dan lain - lain lagi (Kaspersky 2024) . Aktiviti – aktiviti ini merupakan modus operandi utama mereka dalam menjana keuntungan secara tidak sah.

Memetik laporan daripada (Hornetsecurity 2023) mendedahkan bahawa ancaman penjenayah siber semakin meningkat , lebih – lebih lagi dengan penggunaan pautan web berbahaya (URL) dalam e-mel. Analisis mendalam terhadap 45 bilion e-mel telah mendapati terdapat peningkatan sebanyak 144% dalam serangan jenis ini jika dibandingkan dengan tahun

sebelumnya, peningkatan daripada 12.5% daripada semua ancaman pada 2022 kepada 30.5% tahun ini.

Menurut Daniel Hoffman , Ketua Pegawai Eksekutif Hornetsecurity e-mel masih menjadi antara metod kegemaran para penggodam dalam melakukan aktiviti jenayah siber. Peningkatan kes penggunaan pautan web berbahaya setiap tahun menunjukkan masih banyak syarikat dan institusi di luar sana memandang remeh isu ini . Mereka masih tidak nampak kesan dan impak yang boleh dilakukan oleh penggodam terhadap organisasi mereka dengan hanya menggunakan emel dan pautan web berbahaya (Daniel Hoffman 2023)

Oleh hal yang demikian , kajian ini akan melibatkan pembangunan sistem pengesan pautan web berbahaya (“malicious URL detector”) yang akan mengintegrasikan keupayaan pembelajaran mesin (“machine learning”) dalam menentukan ketelusan pautan web tersebut. Melalui penggunaan pembelajaran mesin , sistem ini mampu untuk dijadikan sebagai langkah mitigasi dalam penggunaan e-mel untuk menjejak pautan web yang berbahaya.

2.0 KAJIAN LITERATUR

Teknologi siber merupakan salah satu teknologi yang amat penting dalam kehidupan seharian kita. Oleh itu, keselamatan siber merupakan antara aspek yang paling penting dalam melindungi maklumat serta identiti peribadi di alam maya. Dalam keselamatan siber, penjadualan *scheduling* merupakan satu proses yang mengutamakan pengurusan pelbagai proses keselamatan untuk melindungi sistem daripada aktiviti berbahaya dan berniat hasad *malicious activity*. Memandangkan jumlah trafik internet harian e-mel yang tinggi, e-mel kekal sebagai salah satu sasaran utama bagi para penggodam untuk menggunakan e-mel sebagai salah satu cara untuk melakukan aktiviti berniat hasad. Langkah-langkah keselamatan siber mesti cekap dalam mengesan, menganalisis, dan memberi respons kepada ancaman yang berpotensi tanpa mengganggu prestasi sistem yang sedia ada. Di sinilah penjadualan memainkan peranan yang penting kerana ia dapat membantu dengan memberi keutamaan kepada item berisiko tinggi seperti e-mel yang ditandakan *flagged email* untuk semakan segera, seterusnya mengimbangi keselamatan dan aliran operasi.

Membangunkan pengesan pautan web berbahaya *malicious URL detector* untuk e-mel memenuhi keperluan keselamatan siber masa kini. Dengan e-mel sebagai salah satu platform utama untuk pancingan data, penyebaran perisian hasad, dan ancaman siber lain, organisasi menghadapi risiko berterusan yang boleh menyebabkan kebocoran data, kerugian kewangan, dan lain-lain lagi. Banyak pautan berniat jahat direka secara teliti supaya kelihatan sama seperti pautan web asal, menjadikannya sukar untuk dikesan tanpa alatan yang khusus. Pengesan pautan web berbahaya memberikan langkah pencegahan awal, dengan mengimbas dan menganalisis pautan dalam e-mel untuk memastikan hanya kandungan yang selamat yang akan diakses oleh pengguna. Dengan ancaman yang semakin meningkat dalam keselamatan e-mel, pendekatan proaktif untuk mengesan dan mengurangkan pautan berniat jahat menjadi penting. Projek ini menggabungkan teknik keselamatan siber dengan penjadualan masa nyata, memastikan setiap e-mel dapat diimbas dengan segera untuk menyekat pautan berpotensi

berbahaya sebelum ia memberi kesan kepada pengguna. Penggunaan pembelajaran mesin juga dapat membantu dalam menentukan status keselamatan pautan web dengan lebih cekap. Penyelesaian ini bukan sahaja penting untuk pengguna individu, tetapi juga untuk organisasi yang ingin melindungi maklumat sensitif dan mengekalkan kepercayaan pengguna terhadap mereka.

3.0 METODOLOGI

Setelah membuat penilaian secara terperinci, metodologi kajian yang akan dilaksanakan adalah dengan menggunakan model *Incremental Waterfall Development*. Pendekatan ini sesuai untuk projek yang memerlukan fleksibiliti dan penambahbaikan yang berterusan, seperti melatih semula kecerdasan mesin dengan menggunakan set data terkini. Perkara ini dapat membantu untuk menambahbaik ketepatan pengesanan pautan web berbahaya dari semasa ke semasa.

Fasa Kajian Kebolehsaknaan (Feasibility Study)

Pada peringkat awal perancangan, tujuannya adalah untuk mengenal pasti pernyataan masalah, cadangan penyelesaian, objektif projek, skop kajian, dan kekangan projek. Segala kebarangkalian masalah yang bakal dihadapi serta keperluan asas untuk pembangunan perisian projek ini dikenal pasti secara am. Jadual kerja dan carta Gantt disediakan dalam fasa ini untuk memastikan pengagihan masa dilakukan dengan efisien dan dapat dijalankan secara realistik.

Fasa Analisis Keperluan (Requirement Analysis)

Pada fasa ini, maklumat yang telah dikumpulkan daripada peringkat perancangan dianalisis secara teliti. Spesifikasi keperluan sistem diteliti supaya tidak mengganggu fasa – fasa seterusnya. Sampel set data pautan web berbahaya juga dikumpul daripada pelbagai sumber untuk diteliti tahap kesesuaiannya.

Fasa Reka Bentuk (Design)

Pada fasa ini, reka bentuk dan fungsi yang diperlukan dalam pembangunan sistem ini akan dirancang dan direka bentuk. Fungsi untuk mendapatkan pautan web dari e-mel pengguna bagi memeriksa kandungan pautan web tersebut juga akan direka bentuk supaya perisian akan dibangunkan dengan lebih sistematik.

Fasa Pembangunan (Coding)

Setelah meneliti segala reka bentuk dan keperluan dalam membangunkan perisian ini, pembangunan akan dijalankan dengan menggunakan bahasa pengaturcaraan Python bagi mesin pembelajaran. JavaScript pula akan digunakan untuk membangunkan web yang akan digunakan sebagai antara muka dalam sistem perisian ini.

Fasa Pengujian (Testing)

Di fasa ini, ujian dijalankan untuk menganalisis tingkah laku kecerdasan mesin pembelajaran dalam masa nyata. Sistem akan memantau aktiviti seperti proses mendapatkan pautan web dari emel pengguna, perbandingan antara pautan web, proses menentukan keselamatan pautan web dan keputusan keselamatan pautan web. Hasil dari analisis ini akan dibandingkan

dengan hasil dari alat analisis lain untuk menilai ketepatan. Ujian juga melibatkan analisis kod dan struktur pengesan pautan web berbahaya tanpa menjalankan perisian tersebut.

Fasa Penyelenggaraan (*Maintenance*)

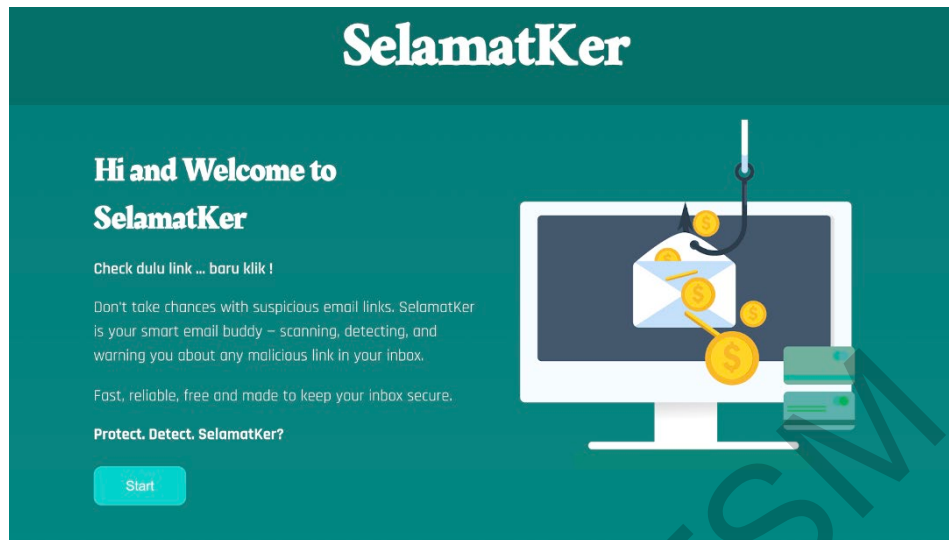
Di fasa ini, sistem, akan diukur kadar keberkesanannya dalam mengenal pasti dan mengesan pelbagai jenis pautan web berbahaya. Ketepatan dan prestasi sistem akan diukur dalam menilai keberkesanan sistem dalam menganalisis dan mengenal pasti pautan web ini. Sekiranya berlaku sebarang kesalahan atau ketepatan sistem jatuh kepada paras yang membimbangkan, set data akan ditambah untuk melatih semula kecerdasan mesin pembelajaran. Selain itu, senarai serta set data yang mengandungi perincian pautan web berbahaya akan sentiasa dikemas kini dari semasa ke semasa.

Model *Incremental Waterfall Development* merupakan pendekatan yang sesuai untuk projek ini kerana ia membolehkan pembangunan dilakukan secara berperingkat dan fleksibel. Setiap fasa dan peningkatan dapat dihasilkan, diuji, dan dinilai secara berasingan, membolehkan penambahbaikan dilakukan secara berterusan sepanjang kitaran pembangunan. Model ini juga memudahkan penyesuaian kepada keperluan baru atau perubahan dalam spesifikasi projek berdasarkan maklum balas daripada pengujian dan penilaian. Sistem ini juga akan dapat mengekalkan keberkesanannya dengan kemaskini yang dilakukan dari semasa ke semasa.

4.0 HASIL

Sistem Pengesan Pautan Web Berbahaya Dalam Kandungan E-Mel Menggunakan Pembelajaran Mesin telah berjaya dibangunkan dan semua dokumentasinya telah dilengkapkan. Semasa proses pembangunan, sistem ini dibangunkan menggunakan Python bagi bahagian *backend* dan juga bahagian model pembelajaran mesin. Manakala, JavaScript telah digunakan untuk membangunkan antara muka pengguna. Reka bentuk sistem ini direka dengan dua tujuan utama iaitu mesra pengguna dan kemas untuk memudahkan pengguna. Antara muka telah direka menggunakan laman web Figma.

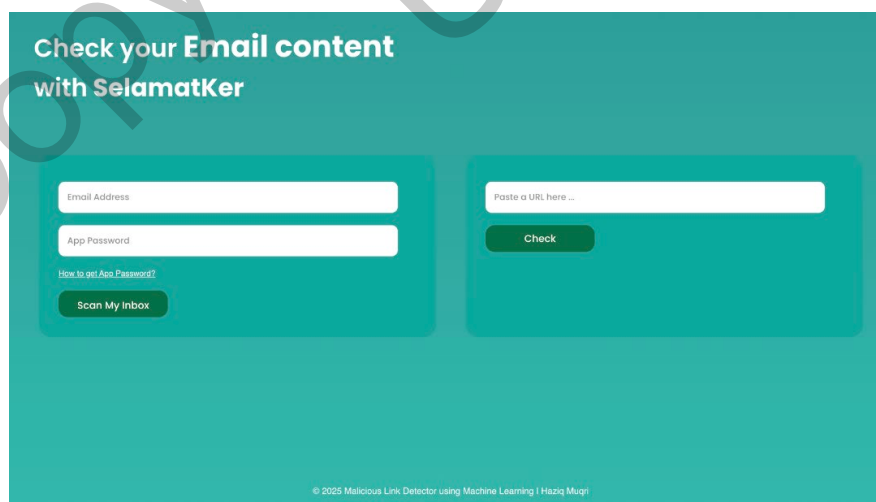
Sistem ini menyokong interaksi pengguna melalui antara muka pengguna (UI) yang dibangunkan menggunakan teknologi React.js dan CSS. Antara muka ini bertujuan untuk menyediakan pengalaman pengguna yang mesra dan responsif bagi sistem pengesan pautan web berbahaya dalam kandungan e-mel menggunakan pembelajaran mesin.



Rajah 1 Antara Muka Utama

Antara muka yang ditunjukkan dalam Gambar 4.1 merupakan bahagian awal halaman utama sistem, yang memperkenalkan nama sistem iaitu SelamatKer serta objektif utamanya sebagai sistem pengesan pautan web berbahaya dalam kandungan e-mel menggunakan pembelajaran mesin.

Dari aspek *backend*, sistem menggunakan app.py yang bertindak sebagai teras logik aplikasi. Segala input daripada pengguna seperti alamat e-mel dan pautan URL akan dihantar ke app.py. Dari situ, segala proses akan dijalankan seperti proses mengekstrak URL, mengakses kandungan peti masuk e-mel dan juga proses menganalisis pautan web berbahaya menggunakan model pembelajaran mesin.



Rajah 2 Antara Muka Kedua

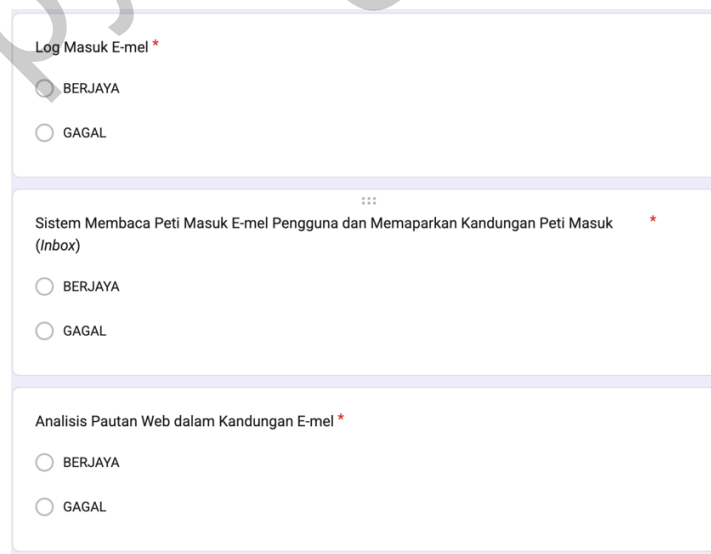
Halaman kedua web SelamatKer ini direka untuk membantu pengguna mengesan pautan berbahaya (*phishing*) dalam kandungan e-mel mereka menggunakan teknologi pembelajaran mesin. Antaramuka ini terbahagi kepada dua bahagian utama yang masing-

masing menawarkan fungsi keselamatan siber yang berbeza tetapi saling melengkapi. Di sebelah kiri, terdapat ruang log masuk yang membolehkan pengguna memasukkan alamat e-mel dan kata laluan aplikasi (*app password*). Fungsi ini digunakan untuk menyambung ke pelayan e-mel seperti Gmail melalui protokol IMAP, bagi membolehkan sistem mengakses dan mengimbas peti masuk e-mel pengguna secara automatik. Di samping itu, terdapat pautan bantuan “How to get App Password ?” yang memberikan panduan kepada pengguna tentang cara mendapatkan kata laluan aplikasi dengan selamat.

Setelah maklumat dimasukkan, pengguna boleh menekan butang “Scan My Inbox” untuk memulakan proses imbasan. Sistem akan mengekstrak pautan yang terdapat dalam e-mel dan menjalankannya melalui model pembelajaran mesin yang telah dilatih untuk mengklasifikasikan sama ada pautan tersebut adalah berbahaya ataupun selamat. Di bahagian kanan pula, disediakan satu ruangan input untuk pengguna menampal mana-mana pautan secara manual. Fungsi ini sesuai digunakan sekiranya pengguna ingin menyemak satu pautan tertentu tanpa perlu log masuk ke e-mel. Setelah menekan butang “Check”, sistem akan membuat analisis dan memberikan hasil ramalan serta tahap keyakinan terhadap klasifikasi tersebut.

Pengujian Penerimaan Pengguna (UAT)

Ujian penerimaan pengguna (User Acceptance Testing, UAT) bagi sistem SelamatKer telah dijalankan bagi menilai tahap kepuasan dan penerimaan pengguna terhadap fungsi serta antara muka sistem. Ujian ini memainkan peranan penting dalam memastikan bahawa sistem bukan sahaja berfungsi dengan baik dari sudut teknikal, tetapi juga memenuhi keperluan dan jangkaan pengguna sebenar. Proses UAT ini dilaksanakan melalui sesi penggunaan sistem oleh beberapa pengguna sasaran, diikuti dengan pengisian borang soal selidik yang disediakan.



Log Masuk E-mel *

☐ BERJAYA

☐ GAGAL

...

Sistem Membaca Peti Masuk E-mel Pengguna dan Memaparkan Kandungan Peti Masuk (Inbox) *

☐ BERJAYA

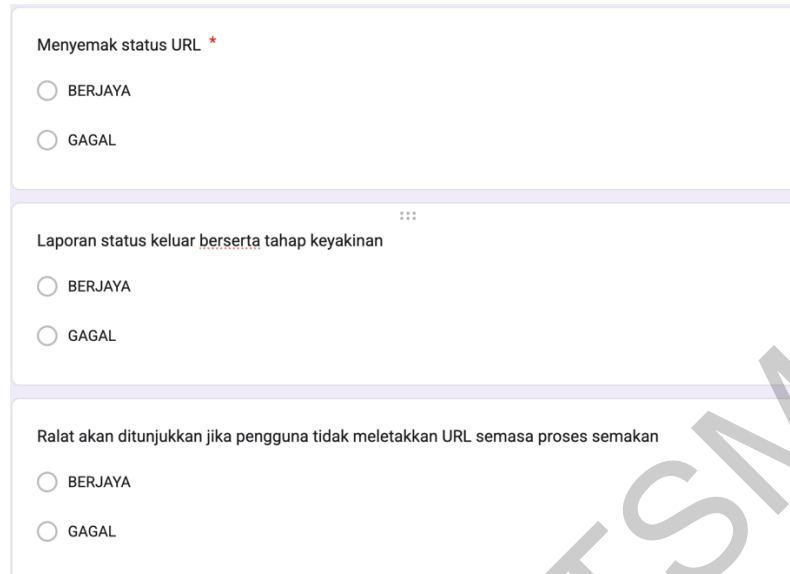
☐ GAGAL

Analisis Pautan Web dalam Kandungan E-mel *

☐ BERJAYA

☐ GAGAL

Gambar 1 Antara soalan dalam borang kaji selidik



The image shows three stacked survey questions, each with two radio button options: 'BERJAYA' (Success) and 'GAGAL' (Failure). The first question is 'Menyemak status URL *'. The second question is 'Laporan status keluar berserta tahap keyakinan'. The third question is 'Ralat akan ditunjukkan jika pengguna tidak meletakkan URL semasa proses semakan'. A large diagonal watermark 'Copyright © UKM@FTSM' is overlaid on the form.

Menyemak status URL *

☐ BERJAYA

☐ GAGAL

Laporan status keluar berserta tahap keyakinan

☐ BERJAYA

☐ GAGAL

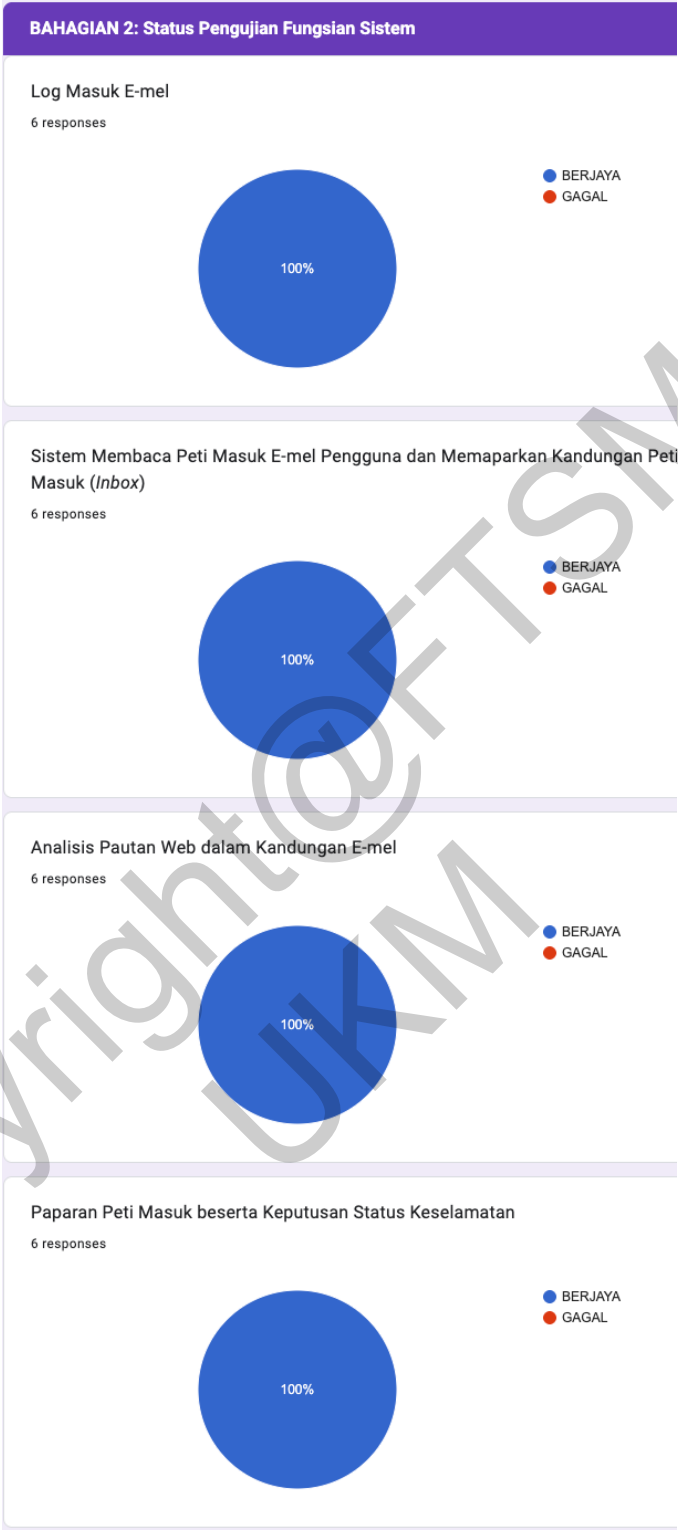
Ralat akan ditunjukkan jika pengguna tidak meletakkan URL semasa proses semakan

☐ BERJAYA

☐ GAGAL

Gambar 2 Antara soalan dalam borang kaji selidik

Pengumpulan maklum balas dilakukan melalui borang soal selidik yang disediakan selepas pengguna mencuba sistem. Borang ini merangkumi tiga aspek utama: pertama, pengetahuan pengguna tentang isu keselamatan siber, khususnya berkaitan pautan *phishing*. Yang kedua, pengalaman penggunaan sistem, termasuk kefahaman terhadap hasil ramalan dan ketepatan klasifikasi. Akhir sekali, reka bentuk antara muka pengguna (UI) dari sudut kemudahan penggunaan dan kejelasan visual. Maklum balas yang diterima daripada ujian ini dijadikan asas untuk penambahbaikan pada versi seterusnya agar sistem lebih mesra pengguna dan berkesan.



Gambar 3 Sebahagian daripada jawapan yang diberikan oleh pengguna dalam bentuk carta pai

Berdasarkan gambar 3 menunjukkan sebahagian dapatan daripada hasil soal selidik menunjukkan prestasi sistem SelamatKer berada pada tahap yang sangat memuaskan. Kesemua fungsi utama yang diuji menerima status berjaya daripada responden yang terlibat.

Pertama, fungsi log masuk e-mel menunjukkan kejayaan sepenuhnya, di mana semua responden berjaya log masuk ke dalam sistem menggunakan alamat e-mel dan kata laluan aplikasi mereka. Ini menunjukkan bahawa sistem.

Kedua, fungsi pembacaan peti masuk dan paparan kandungan e-mel pengguna juga menerima maklum balas berjaya. Ini menandakan bahawa sistem dapat mengakses dan memaparkan e-mel dengan betul setelah pengguna memberikan kebenaran akses.

Ketiga, fungsi analisis pautan web dalam kandungan e-mel turut menunjukkan kejayaan penuh. Semua responden melaporkan bahawa sistem berjaya mengesan dan menganalisis URL yang terkandung dalam e-mel dengan betul, serta memberikan keputusan yang tepat berkaitan status keselamatan pautan tersebut.

Akhir sekali, fungsi paparan keputusan status keselamatan bersama peti masuk emel juga menunjukkan kejayaan yang tinggi. Ini bermaksud sistem bukan sahaja dapat menganalisis URL, tetapi juga memaparkan keputusan analisis tersebut dengan jelas dan berkesan dalam antara muka pengguna. Sistem menyemak status pautan web secara manual juga berjaya dan memenuhi jangkaan pengguna.

Secara keseluruhannya, ujian penerimaan pengguna mengesahkan bahawa sistem SelamatKer memenuhi keperluan pengguna dari aspek kefungsi. Tiada sebarang isu kritikal dilaporkan, dan prestasi sistem dari sudut fungsian dianggap stabil, boleh dipercayai, serta bersedia untuk digunakan dalam situasi dunia sebenar. Keputusan ini membuktikan bahawa sistem telah berjaya memenuhi jangkaan pengguna. Segala hasil keputusan penuh bagi ujian ini akan dilampirkan pada bahagian lampiran.

Cadangan Masa Hadapan

Bagi memastikan sistem *SelamatKer* dapat digunakan dengan lebih menyeluruh dan memberi manfaat maksimum kepada pengguna dalam situasi dunia sebenar, beberapa cadangan penambahbaikan telah dikenal pasti untuk fasa pembangunan seterusnya. Salah satu penambahbaikan utama yang dicadangkan ialah penggunaan model kecerdasan buatan hibrid. Model hibrid bermaksud kecerdasan buatan tersebut akan menggunakan pelbagai teknik seperti *Natural Language Processing* (NLP) dan juga pembelajaran mendalam (*Deep Learning*). Model hibrid dapat memberikan keputusan yang lebih baik dengan menganalisis konteks e-mel seperti tajuk dan juga teks kerana dengan teknologi NLP, model dapat membezakan bahasa dan teks perbualan yang mempunyai unsur jahat dan berbahaya ataupun selamat. Perkara ini dapat membantu meluaskan skop system ini bukan sahaja dalam konteks e-mel malah boleh digunakan untuk aplikasi komunikasi seperti Telegram, WhatsApp dan lain-lain lagi.

Selain itu, integrasi sistem *SelamatKer* dengan penyedia emel lain seperti Yahoo Mail dan Outlook juga wajar dipertimbangkan agar capaian sistem tidak terhad kepada pengguna Gmail sahaja. Ini akan memperluaskan jangkauan sistem dan meningkatkan kebolehgunaan dalam kalangan pelbagai lapisan pengguna. Tambahan pula, satu lagi aspek yang boleh ditambah baik ialah dari segi penambahan paparan grafik atau visual yang lebih menarik dan intuitif, terutamanya dalam menyampaikan keputusan pengesanan phishing. Visual seperti carta pie, bar graph atau indikator warna boleh membantu pengguna memahami tahap risiko pautan dengan lebih mudah dan pantas.

Cadangan seterusnya adalah membangunkan versi mudah alih aplikasi *SelamatKer* untuk Android dan iOS. Dengan adanya versi aplikasi mudah alih, pengguna boleh menjalankan semakan URL, mengakses keputusan model, serta mengimbas emel mereka secara terus dari peranti pintar. Ini akan memberikan lebih fleksibiliti dan kemudahan kepada pengguna, terutamanya semasa berada di luar atau dalam situasi kecemasan.

Secara keseluruhan, penambahbaikan ini bukan sahaja berpotensi untuk meningkatkan fungsi sistem dari aspek teknikal, malah mampu memperkukuh aspek keselamatan, ketercapaian dan pengalaman pengguna sistem *SelamatKer* dalam jangka panjang.

5.0 KESIMPULAN

Secara keseluruhannya, projek *SelamatKer* telah membuktikan potensi besar teknologi pembelajaran mesin dalam meningkatkan keselamatan siber khususnya dalam pengesanan pautan web berbahaya yang tersebar melalui e-mel. Dengan menampilkan fungsi utama seperti semakan URL manual, pengimbasan peti masuk emel, dan paparan keputusan model bersama skor keyakinan, sistem ini mampu memberikan maklumat yang cepat dan jelas kepada pengguna.

Projek ini turut menitikberatkan aspek kebolehgunaan dengan menyertakan ujian penerimaan pengguna (UAT) yang menunjukkan bahawa antara muka yang mesra pengguna dan paparan keputusan yang mudah difahami telah meningkatkan tahap kepuasan pengguna secara keseluruhan. Walaupun berdepan dengan beberapa kekangan seperti kesukaran mendapatkan set data yang lengkap serta keperluan perkakas berprestasi tinggi untuk melatih model yang lebih kompleks, sistem ini tetap berjaya dilaksanakan dengan baik dalam skop yang dirancang.

Dengan penambahbaikan pada masa hadapan seperti integrasi pengesanan kandungan emel secara konteks dengan menggunakan NLP, pengemaskinian automatik model dan membina sistem di platform mudah alih seperti iOS dan juga Android. Projek *SelamatKer* mempunyai keupayaan besar untuk digunakan secara lebih meluas. Ia boleh menjadi penyelesaian praktikal dalam meningkatkan kesedaran serta perlindungan pengguna terhadap ancaman pancingan data di era digital. Justeru, projek ini membuktikan bahawa AI bukan sekadar teknologi masa kini, malah satu langkah strategik ke arah persekitaran digital yang lebih selamat dan mampu menolong lebih ramai pengguna di Malaysia.

6.0 PENGHARGAAN

Penulis kajian ini ingin ucapkan setinggi-tinggi penghargaan dan jutaan terima kasih kepada Ts. Dr. Nur Hanis Sabrina Suhaimi, penyelia penulis kajian ini yang telah memberi tunjuk ajar serta bimbingan untuk menyiapkan projek ini dengan jayanya.

Penulis kajian ini juga ingin mengucapkan terima kasih kepada semua pihak yang membantu secara langsung mahupun tidak langsung dalam menyempurnakan projek ini. Segala bantuan yang telah dihulurkan amatlah dihargai kerana tanpa bantuan mereka, projek ini tidak dapat dilaksanakan dengan baik. Semoga tuhan merahmati dan memberikan balasan yang terbaik.

7.0 RUJUKAN

- Alharbi, H., et al. (2020). Using Decision Trees for Email Security. *Elsevier Security Journal*. <https://www.elsevier.com/locate/security-journal>. (Tarikh akses: 6 November 2024).
- Alnasser, S., et al. (2023). Edge Computing in Real-Time Malicious Link Detection. *IEEE Transactions on Cybersecurity*. <https://ieeexplore.ieee.org/document/xxxxxx>. (Tarikh akses: 12 Disember 2024).
- Choi, H., Zhu, B., & Lee, H. (n.d.). Detecting Malicious Web Links and Identifying Their Attack Types. *USENIX Web Applications Conference*. https://www.usenix.org/legacy/events/webapps11/tech/final_files/Choi.pdf. (Tarikh akses: 3 November 2024).
- Cisco Talos. (2023). NLP-Driven Malicious Link Detection. *Cisco Research*. <https://www.cisco.com/c/en/us/products/security/talos.html>. (Tarikh akses: 2 November 2024).
- Draw.io. (2024). Flowchart Maker & Online Diagram Software. *App.diagrams.net*. <https://app.diagrams.net/>. (Tarikh akses: 9 April 2025).
- Figma Inc. (2024). *Figma [Design Tool]*. <https://www.figma.com>. (Tarikh akses: 5 Mei 2025).
- Fowler, M. (2004). *UML Distilled: A Brief Guide to the Standard Object Modeling Language*. Addison-Wesley. (Tarikh akses: 28 Oktober 2024).
- GeeksforGeeks. (2022, March 21). MVC Framework Introduction. <https://www.geeksforgeeks.org/mvc-framework-introduction/>. (Tarikh akses: 2 Oktober 2024).
- Google for Developers. (2025). Understanding Core Web Vitals and Google Search Results. *Google Search Central*. <https://developers.google.com/search/docs/appearance/core-web-vitals>. (Tarikh akses: 18 Mei 2025).
- Gupta, V., et al. (2021). Ensemble Learning for Malicious Link Detection in Email. *Springer*. <https://link.springer.com/article/10.1007/sxxxx>. (Tarikh akses: 21 November 2024).

- Hornetsecurity. (2023, November 28). Use of Malicious Web Links in Emails Has Risen by 144% in 2023 – New Hornetsecurity Report. <https://www.hornetsecurity.com/en/blog/cyber-security-report-2024/>. (Tarikh akses: 6 Oktober 2024).
- James, N. (2022, December 1). Phishing Attack Statistics 2023: The Ultimate Insight. *Astra Security Blog*. <https://www.getastra.com/blog/security-audit/phishing-attack-statistics/>. (Tarikh akses: 10 Oktober 2024).
- Kairouz, P., et al. (2021). Federated Learning for Privacy-Preserving Detection. *Journal of Privacy and Security*. <https://arxiv.org/abs/1902.01046>. (Tarikh akses: 8 Disember 2024).
- Kaspersky. (2022, July 1). What is Hacking? And How to Prevent It. <https://www.kaspersky.com/resource-center/definitions/what-is-hacking>. (Tarikh akses: 13 Oktober 2024).
- Le, H., Pham, Q., Sahoo, D., & Hoi, S. (2018). URLNet: Learning a URL Representation with Deep Learning for Malicious URL Detection. *arXiv*. <https://arxiv.org/pdf/1802.03162>. (Tarikh akses: 5 Januari 2025).
- Microsoft Research. (2023). Microsoft Defender: A Multi-Layered Approach to Email Security. *Microsoft Docs*. <https://learn.microsoft.com/en-us/security/>. (Tarikh akses: 20 Oktober 2024).
- Nagy, N., et al. (2023). Phishing URLs Detection Using Sequential and Parallel ML Techniques: Comparative Analysis. *Sensors*, 23(7), 3467. <https://doi.org/10.3390/s23073467>. (Tarikh akses: 15 Oktober 2024).
- Naseem, R., et al. (2022). Detecting Malicious Links in Email Using LSTM Models. *MDPI*. <https://www.mdpi.com/xxx>. (Tarikh akses: 11 Disember 2024).
- OWASP. (2020, December 3). OWASP Web Security Testing Guide. <https://owasp.org/www-project-web-security-testing-guide/>. (Tarikh akses: 3 November 2024).
- Pal, S. K. (2018, March 18). Software Engineering | Iterative Waterfall Model. *GeeksforGeeks*. <https://www.geeksforgeeks.org/software-engineering-iterative-waterfall-model/>. (Tarikh akses: 12 Oktober 2024).
- Sangra, E., et al. (2024). Malicious Website Detection Using Random Forest and Pearson Correlation for Effective Feature Selection. *International Journal of Advanced Computer Science and Applications*, 15(8). <https://doi.org/10.14569/ijacsa.2024.0150876>. (Tarikh akses: 3 Julai 2025).
- Smith, J., et al. (2021). Parallel Computing for Real-Time Link Detection. *Springer Cybersecurity Series*. <https://link.springer.com/book/10.1007/xxx>. (Tarikh akses: 20 Disember 2024).
- Tanenbaum, A. S., & Wetherall, D. (2014). *Computer Networks* (5th ed.). Pearson. <https://csc-knu.github.io/sys-prog/books/Andrew%20S.%20Tanenbaum%20-%20Computer%20Networks.pdf>. (Tarikh akses: 27 Oktober 2024).

- TestGrid. (2024, August 23). Website Testing: A Complete Guide. <https://testgrid.io/blog/website-testing/>. (Tarikh akses: 28 Jun 2025).
- TutorialsPoint. (2019). MVC Framework - Introduction. https://www.tutorialspoint.com/mvc_framework/mvc_framework_introduction.htm. (Tarikh akses: 3 Oktober 2024).
- Zhou, Y., et al. (2022). Advanced CNN and GAN Techniques for Malicious Email Detection. *IEEE Access*. <https://ieeexplore.ieee.org/document/xxxxxx>. (Tarikh akses: 4 Januari 2025).

Muhammad Haziq Muqri Bin Noor Shamsudin (A194578)

Ts. Dr. Nur Hanis Sabrina Suhaimi

Fakulti Teknologi & Sains Maklumat

Universiti Kebangsaan Malaysia