

BAB I

PENDAHULUAN

A. Latar Belakang

Transformasi digital yang masif telah menempatkan internet sebagai infrastruktur kritis bagi berbagai aktivitas. Namun, seiring dengan tingginya tingkat adopsi teknologi tersebut, ekosistem digital menghadapi ancaman keamanan siber yang terus berevolusi. Salah satu bentuk ancaman yang paling persisten dan merugikan saat ini adalah serangan *Phishing*. Berbeda dengan serangan konvensional yang mengeksploitasi celah keamanan perangkat lunak, *Phishing* menggunakan pendekatan rekayasa sosial (*social engineering*) yang secara langsung menargetkan aspek psikologis pengguna. Pelaku memanipulasi korban menggunakan situs web palsu yang identik dengan portal resmi untuk mencuri informasi kredensial yang bersifat konfidensial. Saat ini, peretas bahkan menggunakan teknik manipulasi tautan (*Uniform Resource Locator / URL*) tingkat tinggi, seperti *typosquatting* dan *homoglyph attacks*, yang membuat tautan palsu hampir mustahil dibedakan oleh mata telanjang [1], [2].

Secara tradisional, mekanisme pertahanan utama yang digunakan oleh berbagai peramban web dan perangkat lunak keamanan adalah metode *Blacklist URL*. Namun, sistem ini memiliki kelemahan arsitektural yang fatal karena bersifat reaktif. Sistem *blacklist* sepenuhnya tidak berdaya menghadapi serangan *zero-day Phishing*, yaitu URL jahat yang baru saja didaftarkan oleh peretas dan belum sempat diverifikasi ke dalam basis data keamanan global [3].

Untuk menutup celah kerentanan *zero-day* tersebut, sistem keamanan modern mulai mengintegrasikan pendekatan kecerdasan buatan (*Artificial Intelligence*), khususnya cabang *Machine Learning*. Pendekatan ini difokuskan pada analisis karakteristik leksikal dari string URL tanpa perlu memuat konten halaman web secara utuh (*crawling*), sehingga proses deteksi bisa dilakukan jauh lebih cepat dan aman [4].

Algoritma mempelajari pola matematis dari struktur URL seperti panjang domain, jumlah karakter khusus, hingga tingkat keacakan karakter untuk membedakan secara proaktif URL yang sah (*legitimate*) dan anomali [5]. Dalam menangani klasifikasi berbasis ruang fitur berdimensi tinggi seperti struktur

leksikal URL, algoritma *Support Vector Machine* (SVM) menawarkan pendekatan geometris yang sangat tangguh. SVM bekerja dengan mencari *hyperplane* (garis batas tegas) optimal untuk memisahkan kelas data. Meskipun demikian, efektivitas komputasi dan akurasi klasifikasi pada algoritma SVM sangat bergantung pada pemilihan fungsi pemetaan atau *Kernel* yang tepat, terutama ketika menghadapi data URL nyata yang distribusinya sering kali tumpang tindih [6].

Karakteristik ekstraksi leksikal URL memunculkan tantangan analitik tersendiri. Di satu sisi, fitur-fitur tersebut mungkin memiliki korelasi linier yang kuat. Pada kondisi ini, penggunaan *Kernel Linear* menjadi sangat efisien secara komputasi dan sering kali menjadi standar (*baseline*) yang kokoh tanpa risiko *overfitting* [7]. Di sisi lain, evolusi *Phishing* sering kali menghasilkan pola manipulasi teks yang sangat kompleks dan tidak linier. Untuk mengatasi hal ini, *Kernel Sigmoid* menawarkan fungsi aktivasi yang perilakunya menyerupai jaringan saraf tiruan (*Artificial Neural Network*), sehingga berpotensi menangkap pola manipulasi *Phishing* yang lebih rumit pada tingkat dimensi yang berbeda [8].

Selain mencari akurasi tertinggi, tantangan utama dalam implementasi deteksi *Phishing* di dunia nyata adalah menekan tingkat alarm palsu (*False Positive*) yang dapat memblokir situs web sah dan mengganggu kenyamanan pengguna. Oleh karena itu, pengujian performa pengklasifikasi harus mempertimbangkan konfigurasi tingkat kehati-hatian (*threshold*) yang terukur [9].

Berdasarkan urgensi dan landasan teknologi tersebut, penelitian ini mengusulkan Perbandingan *Kernel Linear* dan *Kernel Sigmoid* pada Algoritma SVM untuk Deteksi Website *Phishing*. Melalui pengujian yang komprehensif pada himpunan data berjumlah 30.000 sampel URL seimbang (*balanced dataset*) dan penerapan ambang batas probabilitas sebesar 0.8, penelitian ini bertujuan menemukan fungsi pemetaan (*Kernel*) mana yang paling presisi, akurat, dan optimal sebagai solusi keamanan sistem cerdas pendeteksi *Phishing*.

B. Perumusan Masalah

Bagaimana perbandingan performa antara penggunaan *Kernel Linear* dan *Kernel Sigmoid* pada algoritma *Support Vector Machine* (SVM) dalam mendeteksi *website Phishing* berdasarkan ekstraksi fitur leksikal URL?

C. Batasan Penelitian

Pembahasan dalam penelitian ini dibatasi pada beberapa aspek berikut agar tetap spesifik dan terarah:

1. Proses deteksi *Phishing* sepenuhnya berfokus pada karakteristik leksikal dari string URL tanpa melakukan perayapan (*crawling*) terhadap konten visual atau kode sumber halaman web
2. Algoritma *Machine Learning* yang digunakan sebagai mesin pengambil keputusan dibatasi pada *Support Vector Machine* (SVM) dengan membandingkan dua fungsi *Kernel*: *Linear* dan *Sigmoid*
3. Sistem menerapkan teknik standardisasi skala data (*Feature Scaling*) pada tahap prapemrosesan agar komputasi jarak pada model SVM dapat berjalan secara optimal
4. Data latih dan uji bersumber dari PhishTank untuk URL *Phishing* serta Tranco List untuk URL aman, dengan total 30.000 sampel seimbang (*balanced dataset*)
5. Penentuan akhir kelas *Phishing* pada model menggunakan batas probabilitas (*threshold*) sebesar 0.8 untuk meminimalisir alarm palsu (*False Positive*)
6. Kemampuan deteksi sistem bergantung sepenuhnya pada representasi pola leksikal yang terdapat di dalam data latih dan data uji. Sistem tidak melakukan pencocokan atau validasi kemiripan dengan nama domain *website* resmi, karena model murni berfokus pada anomali struktur matematis URL dan karakteristik *link phishing* yang sangat beragam (seperti penggunaan domain acak).

D. Tujuan Penelitian

Penelitian ini bertujuan untuk mengetahui dan membandingkan tingkat performa (Akurasi, Presisi, *Recall*, dan *F1-Score*) antara *Kernel Linear* dan *Kernel Sigmoid* pada algoritma SVM untuk menemukan model klasifikasi yang paling optimal dalam mendeteksi *website Phishing*.

E. Manfaat Penelitian

Manfaat yang diharapkan dari pelaksanaan penelitian ini adalah sebagai berikut:

1. Manfaat Akademis

Penelitian ini diharapkan dapat memberikan kontribusi signifikan terhadap perkembangan literatur ilmiah di bidang keamanan siber, khususnya mengenai efektivitas sistem deteksi proaktif berbasis kecerdasan buatan (*Machine Learning*). Secara teoritis, hasil penelitian ini akan menjadi referensi penting dalam memahami efektivitas pemilihan fungsi *Kernel* (*Linear* dan *Sigmoid*) pada ruang fitur berdimensi tinggi untuk klasifikasi URL dinamis.

2. Manfaat bagi Penulis

Melalui pengerjaan tugas akhir ini, penulis dapat mengembangkan kompetensi teknis yang lebih dalam, mulai dari tahapan prapemrosesan data skala besar, perancangan arsitektur pemodelan komparatif yang efisien, hingga penguasaan metrik evaluasi performa model yang akurat. Proses ini tidak hanya mengasah kemampuan pemecahan masalah secara sistematis, tetapi juga mempersiapkan penulis dengan keahlian praktis yang relevan dan dibutuhkan dalam industri teknologi informasi dan keamanan digital saat ini.

3. Manfaat bagi Masyarakat

Secara praktis, penelitian ini bertujuan untuk menghasilkan purwarupa sistem perlindungan proaktif yang dapat membantu masyarakat luas dan pengguna internet dalam memitigasi risiko pencurian identitas digital maupun kerugian finansial akibat serangan *Phishing*. Dengan adanya sistem yang mampu mendeteksi ancaman *zero-day Phishing* melalui analisis leksikal yang cepat dan cerdas, pengguna diharapkan dapat berselancar di dunia maya dengan rasa aman yang lebih tinggi tanpa terganggu oleh alarm palsu (*False Positive*). Selain itu, rancang bangun sistem yang efisien ini dapat dijadikan rujukan bagi para praktisi atau pengembang perangkat lunak keamanan untuk membangun alat bantu perlindungan seperti ekstensi peramban atau sistem penyaring pada gerbang jaringan yang lebih tangguh.