

BAB I PENDAHULUAN

1.1 Latar Belakang

Cyber Crime mengalami peningkatan yang signifikan beberapa tahun terakhir. Pada tahun 2024, sistem deteksi mereka menemukan rata-rata 467.000 file berbahaya setiap hari, hasil ini meningkat sebanyak 14% dibandingkan tahun sebelumnya [1]. Salah satu ancaman siber yang paling umum adalah *phishing*, sebuah teknik penipuan yang bertujuan memperoleh informasi sensitif dari halaman login pengguna dan dikirimkannya ke *malicious server* [2]. Serangan ini sering memanfaatkan kelemahan psikologis korban melalui teknik manipulasi digital, seperti mengirimkan pesan yang tampak berasal dari institusi resmi atau individu yang dikenal, untuk mendorong korban mengungkapkan data pribadi atau keuangan mereka. Pada tahun 2023 dan paruh pertama 2024, Indonesia menghadapi 4.046 serangan *phishing* yang terutama menargetkan sektor layanan informasi [3].

Bukan hanya menargetkan individu, akan tetapi juga perusahaan besar, lembaga keuangan, dan instansi pemerintahan. metode *phishing* yang dilakukan para kriminal siber tersebut umumnya membuat situs web palsu yang tidak terlihat mencurigakan dan justru terlihat bisa di percaya sebagai cara menipu korbannya, agar membocorkan kredensial login, dan mencuri informasi sensitif [4]. Kejahatan ini juga telah diakui sebagai sumber utama dari banyaknya penipuan dan pelanggaran data yang dilakukan oleh kriminal siber [5].

Dalam upaya mendeteksi *phishing*, berbagai cara dan teknik telah dikembangkan, mulai dari penggunaan *blacklist*, analisis heuristik, hingga pendekatan berbasis kecerdasan buatan (*artificial intelligence*). Namun, metode konvensional, seperti *blacklist* seringkali gagal dalam mendeteksi serangan *phishing* baru atau *zero-day attacks*, yang belum terdokumentasi dalam basis data keamanan [6]. Oleh karena itu, penelitian ini berfokus pada penerapan metode *machine learning* yang dapat meningkatkan akurasi deteksi *phishing* dengan analisis karakteristik *URL* dan *HTML* dari situs web secara lebih mendalam [7].

Penggabungan teknik *feature selection* berbasis *boosting* menggunakan pendekatan *ensemble learning* dilakukan untuk meningkatkan efisiensi dan akurasi dalam proses klasifikasi *URL phishing*. Proses ini dimulai dengan pemilihan fitur-fitur paling relevan menggunakan teknik pemeringkatan fitur, dimana setiap fitur diberikan peringkat yang didasarkan pada tingkat kepentingannya dalam model. Salah satu metode yang digunakan untuk melakukan pemeringkatan ini adalah *Random Forest Feature Importance*, di mana fitur yang memiliki kontribusi terbesar dalam membagi data pada pohon keputusan diurutkan berdasarkan nilai pentingnya. Dengan cara ini, hanya fitur-fitur yang memiliki pengaruh signifikan terhadap klasifikasi yang dipilih untuk dilatih lebih lanjut.

Setelah seleksi fitur, *lightweight model* dan juga *Random Forest* dilatih menggunakan hanya subset fitur yang sudah terpilih. Pendekatan ini bertujuan untuk mengurangi kompleksitas model dengan menghilangkan fitur yang tidak relevan, mempercepat waktu pelatihan, serta meminimalkan risiko *overfitting* yang bisa terjadi jika terlalu banyak fitur yang tidak berkontribusi secara signifikan tetap digunakan.

Ensemble learning, yang menggabungkan beberapa model pembelajaran untuk memberikan hasil klasifikasi yang lebih stabil dan akurat, bekerja sangat baik dengan seleksi fitur ini. Dalam konteks ini, *Random Forest* sebagai salah satu metode *ensemble learning* juga berperan dalam seleksi fitur, karena ia menggunakan teknik *bagging (Bootstrap Aggregating)* salah satu teknik *ensemble learning* yang bertujuan untuk mengurangi variabilitas model tanpa meningkatkan bias secara signifikan, sehingga menghasilkan model yang lebih stabil dan akurat. Dengan banyak pohon keputusan yang memungkinkan identifikasi fitur yang paling berpengaruh dalam klasifikasi phishing. *XGBoost*, *LightGBM*, dan *CatBoost* yang juga merupakan model berbasis *boosting*, mengandalkan peningkatan kesalahan iteratif yang memungkinkan setiap model belajar dari kesalahan model sebelumnya. Kombinasi *feature subset ranking* dan *ensemble learning* ini saling mendukung dengan memilih fitur yang tepat dan meningkatkan akurasi deteksi

phishing, sekaligus mengurangi risiko *overfitting* dan meningkatkan efisiensi komputasi.

Strategi ini di prediksi dapat meningkatkan performa sistem deteksi *phishing* dengan hasil akurasi yang lebih tinggi dibandingkan metode tradisional. Namun, penelitian ini memiliki pendekatan berbeda dengan memfokuskan pada kombinasi pemeringkatan subset fitur menggunakan *Random Forest* dan penggunaan *lightweight models* secara terpisah maupun kolektif. Berbeda dengan penelitian sebelumnya, pendekatan ini juga mengevaluasi efektivitas masing-masing model secara komparatif pasca dilakukan *feature selection*, serta mengukur dampaknya terhadap efisiensi komputasi dan performa klasifikasi. Pendekatan menggunakan *lightweight models* juga mampu mengungguli teknik-teknik terkini dalam hal akurasi meskipun hanya menggunakan data *labeled* yang lebih sedikit [8]. Dengan demikian, penelitian ini tidak hanya mengadopsi pendekatan serupa, tetapi juga mengoptimalkannya untuk kasus penggunaan yang lebih ringan dan adaptif. Oleh karena itu, penelitian ini diharapkan dapat memberikan kontribusi dalam meningkatkan keamanan siber serta mengurangi risiko serangan *phishing* di masa depan.

1.2 Rumusan Masalah

Berdasarkan latar belakang yang sudah di jelaskan, maka rumusan masalah dalam penelitian ini adalah:

1. Bagaimana penerapan metode *ensemble learning* dan *lightweight models* dapat meningkatkan akurasi deteksi *phishing URL* menggunakan machine learning?
2. Bagaimana penerapan teknik *feature subset ranking* untuk optimalisasi dalam mendeteksi URL *phishing* berdasarkan penggunaan matrik evaluasi?

1.3 Batasan Masalah

Agar penelitian ini tetap fokus dan memiliki ruang lingkup yang jelas, terdapat beberapa batasan masalah yang ditetapkan, yaitu :

1. Penelitian ini hanya berfokus pada deteksi phishing berbasis *URL* dan *HTML*, sehingga tidak mencakup metode phishing lain seperti *voice phishing* atau *spear phishing*.
2. Data yang digunakan dalam penelitian ini bersumber dari dataset publik "*Web page phishing detection* tahun 2021" yang tersedia dan tidak melibatkan data *real-time* dari pengguna individu.
<https://www.kaggle.com/datasets/shashwatwork/web-page-phishing-detection-dataset/data>
3. Teknik yang digunakan dalam penelitian ini adalah pendekatan *ensemble learning* dan *boosting-based feature selection*, tanpa membahas metode keamanan lainnya seperti enkripsi data atau *firewall*.
4. Penelitian ini hanya menggunakan metode *Random Forest Feature Importance* dalam proses seleksi fitur. Pemilihan ini didasarkan pada kemampuan yang efisien dalam menangani data berdimensi tinggi, memberikan interpretasi langsung terhadap pentingnya fitur, dan terintegrasi secara langsung dengan algoritma *ensemble* yang juga digunakan dalam penelitian ini. Metode lain seperti *Recursive Feature Elimination* (RFE) dan *Chi-Square Test* tidak digunakan karena memiliki keterbatasan dalam skalabilitas atau kesesuaian dengan tipe data yang digunakan dalam penelitian ini.
5. Implementasi model yang dikembangkan hanya difokuskan pada sistem berbasis web dan tidak mencakup aplikasi mobile atau perangkat *IoT*.

1.4 Tujuan Penelitian

Penelitian ini bertujuan untuk:

1. Mengembangkan model deteksi *phishing* berbasis *machine learning* dengan tingkat akurasi yang lebih tinggi dibandingkan metode konvensional.
2. Mengidentifikasi fitur-fitur yang paling berpengaruh dalam klasifikasi *phishing* agar proses deteksi dapat dilakukan dengan lebih efisien.

3. Mengevaluasi performa model deteksi phishing menggunakan metrik seperti *akurasi*, *presisi*, *recall*, dan *F1-score* untuk memastikan efektivitasnya

1.5 Manfaat Penelitian

Penelitian ini diharapkan dapat memberikan beberapa manfaat sebagai berikut:

1. Manfaat Akademik

- o Memberikan kontribusi dalam pengembangan teknik deteksi phishing berbasis *ensemble learning* dan *machine learning*.
- o Menambah wawasan dalam penelitian keamanan siber, khususnya dalam analisis fitur pada *phishing detection*.

2. Manfaat Praktis

- o Menghasilkan model deteksi *phishing* yang lebih akurat dan dapat diimplementasikan dalam sistem keamanan web.
- o Memberikan solusi bagi perusahaan dan individu dalam mengidentifikasi ancaman *phishing* dengan lebih cepat dan efektif.

3. Manfaat Sosial

- o Meningkatkan kesadaran masyarakat tentang ancaman *phishing* dan cara menghindarinya.
- o Mengurangi resiko pencurian data dan kejahatan siber yang dapat berdampak buruk pada individu maupun organisasi

1.6 Sistematika Penulisan

Penelitian ini disusun dengan sistematika sebagai berikut ini:

1. BAB I : Pendahuluan

Membahas latar belakang, rumusan masalah, batasan masalah, tujuan penelitian, manfaat penelitian, dan sistematika penulisan.

2. BAB II : Tinjauan Pustaka

Mengulas teori-teori yang berkaitan dengan *phishing*, metode deteksi berbasis *machine learning*, serta hasil penelitian terdahulu yang relevan.

3. BAB III : Metode Penelitian

Menjelaskan metode penelitian yang digunakan, termasuk dataset yang dipakai, teknik pemilihan fitur, algoritma *machine learning* yang digunakan, serta matrik evaluasi yang diterapkan.

4. BAB IV : Hasil dan Pembahasan

Menyajikan hasil eksperimen, analisis kinerja model, serta perbandingan dengan metode lainnya.

5. BAB V : Penutup

Berisi kesimpulan dari penelitian yang telah dilakukan serta saran untuk pengembangan lebih lanjut.

