

Deteksi Serangan *Web Defacement* pada Infrastruktur Kritis Menggunakan *Machine Learning*

(The Detection of Web Defacement Attacks on Critical Infrastructure Using Machine Learning)

Victor Eric Pattiradjawane^{1*}, Doms Upuy²

^{1,2}Program Studi Ilmu Komputer, Fakultas Sains dan Teknologi, Universitas Pattimura
Jl. Ir. M. Putuhena, Ambon, 97233, Indonesia

*Corresponding author's e-mail: * victor.pattiradjawane@lecturer.unpatti.ac.id

Manuscript submitted:
28th February 2025

Manuscript revision:
27th March 2025

Accepted for publication:
29th May 2025

Abstract

The number of threats to the security of websites and web servers, which include things like Web Defacement, is increasing every year. This is a major concern in today's cybersecurity world. It includes websites that are part of critical infrastructure, like government, health, and energy systems. This research study looks at how machine learning models can automatically detect Web Defacement attacks. The main goal is to make sure these models are very accurate. We developed a supervised learning-based classification model using Web Defacement datasets from public archives and simulated mock sites. This research study looks at how well three types of classification models—Random Forest, Support Vector Machine, and Naive Bayes—perform at identifying defaced websites. The results of the experiment show that Random Forest is the best option, with up to 96% accuracy. This research shows that the machine learning approach could be very important in developing a system that can detect cyberattacks early on. This system would protect important infrastructure in the country.

Keywords: Anomaly Detection; Critical Infrastructure; Cybersecurity; Machine Learning; Web Defacement.



This article is an open access article distributed under the terms and conditions of the [Creative Commons Attribution-ShareAlike 4.0 International License](https://creativecommons.org/licenses/by-sa/4.0/).

How to cite this article:

V. Pattiradjawane and D. Upuy, "Deteksi Serangan Web Defacement pada Infrastruktur Kritis Menggunakan Machine Learning", *algorithm*, vol. 1, no. 1, pp. 37-42, May 2025.

Copyright © 2025 Author(s)

Journal homepage: <https://ojs3.unpatti.ac.id/index.php/algorithm>

Research Article [Open Access](#)

1. PENDAHULUAN

Dalam beberapa tahun terakhir, insiden keamanan siber mengalami peningkatan signifikan, baik dari segi frekuensi, kompleksitas, maupun target serangan. Salah satu bentuk serangan yang umum namun berdampak besar adalah *Web Defacement* [1]. Serangan ini mengubah tampilan atau isi halaman web dengan konten yang tidak sah, sering kali bersifat provokatif atau merusak citra organisasi. Ketika situs yang diserang termasuk dalam infrastruktur kritis seperti pemerintahan, kesehatan, dan energi, konsekuensinya mencakup gangguan layanan vital, kerugian ekonomi, serta penurunan kepercayaan publik.

Deteksi serangan *Web Defacement* secara manual atau berbasis signature sering kali tidak mampu mengimbangi kecepatan dan variasi serangan. Untuk itu, pendekatan *machine learning* (ML) muncul sebagai solusi yang adaptif dan cerdas. Model ML dapat dilatih untuk mengenali anomali pada konten halaman web berdasarkan pola-pola historis, sehingga mampu mendeteksi serangan yang sebelumnya belum dikenal (*zero-day attack*).

Penelitian ini bertujuan menerapkan dan mengevaluasi performa beberapa model atau algoritma ML dalam mendeteksi halaman web yang telah di-deface. Hasil dari penelitian ini diharapkan dapat mendukung pengembangan sistem deteksi dini berbasis ML yang efisien dan akurat untuk perlindungan infrastruktur kritis di Indonesia.

Berbagai penelitian telah dilakukan terkait deteksi serangan *Web Defacement*. Metode berbasis signature tidak mampu mendeteksi serangan baru secara efektif [2]. Pendekatan visual berbasis CNN telah digunakan untuk mendeteksi perubahan tampilan halaman web [3], sedangkan perbandingan algoritma klasifikasi menunjukkan efektivitas yang beragam dalam mendeteksi web berbahaya [4].

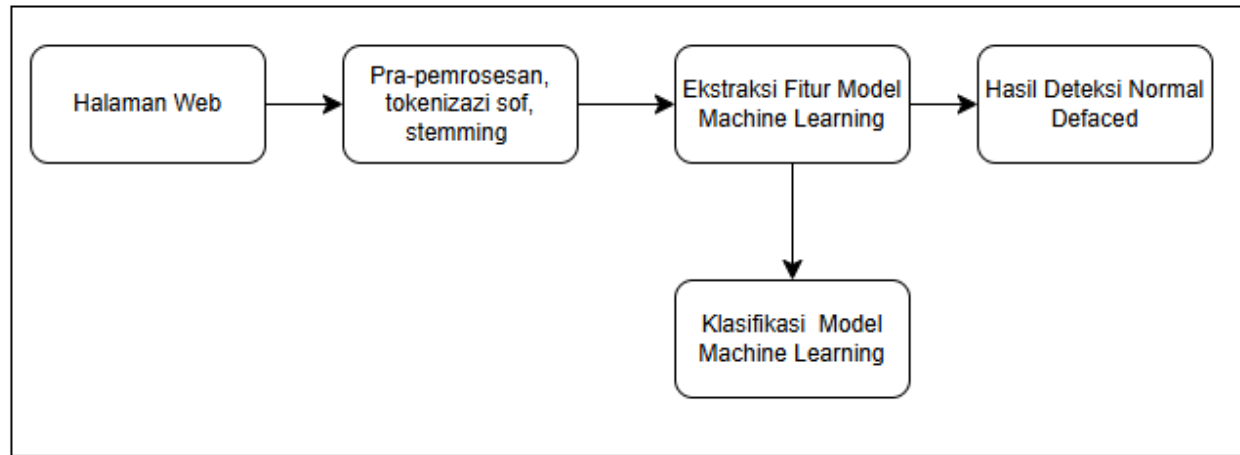
Metode deteksi anomali relevan untuk diterapkan dalam konteks *defacement* [5], [6]. Ekstraksi fitur HTML juga dianggap penting dalam membedakan konten sah dan konten yang telah diubah [7]. Pendekatan *hybrid* yang menggabungkan fitur log dan konten juga telah diusulkan untuk mendeteksi serangan web [8]. Dengan mengombinasikan keunggulan berbagai pendekatan ini, sistem deteksi diharapkan dapat mencapai tingkat akurasi yang lebih tinggi serta mengurangi *false positive*.

Pentingnya kualitas data dalam pelatihan model ML untuk keamanan siber telah banyak disorot [9]. Model *ensemble* seperti *Random Forest* dan *Gradient Boosting* menunjukkan efektivitas dalam mendeteksi intrusi [10]. Pendekatan berbasis AI dikembangkan untuk menghasilkan *threat intelligence* secara otomatis [11]. Sejalan dengan itu, penelitian ini menekankan kontribusi praktis dalam menghasilkan model deteksi *Web Defacement* yang tidak hanya efektif, tetapi juga mampu beradaptasi dengan dinamika serangan siber yang terus berkembang.

Dengan memperhatikan perkembangan tersebut, penelitian ini diharapkan mampu memberikan kontribusi nyata dalam bidang keamanan siber, khususnya pada aspek deteksi dini serangan *Web Defacement*. Pendekatan berbasis *machine learning* dengan kombinasi teknik ekstraksi fitur yang tepat diharapkan tidak hanya meningkatkan keakuratan sistem deteksi, tetapi juga mampu beradaptasi dengan pola serangan yang dinamis. Lebih jauh, penelitian ini juga berpotensi menjadi dasar bagi pengembangan sistem keamanan yang bersifat *real-time*, skalabel, dan dapat diintegrasikan ke berbagai jenis infrastruktur digital yang kritis.

2. METODE PENELITIAN

Penelitian ini menerapkan pendekatan berbasis supervised *machine learning* untuk mendeteksi serangan *Web Defacement* pada infrastruktur kritis. Metodologi terdiri dari beberapa tahapan sistematis mulai dari akuisisi data, pra-pemrosesan, ekstraksi fitur, pelatihan model, hingga evaluasi kinerja. Diagram arsitektur sistem ditunjukkan pada Gambar 1.



Gambar 1. Arsitektur Sistem Deteksi Serangan *Web Defacement*

2.1. Akuisisi dan Klasifikasi Dataset

Dataset yang digunakan terdiri dari dua kategori utama yaitu halaman web normal dan halaman web yang telah mengalami *defacement*. Sumber data diperoleh dari arsip zone-h berupa situs-situs publik yang terdampak *defacement* (<https://www.zone-h.org/archive>), dan simulasi *defacement* lokal berdasarkan skenario umum seperti penggantian teks atau penyisipan pesan politik. Setiap halaman web disimpan dalam bentuk HTML atau plain text dan dilabeli sesuai kategorinya. Total data yang digunakan adalah sebanyak 100 sampel, dengan distribusi seimbang antara dua kelas.

2.2. Pra-Pemrosesan Data

Langkah ini bertujuan untuk membersihkan dan normalisasi data sebelum digunakan dalam model pembelajaran mesin. Tahapan pra-pemrosesan meliputi: 1) Tokenisasi: Memecah dokumen menjadi unit kata (tokens); 2) Stopword Removal: Menghapus kata-kata umum yang tidak memberikan nilai informatif (contoh: “yang”, “dan”, “dengan”) ; 3) Stemming: Mengubah kata ke bentuk dasarnya untuk menyatukan variasi morfologis. 4) Lowercasing: Menstandarkan semua kata menjadi huruf kecil. Setelah proses ini, setiap dokumen diwakili oleh kumpulan kata yang siap untuk diekstraksi fiturnya.

2.3. Ekstraksi Fitur

Fitur teks diekstraksi menggunakan teknik TF-IDF (*Term Frequency-Inverse Document Frequency*). TF-IDF mengukur pentingnya suatu kata dalam dokumen relatif terhadap kumpulan dokumen. TF mencerminkan frekuensi kata dalam dokumen, sedangkan IDF memberikan bobot lebih tinggi pada kata yang jarang muncul di seluruh dokumen. Output dari tahap ini adalah vektor fitur numerik berdimensi tinggi yang dapat digunakan oleh algoritma klasifikasi.

2.4. Pelatihan Model *Machine learning*

Tiga algoritma *machine learning* diuji dan dibandingkan dalam penelitian ini, yaitu: *Random Forest* (RF): algoritma ensemble berbasis pohon keputusan yang membentuk beberapa pohon acak dan menggabungkan hasilnya untuk klasifikasi akhir; *Support Vector Machine* (SVM): model berbasis margin maksimal yang mencoba menemukan hyperplane terbaik untuk memisahkan dua kelas; dan Naive Bayes (NB): pendekatan probabilistik sederhana berdasarkan Teorema Bayes dengan asumsi independensi antar fitur. Setiap model dilatih dengan data TF-IDF dan dikonfigurasi menggunakan teknik validasi silang *10-fold* untuk menghindari overfitting dan mendapatkan hasil generalisasi yang akurat.

2.5. Evaluasi Kinerja Model

Model yang telah dilatih dievaluasi menggunakan metrik evaluasi sebagai berikut: 1) Akurasi: Persentase klasifikasi yang benar terhadap total data. 2) *Precision*: Kemampuan model dalam mengklasifikasikan data positif secara benar. 3) *Recall* (Sensitivity): Kemampuan model untuk menangkap seluruh data positif. 4) *F1-score*: Harmonik rata-rata antara *precision* dan *recall*. Evaluasi dilakukan menggunakan *confusion matrix* untuk memahami distribusi prediksi model secara lebih mendalam.

3. HASIL DAN PEMBAHASAN

Hasil eksperimen menunjukkan bahwa model **Random Forest** unggul dengan akurasi sebesar 96%, *precision* 95,8%, *recall* 97%, dan nilai *F1-score* 96,4%. Nilai ROC-AUC yang tinggi yaitu 0,97 juga memperkuat performa model ini dalam mengklasifikasikan serangan *Web Defacement*. Sementara itu, model **Support Vector Machine (SVM)** memperoleh akurasi 91,2%, *precision* 90,5%, *recall* 92,4%, dan *F1-score* 91,4%. Hasil ini menunjukkan bahwa meskipun kinerja SVM cukup baik, model ini masih berada di bawah *Random Forest*. Model Naive Bayes menghasilkan akurasi 86,7%, *precision* 86,1%, *recall* 87,2%, dan *F1-score* 86,6%, yang relatif lebih rendah dibanding dua algoritma lainnya (Tabel 1, Gambar 2).

Perbandingan performa ini memperlihatkan bahwa *Random Forest* memiliki stabilitas yang lebih baik terhadap variasi struktur konten HTML. Hal ini disebabkan oleh sifat ensemble *Random Forest* yang menggabungkan banyak pohon keputusan, sehingga lebih robust dalam menangani data yang kompleks dan beragam. Sementara itu, SVM memiliki keunggulan dalam mendeteksi pola dengan margin yang jelas, namun sensitif terhadap parameter kernel dan membutuhkan waktu pelatihan yang relatif lebih lama. Naive Bayes, meskipun sederhana dan efisien secara komputasi, menunjukkan keterbatasan karena asumsi independensi antar fitur sering tidak terpenuhi dalam kasus deteksi *Web Defacement* yang kompleks.

Tabel 1. Hasil Evaluasi Model Berdasarkan Empat Metrik Utama: Akurasi, Precision, Recall, Dan F1-Score

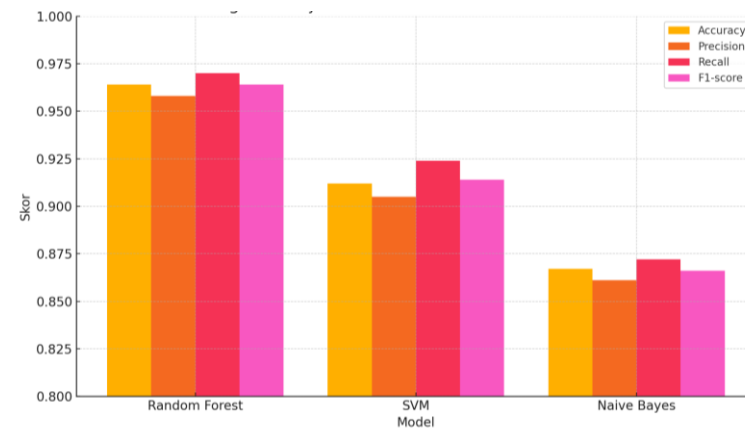
Algoritma	Akurasi	Precision	Recall	F1-score
<i>Random Forest</i>	0.964	0.958	0.970	0.964
Support Vector Machine	0.912	0.905	0.924	0.914
Naive Bayes	0.867	0.861	0.872	0.866

Penelitian ini menunjukkan bahwa pendekatan *machine learning*, khususnya *Random Forest*, mampu mengidentifikasi serangan *Web Defacement* dengan tingkat akurasi yang sangat tinggi. Kombinasi dari fitur TF-IDF dan arsitektur memberikan keunggulan dalam mendeteksi perubahan anomali konten web. Selain itu, pendekatan ini bersifat generik dan dapat diadaptasi untuk berbagai bahasa atau struktur HTML, sehingga menjadikannya fleksibel untuk diintegrasikan ke dalam sistem keamanan *real-time*.

Visualisasi pada Gambar 2 juga memperkuat temuan ini, di mana batang hasil evaluasi untuk *Random Forest* konsisten lebih tinggi dibanding SVM dan Naive Bayes pada semua metrik evaluasi. Hal ini menegaskan bahwa pendekatan berbasis ensemble lebih mampu menggeneralisasi data, sehingga menghasilkan kinerja yang lebih stabil pada data uji.

Selain itu, hasil ini juga menunjukkan bahwa pemilihan algoritma yang tepat sangat berpengaruh pada efektivitas sistem deteksi. *Random Forest* dapat menjadi pilihan utama untuk implementasi sistem keamanan siber berbasis *machine learning* karena tidak hanya memberikan akurasi tinggi, tetapi juga mampu beradaptasi dengan variasi data yang dinamis. Sementara itu,

SVM tetap relevan untuk kasus dengan jumlah data yang lebih kecil dan pola klasifikasi yang lebih sederhana. Naive Bayes, meskipun memiliki performa paling rendah dalam penelitian ini, tetap memiliki nilai praktis karena efisiensinya, khususnya pada sistem dengan keterbatasan komputasi.



Gambar 2. Perbandingan Visual Dari Hasil Evaluasi

Secara keseluruhan, penelitian ini membuktikan bahwa pendekatan *machine learning*, khususnya Random Forest, mampu mengidentifikasi serangan *Web Defacement* dengan tingkat akurasi yang sangat tinggi. Kombinasi dari fitur TF-IDF dan struktur arsitektur algoritma memberikan keunggulan dalam mendeteksi perubahan anomali konten web. Selain itu, pendekatan ini bersifat generik dan dapat diadaptasi untuk berbagai bahasa atau struktur HTML, sehingga menjadikannya fleksibel untuk diintegrasikan ke dalam sistem keamanan *real-time*.

4. KESIMPULAN DAN SARAN

Penelitian ini membuktikan bahwa penerapan *machine learning*, khususnya algoritma *Random Forest*, efektif dalam mendeteksi serangan *Web Defacement*. Model ini memiliki potensi untuk diimplementasikan dalam sistem deteksi dini pada infrastruktur kritis nasional. Uji coba pada situs dummy yang meniru halaman pemerintahan menunjukkan model algoritma *Random Forest* (RF) mampu mendeteksi *defacement* dengan latensi rendah. Namun, tantangan muncul pada halaman dengan konten dinamis atau berita yang mengandung kata-kata sensitif yang menyerupai konten *defacement*.

Saran untuk penelitian selanjutnya adalah menggabungkan analisis konten visual dan metadata halaman serta menerapkan pendekatan *real-time* berbasis *edge computing* dengan jumlah data yang lebih besar. Perlu juga dilakukan pengujian pada sistem nyata untuk mengevaluasi performa model dalam kondisi operasional sesungguhnya.

REFERENSI

- [1] <https://www.bssn.go.id/wp-content/uploads/2024/03/Lanskap-Keamanan-Siber-Indonesia-2023.pdf>
- [2] M. Hassen et al., 'Web Defacement Detection Using Signature-Based Methods', International Journal of Cyber Security, 2019.
- [3] J. Kim et al., 'Image-based Deep Learning for Defacement Detection', IEEE CyberSec Conference, 2020.
- [4] R. Ahmad et al., 'Comparative Study of Machine learning Algorithms', Journal of Information Security, 2021.

- [5] V. Chandola et al., 'Anomaly Detection: A Survey', ACM Computing Surveys, vol. 41, no. 3, 2009.
- [6] A. Patcha and J. Park, 'An Overview of Anomaly Detection Techniques', Computer Networks, vol. 51, no. 12, pp. 3448–3470, 2007.
- [7] G. Vrbančič et al., 'HTML Content Feature Extraction for Web Security', Information Sciences, vol. 484, 2019.
- [8] Y. Zhang et al., 'Hybrid ML Approaches for Web Attack Detection', Computers & Security, vol. 74, 2018.
- [9] R. Sommer and V. Paxson, 'Outside the Closed World: On Using ML for Security', IEEE S&P, 2010.
- [10] R. Vinayakumar et al., 'Deep Learning Approaches for Cybersecurity', Future Generation Computer Systems, 2019.
- [11] M. Alazab et al., 'AI-Driven Cyber Threat Intelligence', IEEE Access, 2020.