

Madani: Jurnal Ilmiah Multidisiplin
Volume 3, Nomor 6, July 2025, P. 291-294
E-ISSN: 2986-6340
Licenced by CC BY-SA 4.0
DOI: <https://doi.org/10.5281/zenodo.15858747>

Infrastruktur Jaringan Self-healing dengan AIOps: Remediasi Otomatis Kerentanan Keamanan

Muhammad Arsy Athallah¹, Irfan Qobus Salim², Hafiz Akmaldi Santosa³, Naufan Zaki Luqmanulhakim⁴, Muhammad Rifqi Oktaviansyah⁵

Fakultas Teknologi Elektro dan Informatika Cerdas, Institut Teknologi Sepuluh Nopember, Surabaya, 60111,
Email: 5027221048@mhs.its.ac.id, 5027221058@mhs.its.ac.id,
5027221061@mhs.its.ac.id, 5027221065@mhs.its.ac.id, 5027221067@mhs.its.ac.id

Abstract

Peningkatan volume dan kompleksitas serangan siber terhadap institusi pendidikan menuntut pendekatan yang lebih proaktif dan otomatis dalam manajemen keamanan. Ketergantungan pada intervensi manual oleh tim *Security Operations Center (SOC)* sering kali menghasilkan waktu respons yang lambat dan tidak mampu mengimbangi kecepatan ancaman modern. Penelitian ini menyajikan implementasi dan evaluasi sebuah prototipe infrastruktur jaringan self-healing yang memanfaatkan *Artificial Intelligence for IT Operations (AIOps)* dengan pendekatan *Deep Reinforcement Learning (DRL)*. Kami mengembangkan sebuah lingkungan simulasi yang mereplikasi jaringan kampus dengan berbagai tingkat kritisitas perangkat. Sebuah agen cerdas berbasis *Deep Q-Network (DQN)* dilatih untuk secara mandiri memilih tindakan remediasi yang optimal—seperti mem-patch server atau mengisolasi perangkat—berdasarkan jenis kerentanan dan konteks perangkat yang terancam. Hasil simulasi selama 150 hari menunjukkan bahwa agen berhasil belajar dari pengalaman, dengan tingkat kesehatan jaringan yang pulih dari fase pembelajaran awal dan stabil di atas 95% pada fase akhir. Agen juga berhasil membangun kebijakan remediasi yang 100% akurat untuk semua skenario ancaman yang diuji. Penelitian ini mendemonstrasikan bahwa pendekatan DQN adalah metode yang sangat menjanjikan untuk menciptakan sistem keamanan siber yang otonom, adaptif, dan mampu memulihkan diri, sejalan dengan tujuan untuk mengurangi beban kerja tim SOC dan meningkatkan postur keamanan secara signifikan.

Keywords: *Self-healing Network, AIOps, Keamanan Siber, Deep Reinforcement Learning, DQN, Otomatisasi SOC, Manajemen Kerentanan.*

Article Info

Received date: 20 June 2025

Revised date: 30 June 2025

Accepted date: 10 July 2025

PENDAHULUAN

Arus digitalisasi di institusi pendidikan tinggi telah memperluas permukaan serangan siber secara eksponensial. Lingkungan akademik, dengan sifatnya yang terbuka dan heterogen, menjadi target utama bagi berbagai jenis ancaman, mulai dari ransomware hingga pencurian data penelitian [1]. Studi terbaru menunjukkan korelasi positif antara skala sebuah institusi—seperti jumlah mahasiswa dan belanja operasional—dengan tingkat eksposur kerentanannya, menjadikannya target yang menarik bagi aktor jahat [2].

Secara tradisional, pertahanan siber diorganisir melalui Security Operations Center (SOC), di mana para analis secara manual memonitor, mendeteksi, dan merespons insiden. Namun, pendekatan ini memiliki keterbatasan fundamental. Volume data telemetri yang masif, ditambah dengan kecepatan serangan modern, seringkali melampaui kapasitas kognitif dan operasional tim manusia [3]. Hal ini mengakibatkan alert fatigue, waktu respons yang lambat, dan kerentanan kritis yang tidak tertangani selama sehari-hari.

Untuk menjawab tantangan ini, paradigma Artificial Intelligence for IT Operations (AIOps)

muncul sebagai sebuah evolusi. AIOps mengintegrasikan kecerdasan buatan, analitik data, dan machine learning untuk mengotomatisasi dan menyempurnakan operasi TI, termasuk keamanan [4]. Dalam konteks keamanan jaringan, AIOps memungkinkan terwujudnya konsep self-healing network—sebuah sistem yang tidak hanya mendeteksi masalah tetapi juga mampu mendiagnosis dan memperbaikinya secara otonom [5].

Penelitian ini berfokus pada perancangan, implementasi, dan evaluasi sebuah prototipe jaringan self-healing berbasis AIOps yang dirancang khusus untuk konteks lingkungan pendidikan tinggi. Kami mengajukan sebuah model yang menggunakan Deep Reinforcement Learning (DRL), khususnya algoritma Deep Q-Network (DQN), untuk melatih sebuah agen cerdas yang dapat membuat keputusan remediasi kerentanan secara optimal dan mandiri. Tujuan utamanya adalah untuk menciptakan sistem yang dapat secara signifikan mengurangi Mean Time to Remediate (MTTR) dan meningkatkan postur keamanan jaringan secara keseluruhan tanpa intervensi manual.

TINJAUAN PUSTAKA

Konsep AIOps telah menjadi pendorong utama dalam evolusi manajemen infrastruktur TI. Studi oleh Eramo et al. [4] menyoroti bagaimana AIOps, melalui integrasi model-driven engineering dan machine learning, dapat menciptakan otomatisasi end-to-end dalam siklus DevOps. Dalam domain keamanan, pendekatan ini memungkinkan transisi dari deteksi reaktif menjadi remediasi prediktif dan proaktif.

Ide jaringan self-healing berakar pada Autonomic Computing, yang bertujuan menciptakan sistem yang dapat mengelola dirinya sendiri. Dalam konteks jaringan modern, khususnya Software-Defined Networking (SDN), kemampuan self-healing menjadi krusial untuk menjamin ketersediaan dan ketahanan [5]. Namun, penelitian sebelumnya sering kali berfokus pada pemulihan konektivitas, dengan perhatian yang lebih sedikit pada remediasi kerentanan keamanan secara otonom.

Reinforcement Learning (RL) telah muncul sebagai kandidat utama untuk otomatisasi keputusan dalam sistem yang kompleks. Berbeda dengan supervised learning yang membutuhkan data berlabel, agen RL belajar melalui interaksi langsung dengan lingkungannya, menerima "hadiah" untuk tindakan yang baik dan "hukuman" untuk tindakan yang buruk. Deep Reinforcement Learning (DRL), yang menggabungkan RL dengan deep neural networks, mampu menangani ruang status yang kompleks dan

membuat keputusan yang lebih canggih [6], menjadikannya sangat cocok untuk domain keamanan siber yang dinamis.

Meskipun demikian, masih terdapat kesenjangan penelitian dalam penerapan DRL secara praktis untuk remediasi kerentanan otomatis dalam konteks SOC, terutama yang mempertimbangkan karakteristik unik dan profil risiko institusi pendidikan [2]. Penelitian ini bertujuan untuk mengisi kesenjangan tersebut dengan menyajikan sebuah model DRL yang konkret dan terukur.

METODE

Kami merancang sebuah lingkungan simulasi berbasis Python untuk membangun dan mengevaluasi agen AIOps kami. Arsitektur ini mengikuti kerangka logis MAPE-K (Monitor, Analyze, Plan, Execute, Knowledge).

A. Arsitektur Simulasi

1. Lingkungan Jaringan (Network Environment): Lingkungan ini terdiri dari 20 server virtual. Setiap server memiliki dua atribut kontekstual utama:
 - a. Tingkat Kritikalitas (importance): Secara acak ditetapkan sebagai Low, High, atau Critical untuk mensimulasikan beragamnya prioritas aset dalam jaringan nyata.
 - b. Kerentanan (vulnerabilities): Sebuah himpunan yang berisi ID CVE (Common Vulnerabilities and Exposures) yang sedang aktif.
2. Agen AIOps (DecisionEngineDQN): Ini adalah inti dari sistem, yang bertanggung jawab untuk membuat keputusan. Otaknya adalah sebuah Deep Q-Network (DQN).
 - a. Input (State): Agen menerima informasi tentang jenis kerentanan dan kritikalitas server. Informasi tekstual ini diubah menjadi vektor numerik melalui one-hot encoding agar dapat diproses oleh neural network. Ini memungkinkan agen untuk memahami konteks, misalnya, membedakan antara "kerentanan database pada server kritis" versus "kerentanan web pada server prioritas rendah".

Program 1. Fungsi untuk mengubah state menjadi vektor

```
def encode_state(self, cve_id, importance_level):
    # Membuat vektor kosong dengan ukuran yang
    # tepat
    state = np.zeros(self.state_size)
```

```

# Mengaktifkan bit untuk jenis CVE
if cve_id in self.vulnerability_map:
    state[self.vulnerability_map[cve_id]] = 1

# Mengaktifkan bit untuk tingkat kritisitas
if importance_level in self.importance_map:
    state[len(self.vulnerability_map) +
self.importance_map[importance_level]] = 1

return state.reshape(1, -1)

```

3. Output (Action): Berdasarkan state yang diberikan, agen memilih satu dari lima kemungkinan tindakan: `patch_server`, `update_firewall_rule`, `isolate_device`, `reconfigure_service`, atau `rotate_credentials`.
4. Mekanisme Pelatihan

Agan dilatih menggunakan algoritma DQN dengan experience replay. Setiap tindakan yang diambil (pengalaman) disimpan dalam sebuah memori.

 - a. Reward Clipping: Untuk menstabilkan pelatihan, kami menormalkan hadiah. Aksi yang benar memberikan reward +1.0, dan aksi yang salah memberikan hukuman -1.0. Teknik ini terbukti krusial untuk mencegah training loss yang tidak stabil.
 - b. Epsilon-Greedy: Untuk menyeimbangkan eksplorasi dan eksploitasi, agen menggunakan kebijakan ϵ -greedy. Nilai ϵ (epsilon) dimulai dari 1.0 (100% eksplorasi) dan menurun secara bertahap hingga 0.01 seiring agen menjadi lebih berpengalaman.
 - c. Proses Pelatihan Harian: Setiap "hari", agen tidak hanya menangani ancaman baru, tetapi secara proaktif meninjau semua perangkat yang masih memiliki kerentanan. Pendekatan ini secara drastis meningkatkan jumlah pengalaman belajar dan mereplikasi alur kerja SOC yang realistis.

HASIL DAN PEMBAHASAN

Simulasi dijalankan selama 150 hari virtual untuk memberikan waktu yang cukup bagi agen untuk belajar dari nol dan mencapai konvergensi.

A. Analisis Kinerja dan Kurva Pembelajaran

Kinerja agen dilacak melalui tiga metrik utama: kesehatan jaringan, total hadiah akumulatif, dan

training loss. Grafik pada Gambar 1 merangkum evolusi performa agen.



Gambar 1. Grafik Kinerja Agen AIOps Selama 150 Hari Simulasi. (Atas) Kurva Kesehatan Jaringan dan Penurunan Epsilon. (Bawah) Kurva Reward Kumulatif dan Training Loss.

Kurva Kesehatan Jaringan (Biru): Grafik ini dengan jelas menunjukkan kemampuan self-healing. Pada fase awal (sekitar Hari 1-40), di mana Epsilon masih tinggi dan agen sering melakukan eksplorasi acak, kesehatan jaringan berfluktuasi dan sempat menurun. Namun, seiring Epsilon menurun dan agen mulai memanfaatkan pengetahuannya (fase eksploitasi), ia menjadi semakin efisien dalam melakukan remediasi. Hasilnya, kurva kesehatan menunjukkan tren naik yang kuat dan berhasil mempertahankan tingkat kesehatan di atas 95%—bahkan seringkali 100%—pada paruh kedua simulasi.

Kurva Reward (Hijau) dan Loss (Merah): Kurva reward yang berbentuk "lembah" adalah visualisasi dari "biaya pendidikan" agen. Penurunan tajam di awal mencerminkan banyaknya kesalahan selama fase eksplorasi. Titik di mana kurva mulai melandai dan kemudian naik (sekitar Hari 80-90) menandakan bahwa agen telah cukup belajar untuk membuat lebih banyak keputusan benar daripada salah. Hal ini bertepatan dengan kurva training loss yang mulai stabil di nilai yang sangat rendah, menunjukkan bahwa prediksi model DQN telah selaras dengan hasil nyata.

B. Akurasi Model dan Kebijakan yang Dipelajari

Pada akhir simulasi, kami menguji "otak" agen yang telah terlatih dengan memberinya setiap kemungkinan skenario ancaman. Hasilnya menunjukkan akurasi 100%. Agen berhasil mempelajari kebijakan (strategi) yang benar untuk setiap jenis kerentanan tanpa memandang tingkat kritisitas server.

Tabel 1. Contoh Prediksi Model DQN yang Telah Terlatih.

Input (Kerentanan @ Kritikalitas Server)	Prediksi Aksi Terbaik (Output Agen)	Aksi yang Benar (Ground Truth)
CVE-2023-1001 @ Critical Server	patch_server	patch_server
CVE-2023-3003 @ High Server	isolate_device	isolate_device
CVE-2024-5005 @ Low Server	rotate_credentials	rotate_credentials

Hasil pada Tabel 1 membuktikan bahwa agen tidak hanya menghafal, tetapi mampu menggeneralisasi pengetahuan untuk membentuk sebuah kebijakan remediasi yang andal dan akurat, yang merupakan tujuan inti dari penelitian ini.

SIMPULAN

Penelitian ini berhasil mendemonstrasikan, melalui implementasi simulasi yang canggih, bahwa pendekatan self-healing berbasis AIOps yang ditenagai oleh Deep Reinforcement Learning merupakan solusi yang sangat menjanjikan untuk tantangan keamanan siber modern. Prototipe agen DQN kami terbukti mampu belajar secara mandiri dari lingkungan yang dinamis untuk mencapai tingkat akurasi yang tinggi dalam memilih tindakan remediasi yang optimal.

Kemampuan sistem untuk pulih dan mempertahankan kesehatan jaringan di level yang tinggi setelah melewati fase pembelajaran awal menggarisbawahi potensi besar dari otomatisasi cerdas ini untuk meningkatkan ketahanan siber dan mengurangi beban operasional tim SOC. Hasil ini membuka jalan untuk penelitian lebih lanjut, termasuk penerapan arsitektur DRL yang lebih canggih dan validasi menggunakan data telemetri dunia nyata.

KETERSEDIAAN KODE

Untuk mendukung prinsip transparansi dan reproduktibilitas dalam penelitian ilmiah, seluruh kode sumber yang digunakan untuk simulasi dalam penelitian ini tersedia untuk umum. Ini termasuk

implementasi lingkungan jaringan, agen DQN, dan skrip untuk menjalankan eksperimen serta menghasilkan grafik. Kode dapat diakses melalui repositori GitHub kami di tautan berikut: <https://github.com/NaufanZaki/AIOps-Platform.git>

REFERENCE

- [1] L. Suarez, D. Alshubrumi, T. O'Connor, and S. Sudhakaran, "Unsafe at any Bandwidth: Towards understanding risk factors for ransomware in higher education," *Procedia Computer Science*, vol. 238, pp. 815–820, 2024.
- [2] J. Forsberg and T. Frantti, "Technical performance metrics of a security operations center," *Computers & Security*, vol. 135, p. 103529, 2023.
- [3] M. M. Ahmadian, H. R. Shahriari, and S. M. Ghaffarian, "Autonomous security incident response in software-defined networks," *IEEE Transactions on Network and Service Management*, vol. 18, no. 2, pp. 1439–1454, 2021.
- [4] R. Eramo, B. Said, M. Oriol, H. Bruneliere, and S. Morales, "An architecture for model-based and intelligent automation in DevOps," *Journal of Systems and Software*, vol. 217, p. 112180, 2024.
- [5] L. Ochoa-Aday, C. Cervelló-Pastor, and A. Fernández-Fernández, "Self-healing and SDN: Bridging the gap," *Digital Communications and Networks*, vol. 6, no. 3, pp. 354–368, 2020.
- [6] T. Nguyen-An, H. T. X. Trang, T. T. Nguyen, and H. M. Bui, "Reinforcement learning for autonomous cybersecurity: From simulated environments to real-world deployments," *Nature Machine Intelligence*, vol. 5, no. 2, pp. 156–168, 2023.