

BAB V PENUTUP

5.1 Kesimpulan

Berdasarkan hasil penelitian yang telah dilakukan melalui studi literatur mendalam, eksperimen teknis, dan survei terhadap 45 praktisi keamanan siber, penelitian ini menarik beberapa kesimpulan utama sebagai jawaban atas rumusan masalah:

a. Bentuk-Bentuk Penyalahgunaan AI dalam Serangan Siber

Penelitian ini berhasil mengidentifikasi dan menganalisis enam bentuk utama penyalahgunaan teknologi AI yang secara signifikan mengubah lanskap ancaman digital:

1. Teknologi *Deepfake*: Manipulasi konten video dan audio untuk penipuan identitas.
2. *Phishing* Bertenaga AI: Pembuatan pesan penipuan yang sangat personal dengan tata bahasa sempurna, meningkatkan tingkat keberhasilan serangan.
3. *Adaptive Malware*: *Malware* yang mampu belajar dan beradaptasi untuk menghindari deteksi sistem keamanan tradisional.
4. AI-Enhanced DDoS: Optimalisasi serangan *Distributed Denial of Service* untuk efektivitas gangguan yang lebih tinggi.
5. Eksploitasi Kerentanan Otomatis: Penggunaan AI untuk memindai dan mengeksploitasi celah keamanan secara mandiri dan cepat.
6. Rekayasa Sosial Berbasis AI: Manipulasi psikologis skala besar yang dipersonalisasi melalui bot atau asisten virtual.

Penyalahgunaan ini meningkatkan skala, kecepatan, dan kompleksitas serangan, sekaligus menurunkan hambatan masuk (*barrier to entry*) bagi pelaku kejahatan siber.

b. Dampak Serangan Berbasis AI Terhadap CIA Triad

Ancaman berbasis AI memberikan dampak multidimensi terhadap pilar keamanan informasi:

1. Kerahasiaan (*Confidentiality*): Dampak paling kritis (skor 4,44/5). Penggunaan AI dalam *phishing* dan *deepfake* secara signifikan

mempermudah pencurian *kredensial* dan identitas.

2. Integritas (*Integrity*): Dampak tinggi (skor 4.26/5). AI memungkinkan manipulasi data secara halus yang sulit dideteksi oleh inspeksi manual, mengancam keaslian informasi publik dan data organisasi.
3. Ketersediaan (*Availability*): Dampak signifikan (skor 4.13/5). Serangan DDoS yang dioptimalkan AI meningkatkan efektivitas serangan dan memperlama durasi *downtime* sistem.
4. Strategi Mitigasi yang Efektif

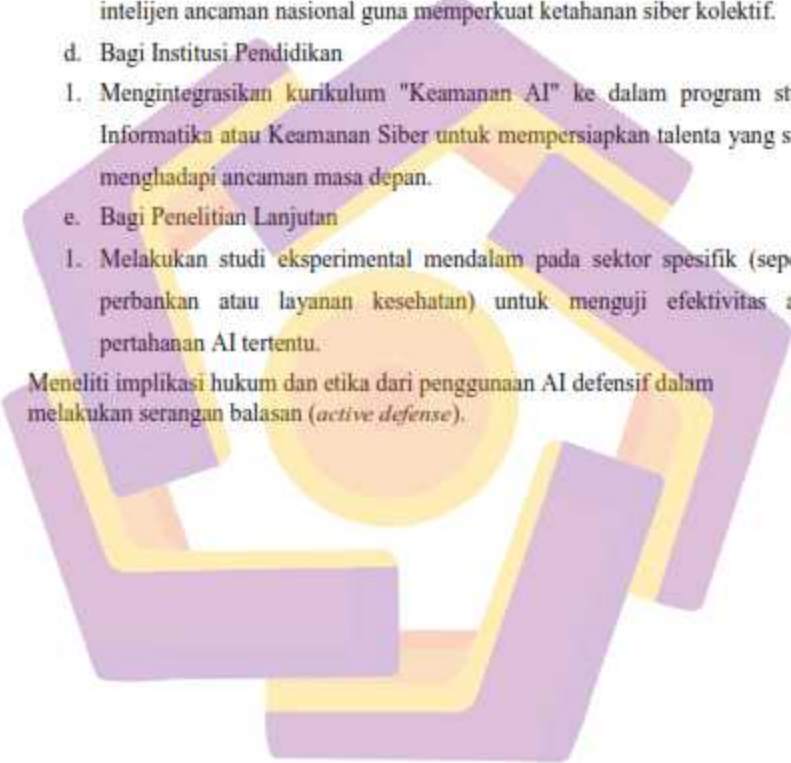
Strategi pertahanan harus dilakukan secara komprehensif melalui empat pilar utama:

1. Strategi Teknis: Implementasi AI defensif, arsitektur *Zero Trust*, sistem deteksi *deepfake*, dan otentikasi multi-faktor (MFA).
2. Strategi Prosedural: Pembaruan kebijakan keamanan yang mencakup ancaman AI dan prosedur respons insiden yang lebih responsif.
3. Strategi Berpusat pada Manusia: Pelatihan kesadaran keamanan (*security awareness*) berkelanjutan sebagai garis pertahanan utama.
4. Strategi Kolaboratif: Berbagi intelijen ancaman (*threat intelligence*) lintas organisasi dan sektor.

5.2 Saran

Berdasarkan temuan penelitian, berikut adalah saran-saran yang direkomendasikan kepada berbagai pemangku kepentingan:

- a. Bagi Organisasi dan Perusahaan
 1. Segera mengadopsi prinsip *Zero Trust Architecture* dan memperkuat sistem otentikasi menggunakan MFA.
 2. Meningkatkan frekuensi pelatihan keamanan bagi seluruh karyawan, tidak hanya tim IT, guna mengantisipasi serangan rekayasa sosial bertenaga AI.
 3. Mengalokasikan anggaran khusus (direkomendasikan minimal 8-10% dari anggaran TI) untuk investasi teknologi keamanan berbasis AI yang mampu melakukan deteksi dini secara otomatis.
- b. Bagi Praktisi Keamanan Siber
 1. Meningkatkan kompetensi teknis terkait AI/ML melalui sertifikasi dan pembelajaran mandiri untuk memahami cara kerja serangan AI.

- 
2. Mengadopsi pola pikir *Red Teaming* (berpikir seperti penyerang) untuk memprediksi skenario serangan masa depan.
 - c. Bagi Pemerintah dan Regulator
 1. Mengembangkan kerangka regulasi yang spesifik mengatur etika dan penggunaan AI untuk mencegah penyalahgunaan teknologi di tingkat nasional.
 2. Mendorong kolaborasi publik-swasta dalam membangun platform berbagi intelijen ancaman nasional guna memperkuat ketahanan siber kolektif.
 - d. Bagi Institusi Pendidikan
 1. Mengintegrasikan kurikulum "Keamanan AI" ke dalam program studi Informatika atau Keamanan Siber untuk mempersiapkan talenta yang siap menghadapi ancaman masa depan.
 - e. Bagi Penelitian Lanjutan
 1. Melakukan studi eksperimental mendalam pada sektor spesifik (seperti perbankan atau layanan kesehatan) untuk menguji efektivitas alat pertahanan AI tertentu.

Meneliti implikasi hukum dan etika dari penggunaan AI defensif dalam melakukan serangan balasan (*active defense*).