

KAJIAN ETIKA DAN MANAJEMEN RISIKO INFORMASI DALAM PEMANFAATAN GENERATIVE AI UNTUK PENGEMBANGAN PERANGKAT LUNAK

Muhammad Gading Herlambang Sopian¹, Ahmad Daffa Herdian², Muhammad Wahyu Ramadhan³, Gloria Ageng Larasati⁴, Mohamad Kendis Prasetya⁵, Sri Mashtufah Anjani⁶,
Ahmad Nursodiq⁷

mgading3101@gmail.com¹, daffaahmad0517@gmail.com², whyrmdbh558@gmail.com³,
glorialarasati04@gmail.com⁴, kubiilprasetya@gmail.com⁵, mashtufahanjani@gmail.com⁶,
dosen02526@unpam.ac.id⁷

Universitas Pamulang

Abstrak

Generative Artificial Intelligence (GenAI) telah memicu transformasi radikal dalam siklus hidup pengembangan perangkat lunak (Software Development Life Cycle/SDLC), menawarkan efisiensi melalui otomatisasi kode dan debugging. Namun, adopsi teknologi ini menghadirkan kompleksitas baru terkait etika dan keamanan informasi, termasuk bias algoritmik, halusinasi kode, pelanggaran hak kekayaan intelektual, dan kebocoran data sensitif. Penelitian ini bertujuan untuk menganalisis implikasi etis dan merumuskan strategi manajemen risiko informasi dalam pemanfaatan GenAI bagi pengembang perangkat lunak. Metode penelitian yang digunakan adalah tinjauan literatur sistematis (Systematic Literature Review) dengan menganalisis 25 artikel ilmiah dan dokumen regulasi periode 2020–2025. Hasil penelitian menunjukkan bahwa meskipun GenAI, seperti model Large Language Models (LLM), meningkatkan produktivitas hingga 45%, penggunaannya tanpa tata kelola yang ketat meningkatkan risiko kerentanan keamanan siber dan masalah legalitas. Studi ini merekomendasikan kerangka kerja mitigasi risiko yang mencakup pendekatan Human-in-the-Loop (HITL), audit algoritma berkala, serta kepatuhan ketat terhadap Undang-Undang Perlindungan Data Pribadi (UU PDP). Kesimpulannya, integrasi GenAI dalam rekayasa perangkat lunak harus menyeimbangkan inovasi teknis dengan tanggung jawab etis dan kepatuhan regulasi.

Kata Kunci: Generative Ai, Etika Ai, Manajemen Risiko Informasi, Pengembangan Perangkat Lunak, Perlindungan Data, Tata Kelola Ai.

Abstract

Generative Artificial Intelligence (GenAI) has triggered a radical transformation in the Software Development Life Cycle (SDLC), offering efficiency through code automation and debugging. However, the adoption of this technology introduces new complexities regarding ethics and information security, including algorithmic bias, code hallucinations, intellectual property infringement, and sensitive data leakage. This Study aims to analyze ethical implications and formulate information risk management strategies in utilizing GenAI for software developers. The research method used is a Systematic Literature Review (SLR) by analyzing 25 scientific articles and regulatory documents from 2020–2025. The results indicate that while GenAI, such as Large Language Models (LLM), increases productivity by up to 45%, its use without strict governance increases cybersecurity vulnerabilities and legal issues. This study recommends a risk mitigation framework that includes a Human-in-the-Loop (HITL) approach, periodic algorithm audits, and strict compliance with the Personal Data Protection Law (UU PDP). In conclusion, integrating GenAI into software engineering must balance technical innovation with ethical responsibility and regulatory compliance.

Keyword: Generative Ai, Computer Ethics, Information Risk Management, Software Engineering, Data Protection Law.

1. PENDAHULUAN

Perkembangan pesat Generative AI (GenAI) telah membuka era baru dalam rekayasa perangkat lunak. Model AI generatif mampu menghasilkan potongan kode, menyarankan fungsi, dan membantu debugging dengan akurasi yang semakin tinggi, didukung oleh kemajuan dalam arsitektur Large Language Models (LLM) [1]. Kehadiran teknologi ini menawarkan efisiensi waktu dan biaya yang signifikan, mendorong adopsi cepat di kalangan pengembang profesional [2].

Namun, adopsi masif GenAI menghadirkan isu etika dan tata kelola yang belum terselesaikan. Salah satu tantangan mendasar adalah isu bias algoritmik, di mana model cenderung mereplikasi atau bahkan memperkuat bias yang inheren dalam data latihnya, yang dapat berujung pada pengembangan sistem yang tidak adil [3]. Selain itu, GenAI juga menimbulkan masalah HKI, karena model sering dilatih pada repositori kode dengan lisensi yang beragam, meningkatkan risiko pelanggaran hak cipta pada output kode yang dihasilkan.

Dari perspektif keamanan informasi, GenAI rentan menghasilkan kode yang mengandung kerentanan (vulnerabilities). Penelitian menunjukkan bahwa kode yang dihasilkan AI, meskipun berfungsi, seringkali kurang memiliki konstruksi pemrograman defensif, yang dapat dieksploitasi oleh pihak tidak bertanggung jawab [4]. Risiko ini diperparah oleh potensi kebocoran data sensitif ketika pengembang memasukkan rahasia dagang atau data pelanggan ke dalam prompt [5].

Di ranah hukum, kepatuhan terhadap regulasi nasional menjadi sangat penting. Undang-Undang Nomor 27 Tahun 2022 tentang Perlindungan Data Pribadi (UU PDP) di Indonesia mewajibkan pengendali data untuk menjamin keamanan pemrosesan data

pribadi, yang secara langsung memengaruhi cara pengembang menggunakan GenAI untuk memproses data pengguna [6].

Oleh karena itu, penelitian ini bertujuan untuk merumuskan kerangka kajian etika dan strategi mitigasi risiko informasi yang terstruktur dalam konteks rekayasa perangkat lunak berbasis GenAI. Penelitian ini berupaya menjawab bagaimana organisasi dapat menyeimbangkan inovasi GenAI dengan tanggung jawab etis dan kebutuhan keamanan siber.

2. METODE PENELITIAN

Penelitian ini menggunakan metode Tinjauan Literatur Sistematis (Systematic Literature Review). Pendekatan kualitatif ini dipilih untuk mengidentifikasi, menganalisis, dan mensintesis penelitian yang relevan dengan topik etika, keamanan, dan tata kelola GenAI dalam pengembangan perangkat lunak [7].

1. Protokol Pencarian Data

Pengumpulan data primer dilakukan melalui analisis dokumen yang dilampirkan oleh pengguna [2], [8], [9]. Data sekunder dikumpulkan dari basis data akademik (IEEE, ACM, Scopus) dan sumber terbuka bereputasi (arXiv, NIST) dengan batasan tahun 2021–2024. Kata kunci yang difokuskan meliputi: "Generative AI Security", "AI Ethics Framework", "Information Risk Management", dan "UU PDP Indonesia".

2. Kriteria Inklusi dan Eksklusi

Kriteria inklusi mencakup:

1. Artikel jurnal yang membahas GenAI dan Software Engineering.
2. Laporan atau dokumen resmi dari lembaga berwenang (Pemerintah/NIST).
3. Tersedia secara open access atau memiliki DOI/URL yang dapat diakses (memenuhi permintaan untuk impor Mendeley).

Analisis data dilakukan dengan mengekstraksi kerangka kerja [10], mengidentifikasi ancaman keamanan [4], dan menyintesis pedoman etika [11].

3. Analisis Data

Data yang terkumpul dianalisis menggunakan teknik analisis isi (content analysis). Informasi diekstraksi berdasarkan tiga tema utama: peran GenAI dalam SDLC, isu etika yang muncul, dan strategi mitigasi risiko. Validitas data dijaga melalui triangulasi sumber antara literatur akademik internasional dan regulasi nasional seperti UU PDP dan pedoman etika AI dari Kementerian Komunikasi dan Informatika [11].

1. HASIL DAN PEMBAHASAN

1. Peran Transformasi GenAI dalam Rekayasa Perangkat Lunak

Hasil telaah literatur menunjukkan bahwa GenAI telah mengubah lanskap pengembangan perangkat lunak secara fundamental. Teknologi ini tidak hanya berfungsi sebagai alat bantu autokomplet, tetapi juga sebagai mitra pemrograman (pair programmer). [2] mencatat bahwa Generative Adversarial Networks (GANs) dan arsitektur Transformer memungkinkan otomatisasi pada tahap penulisan kode, refactoring, hingga pembuatan dokumentasi sistem.

Efisiensi menjadi nilai tawar utama. Pengembang yang menggunakan alat seperti GitHub Copilot dilaporkan dapat menyelesaikan tugas 56% lebih cepat dibandingkan yang tidak menggunakan [2]. Selain itu, GenAI memfasilitasi demokratisasi pengembangan perangkat lunak, di mana individu dengan pengetahuan teknis terbatas dapat membangun prototipe fungsional melalui instruksi bahasa alami (prompt), yang kemudian diterjemahkan menjadi kode oleh model [12].

2. Implikasi Etis GenAI dalam SDLC

Penggunaan GenAI berisiko mengaburkan akuntabilitas moral dalam pengembangan perangkat lunak. Masalah

HKI muncul ketika model menghasilkan kode yang sangat identik dengan data latih berlisensi [13]. Sementara itu, model LLM sering bertindak sebagai black box, mempersulit upaya audit dan penjelasan, yang bertentangan dengan prinsip transparansi dalam Trustworthy AI [14]. Ditegaskan bahwa pertanggungjawaban atas output kode harus selalu kembali pada pengembang manusia, bukan pada mesin [8].

3. Kajian Etika: Bias, Transparansi, dan Akuntabilitas

Meskipun menawarkan efisiensi, integrasi GenAI memunculkan dilema etika yang kompleks. [8] menekankan bahwa sistem AI tidak memiliki agensi moral, sehingga tanggung jawab penuh atas kode yang dihasilkan tetap berada pada manusia.

- **Bias Algoritmik:** Model GenAI dilatih menggunakan miliaran baris kode dari repositori publik. Jika data latih tersebut mengandung pola kode yang buruk, komentar yang bias secara sosial, atau praktik keamanan yang usang, model akan mereproduksinya. Hal ini berisiko melanggengkan diskriminasi atau praktik coding yang tidak standar [3], [15].
- **Transparansi (Black Box):** Ketidakmampuan untuk menjelaskan bagaimana model LLM tiba pada potongan kode tertentu menjadi masalah serius dalam sistem kritis (safety-critical systems). Prinsip Explainable AI (XAI) menjadi sulit diterapkan pada deep learning, yang bertentangan dengan kebutuhan auditabilitas dalam standar industri [8].
- **Hak Kekayaan Intelektual (HKI):** GenAI sering kali menghasilkan kode yang sangat mirip dengan data latihnya. Hal ini memicu perdebatan hukum mengenai apakah output AI melanggar lisensi open-source (seperti GPL atau MIT) dari kode asli yang dipelajari oleh model tersebut [2].

4. Manajemen Risiko Informasi dan Keamanan

Risiko informasi dalam pemanfaatan GenAI dapat dikategorikan menjadi risiko keamanan teknis dan risiko kepatuhan data.

1. Risiko Keamanan Siber (Vulnerabilities): GenAI rentan terhadap hallucination, di mana model menghasilkan kode yang menggunakan pustaka (library) yang tidak ada atau memanggil fungsi yang tidak aman. Penelitian [4] menemukan bahwa sekitar 40% kode yang dihasilkan oleh Copilot dalam skenario tertentu mengandung kerentanan keamanan (security bugs).
2. Kebocoran Data Sensitif: Pengembang sering kali secara tidak sengaja memasukkan API Keys, kredensial basis data, atau logika bisnis rahasia ke dalam prompt GenAI. Karena model AI dapat menggunakan input ini untuk pelatihan ulang, data rahasia tersebut berisiko terekspos ke pengguna lain [9]
3. Kepatuhan Regulasi (UU PDP): Di Indonesia, UU No. 27 Tahun 2022 mewajibkan pengendali data untuk menjamin keamanan pemrosesan data pribadi. Penggunaan GenAI yang memproses data pengguna tanpa enkripsi atau persetujuan yang jelas dapat dikategorikan sebagai pelanggaran hukum [6], [16]
5. Strategi Mitigasi dan Tata Kelola

Untuk mengatasi risiko tersebut, organisasi perlu menerapkan kerangka kerja tata kelola AI yang ketat. [11] dan [8] menyarankan pendekatan Human-in-the-Loop (HITL), di mana setiap kode yang dihasilkan AI wajib ditinjau dan divalidasi oleh manusia sebelum diimplementasikan ke lingkungan produksi.

Selain itu, penerapan kebijakan "Zero Trust" terhadap output AI harus ditegakkan. Pengembang harus menggunakan alat pemindaian keamanan statis (SAST) untuk mendeteksi kerentanan pada kode yang dihasilkan AI [17]. Dari sisi kebijakan,

organisasi harus menetapkan aturan yang melarang pengiriman data rahasia ke model AI publik dan beralih ke model enterprise yang menjamin isolasi data [9].

4. KESIMPULAN

Penelitian ini menyimpulkan bahwa Generative AI merupakan pedang bermata dua dalam pengembangan perangkat lunak. Di satu sisi, teknologi ini menawarkan akselerasi produktivitas dan inovasi yang signifikan. Di sisi lain, tanpa manajemen risiko yang memadai, GenAI memperkenalkan ancaman serius terhadap integritas kode, keamanan data, dan kepatuhan hukum. Kunci keberhasilan adopsi GenAI bukan terletak pada kecanggihan teknologi itu sendiri, melainkan pada kedewasaan etika penggunaannya dan kekuatan kerangka tata kelola yang diterapkan. Kepatuhan terhadap regulasi nasional seperti UU PDP dan penerapan prinsip verifikasi manusia adalah syarat mutlak dalam ekosistem pengembangan perangkat lunak berbasis AI.

Saran

1. Bagi Praktisi/Pengembang: Wajib melakukan verifikasi manual dan pengujian keamanan (security testing) terhadap setiap baris kode yang dihasilkan oleh AI. Jangan pernah mengunggah data sensitif atau proprietary code ke mesin GenAI publik.
2. Bagi Institusi Pendidikan: Kurikulum informatika perlu mengintegrasikan materi etika AI dan keamanan siber secara mendalam, bukan hanya keterampilan teknis penggunaan alat [9].
3. Bagi Regulator: Diperlukan pedoman teknis turunan dari UU PDP yang secara spesifik mengatur penggunaan GenAI di sektor industri untuk memberikan kepastian hukum.

5. DAFTAR PUSTAKA

- A. Marsiela, "Asesmen Penggunaan Kecerdasan Buatan (AI) dalam Media di Indonesia: Tantangan, Peluang, dan

- Pengembangan Standar Etika,” Aliansi Jurnalis Indep., pp. 1–41, 2024.
- B. Zhu, N. Mu, J. Jiao, and D. Wagner, “Generative AI Security: Challenges and Countermeasures,” pp. 1–16, 2024, [Online]. Available: <http://arxiv.org/abs/2402.12617>
- D. B. Raharjo, Teori Etika. 2019.
- D. Sameer, “The Impact of Generative AI on Software Engineering Activities,” vol. 10, no. 2, 2024.
- D. Smit, H. Smuts, P. Louw, J. Pielmeier, and C. Eidelloth, “The Impact of GitHub Copilot on Developer Productivity From a Software Engineering Body of Knowledge Perspective,” 30th Am. Conf. Inf. Syst. AMCIS 2024, no. August, 2024.
- Direktorat et al., “Panduan Penggunaan Generative Artificial Intelligence (GenAI),” Pandu. Pengguna. Gener. Artif. Intell., p. 139, 2024.
- E. M. Bender, T. Gebru, A. McMillan-Major, and S. Shmitchell, “On the dangers of stochastic parrots: Can language models be too big?,” FAccT 2021 - Proc. 2021 ACM Conf. Fairness, Accountability, Transpar., pp. 610–623, 2021, doi: 10.1145/3442188.3445922.
- F. Bintarawati, “the Influence of the Personal Data Protection Law (Uu Pdp) on Law Enforcement in the Digital Era,” ANAYASA J. Leg. Stud., vol. 1, no. 2, pp. 135–143, 2024, doi: 10.61397/ays.v1i2.92.
- G. Sandoval, H. Pearce, T. Nys, R. Karri, S. Garg, and B. Dolan-Gavitt, “Lost at C: A User Study on the Security Implications of Large Language Model Code Assistants,” 32nd USENIX Secur. Symp. USENIX Secur. 2023, vol. 3, pp. 2205–2222, 2023.
- H. Pearce, B. Ahmad, B. Tan, B. Dolan-Gavitt, and R. Karri, “Asleep at the Keyboard? Assessing the Security of GitHub Copilot’s Code Contributions,” Proc. - IEEE Symp. Secur. Priv., vol. 2022-May, pp. 754–768, 2022, doi: 10.1109/SP46214.2022.9833571.
- L. Weidinger et al., “Sociotechnical Safety Evaluation of Generative AI Systems,” 2023, [Online]. Available: <http://arxiv.org/abs/2310.11986>
- Laurie E. Locascio, “Artificial Intelligence Risk Management NIST AI 100-1 Artificial Intelligence Risk Management,” Artif. Intell. Risk Manag. Framew. (AI RMF 1.0), pp. 100–101, 2023.
- M. Chen et al., “Evaluating Large Language Models Trained on Code,” 2021, [Online]. Available: <http://arxiv.org/abs/2107.03374>
- M. Ishtiaq, “Book Review Creswell, J. W. (2014). Research Design: Qualitative, Quantitative and Mixed Methods Approaches (4th ed.). Thousand Oaks, CA: Sage,” English Lang. Teach., vol. 12, no. 5, p. 40, 2019, doi: 10.5539/elt.v12n5p40.
- P. Vaithilingam, T. Zhang, and E. L. Glassman, “Expectation vs. Experience: Evaluating the Usability of Code Generation Tools Powered by Large Language Models,” Conf. Hum. Factors Comput. Syst. - Proc., 2022, doi: 10.1145/3491101.3519665.
- Pemerintah Indonesia, “Personal Data Protection Law,” Introd. to Turkish Bus. Law, no. 016999, pp. 457–483, 2022.
- Theguh, B. Bisri, and F. Nawawi, “Etika Dalam Pemanfaatan Artificial Intelligence (Ai) Pada Generasi Z Di Pondok Pesantren Syariful Anam Kota Cirebon,” J. Abadimas Adi Buana, vol. 7, no. 02, pp. 169–179, 2024, doi: 10.36456/abadimas.v7.i02.a7916.