



Solusi AI Untuk Melindungi Privasi dan Keamanan Data Pengguna *AI Solutions to Protect User Privacy and Data Security*

Zen Munawar¹, Sri Sutjiningtyas², Novianti Indah Putri³, Milla Marlina⁴, Rita Komalasari⁵, Herru Soerjono⁶

^{1,5}Manajemen Informatika, Politeknik LP3I

²Teknik Informatika, Ilmu Komputer dan Informatika, Universitas Nurtanio

³Sistem Informasi, Ilmu Komputer dan Sistem Informasi, Universitas Kebangsaan Republik Indonesia

^{4,6}Administrasi Bisnis, Politeknik LP3I

¹munawarzen@gmail.com, ²srisutjiningtyas70@gmail.com, ³noviantiindahputri2021@gmail.com, ⁴millamarlina@pln.ac.id,

⁵ritakomalasari@plb.ac.id, ⁶herrusoerjono2022@gmail.com.

Abstract

The integration of AI into massive data poses serious privacy risks. This study analyzes technical solutions such as differential privacy to balance data protection with AI system performance. This study aims to analyze crucial issues related to data privacy and security arising from the massive integration of Artificial Intelligence (AI) across various industries and to evaluate technical approaches that can maintain the confidentiality of user information without compromising system performance. The main focus of this study is to assess the effectiveness of data protection mechanisms such as differential privacy, federated learning, homomorphic encryption, and secure multi-party computation in mitigating the risks of privacy breaches and unauthorized access. This study uses a qualitative approach with a systematic literature review method to evaluate privacy solutions in AI systems. The results show that techniques such as federated learning and homomorphic encryption are effective in protecting personal data in AI systems. However, their implementation faces trade-off challenges in the form of decreased computational speed and model accuracy. Design recommendations focus on the balance between security and performance to support safe and scalable AI regulations. Data protection techniques such as federated learning are effective in maintaining privacy, but they risk degrading AI computational performance. In conclusion, a balance between security and efficiency is needed. The recommendation is that developers should adopt privacy-by-design and that governments should develop technical regulations that support safe and scalable AI innovation for the public.

Keywords: AI, Data Protection, Differential Privacy, Cybersecurity

Abstrak

Integrasi AI pada data masif memicu risiko privasi serius. Penelitian ini menganalisis solusi teknis seperti privasi diferensial untuk menyeimbangkan perlindungan data dengan kinerja sistem AI. Penelitian ini bertujuan untuk menganalisis isu-isu krusial terkait privasi dan keamanan data yang muncul akibat integrasi masif Kecerdasan Buatan (AI) di berbagai industri, serta mengevaluasi pendekatan teknis yang dapat menjaga kerahasiaan informasi pengguna tanpa mengorbankan kinerja sistem. Fokus utama penelitian ini adalah mengkaji efektivitas mekanisme perlindungan data seperti privasi diferensial, pembelajaran federasi, enkripsi homomorfik, dan komputasi multi-pihak yang aman dalam memitigasi risiko pelanggaran privasi dan akses ilegal. Penelitian ini menggunakan pendekatan kualitatif dengan metode studi literatur sistematis untuk mengevaluasi solusi privasi pada sistem kecerdasan buatan. Hasil penelitian menunjukkan bahwa teknik seperti federated learning dan enkripsi homomorfik efektif melindungi data pribadi dalam sistem AI. Namun, implementasinya menghadapi tantangan trade-off berupa penurunan kecepatan komputasi dan akurasi model. Rekomendasi desain berfokus pada keseimbangan antara keamanan dan performa guna mendukung regulasi AI yang aman serta terukur. Teknik perlindungan data seperti federated learning efektif menjaga privasi, namun berisiko menurunkan performa komputasi AI. Kesimpulannya, diperlukan keseimbangan antara keamanan dan efisiensi. Rekomendasinya, pengembang harus mengadopsi desain privacy-by-design dan pemerintah perlu menyusun regulasi teknis yang mendukung inovasi AI yang aman serta terukur bagi masyarakat.

Kata kunci: AI, Perlindungan Data, Privasi Diferensial, Keamanan Siber

1. Pendahuluan

Di era digital saat ini, evolusi teknologi yang pesat telah menyebabkan peningkatan ancaman siber yang signifikan, sehingga memerlukan langkah-langkah cangguh untuk melindungi informasi sensitif. Seiring

dengan terus berkembangnya Kecerdasan Buatan (AI) di berbagai bidang, integrasinya ke dalam kewanaman data pribadi menghadirkan potensi yang tak tertandingi untuk mendeteksi, mencegah, dan menanggapi ancaman siber. AI mentransformasi beberapa sektor karena menghadirkan instrumen baru untuk

menganalisis data, membuat keputusan, dan bahkan mengotomatiskan proses. AI merambah ke setiap sektor, mulai dari perawatan kesehatan, keuangan, dan hampir setiap sektor lainnya melalui inovasi untuk memberikan layanan yang paling personal dan mengoptimalkan operasi. Namun, implementasi sosial sistem AI berdasarkan intensitas data juga menimbulkan kekhawatiran serius terkait privasi dan keamanan [1].

Model kecerdasan buatan dapat dirancang menggunakan sejumlah besar data pengguna, yang meliputi data identitas, kesehatan, keuangan, dan perilaku, dan ini membuat model tersebut rentan terhadap pelanggaran privasi, pencurian data, atau peretasan [2]. Ada dua argumen untuk melindungi data dalam sistem AI: pertama, kecenderungan badan pengatur untuk lebih memperhatikan manajemen data; kedua, implikasi penyalahgunaan data. Seperti yang diilustrasikan oleh Peraturan Perlindungan Data Pribadi, ada standar tinggi yang harus dipenuhi organisasi saat menangkap, menyimpan, dan menggunakan informasi pribadi. Karena organisasi harus mematuhi peraturan ini, ditambah dengan tuntutan teknologi yang diperlukan agar AI dapat berkinerja optimal saat berkembang, hal inilah yang membuat perlindungan privasi menjadi sangat sulit dalam sistem AI. Penelitian ini membahas isu-isu yang berkaitan dengan privasi dan keamanan data dalam AI, meninjau langkah-langkah yang dapat diadopsi untuk meningkatkan keamanan data pengguna, dan menganalisis pengambilan keputusan antara privasi, efisiensi, dan kapasitas sistem. Penelitian ini menyajikan pedoman tentang bagaimana pengembang AI dapat mencapai keseimbangan antara faktor-faktor tersebut.

2. Metode Penelitian

2.1. Pendekatan dan Pengumpulan Data Penelitian

Penelitian ini menggunakan pendekatan kualitatif dengan metode studi literatur sistematis untuk mengevaluasi solusi privasi pada sistem kecerdasan buatan. Data dikumpulkan melalui penelusuran literatur sekunder yang bersumber dari basis data ilmiah bereputasi, termasuk jurnal internasional, prosiding konferensi, serta laporan teknis terkini yang berkaitan dengan keamanan data dan AI. Proses pemilihan literatur dilakukan dengan menggunakan kata kunci spesifik seperti "AI privacy", "federated learning", "differential privacy", dan "secure multi-party computation" guna memastikan relevansi materi dengan fokus pembahasan.

Tahapan analisis data dilakukan melalui tiga fase utama. Identifikasi dan Kategorisasi: Memetakan berbagai ancaman privasi dalam siklus hidup data pada model AI. Komparasi Teknis: Membandingkan

efektivitas dan batasan dari masing-masing metode perlindungan data (privasi diferensial, enkripsi homomorfik, dan lain-lain) terhadap performa komputasi. Sintesis Strategis: Menarik kesimpulan dari temuan-temuan tersebut untuk merumuskan rekomendasi desain solusi AI yang aman dan terukur. Pendekatan ini dipilih agar dapat memberikan gambaran komprehensif mengenai tantangan teknis dan regulasi dalam mengimplementasikan AI yang berorientasi pada perlindungan privasi.

2.2. Tantangan dalam Privasi dan Keamanan Data untuk Sistem AI

Sejumlah besar input atau kumpulan data merupakan dasar dari sifat sistem AI karena membantu dalam memprediksi hasil, rekomendasi yang diharapkan, dan mengotomatiskan proses. Namun, penggunaan kumpulan data yang sangat besar ini membawa risiko besar pelanggaran privasi pengguna [3]. Tantangan privasi dan keamanan data dalam Sistem AI adalah, untuk menghindari kebocoran informasi dalam sistem dan orang-orang; memastikan bahwa data yang dikumpulkan oleh sistem tertentu berkualitas tinggi tanpa melanggar hak privasi orang; kepatuhan terhadap undang-undang privasi yang ditetapkan pada peraturan perlindungan data pribadi; risiko yang ditimbulkan oleh berbagi data dan akses pihak ketiga; dan, beban komputasi dan efek kinerja dari teknik privasi. Semua masalah ini menimbulkan masalah terkait keamanan dan kepercayaan yang diberikan dalam memenuhi kebutuhan peserta didik untuk pembelajaran yang tepat.

2.3. Sensitivitas Data dan Risiko Privasi

Sebagian besar aplikasi AI perlu bekerja dengan informasi pribadi, sensitif, atau rahasia [4]. Data ini pada dasarnya seringkali cukup sensitif, terutama di sektor-sektor seperti perawatan kesehatan, keuangan, dan e-commerce jika sampai terekspos. Sebagai contoh, model AI dapat digunakan untuk memprediksi kondisi kesehatan dalam kasus penyakit dan gangguan kesehatan, tetapi untuk sampai ke sana kita perlu mengumpulkan catatan pasien, diagnosis, dan detail kesehatan pribadi lainnya. Ancaman potensial yang terkait dengan kebocoran data tersebut adalah pelanggaran privasi yang sangat sensitif termasuk pencurian identitas, penipuan asuransi, dan kebocoran informasi medis yang distigmatisasi [5]. Namun, bahkan jika dianonimkan, data tersebut tetap berisiko diidentifikasi kembali dengan teknik lain. Hal ini bahkan lebih mengkhawatirkan dengan aplikasi seperti sistem rekomendasi yang, misalnya, mengungkapkan banyak hal tentang gaya hidup dan preferensi seseorang.

2.4. Pelanggaran Data dan Akses Tidak Sah

Penyimpanan model dan data AI di lingkungan cloud adalah arah utama ancaman siber. Hal terburuk dapat terjadi ketika seorang peretas berhasil menembus

jaringan tertentu, mereka mendapatkan akses tidak hanya ke data mentah tetapi juga kekayaan intelektual, algoritma, dan aset penting lainnya [6]. Volume data juga meningkat seiring dengan perluasan ukuran dan aplikasi sistem AI, dan data yang sangat besar ini menjadi tempat berlindung bagi peretas. Contoh sistem AI adalah sistem pengemudian otonom yang akan rentan terhadap jenis serangan ini; serangan yang memengaruhi fungsinya, seperti pencurian data pengguna atau bahkan manipulasi prediksi, akan memiliki konsekuensi yang serius. Masalah signifikan lainnya terkait dengan pihak internal: di antara ancaman potensial, pihak internal sangat berbahaya [7]. Kesalahan penanganan informasi oleh karyawan atau kontraktor yang bertanggung jawab menangani data ini kemungkinan terjadi karena berbagai alasan, seperti niat jahat yang disengaja terhadap data tersebut. Karena model AI membutuhkan akses konstan ke data pengguna untuk melatih modelnya, kerentanan seperti itu merupakan masalah keamanan yang nyata.

2.4. Kurangnya Transparansi, Kemampuan Penjelasan, dan Skalabilitas versus Privasi

Jumlah masalah privasi dan keamanan data semakin diperparah oleh fakta bahwa banyak model AI, termasuk model pembelajaran mendalam, berfungsi seperti 'kotak hitam'. Model-model ini sampai pada kesimpulan setelah menggunakan sejumlah fungsi dengan sedikit atau tanpa penjelasan tentang bagaimana tepatnya keputusan untuk mengambil satu arah atau arah lain diambil. Kurangnya interpretasi ini khususnya menjadi penyebab kekhawatiran dari perspektif privasi karena pengguna tidak pernah memahami bagaimana data mereka diproses [8]. Transparansi sangat penting di beberapa bidang, termasuk peradilan pidana atau perawatan kesehatan, di mana adopsi keputusan AI dapat menyebabkan hasil yang ekstrem. Ada kategori pendekatan yang, jika kurang memiliki transparansi dan interpretasi yang memadai, akan gagal mendapatkan penerimaan publik dan dapat menjadi subjek regulasi karena akan dianggap tidak transparan dan tidak adil.

Hal ini paling jelas terlihat karena salah satu masalah utama sistem AI adalah volume data yang besar yang harus dikelola sambil menghormati hak privasi individu [9]. Setiap AI, khususnya model pembelajaran mendalam, membutuhkan kumpulan data besar untuk pelatihan, guna mencapai tingkat akurasi kinerja tinggi. Namun demikian, satu kelemahan utama muncul dari fakta bahwa meskipun sejumlah besar data dapat diproses melalui kumpulan data besar ini, privasi individu dalam data besar tersebut perlu dilindungi saat memproses informasi. Langkah-langkah keamanan yang diperlukan untuk melindungi privasi, meskipun penting, tidak dapat diimplementasikan tanpa biaya komputasi tambahan, yang membuatnya sulit untuk meningkatkan skala sistem AI [10].

Sebagai contoh, privasi diferensial yang mengatur noise untuk menjamin privasi dapat menurunkan akurasi umum model. Konfigurasi jaringan yang disajikan memperkenalkan noise ke dalam sistem, atau dengan kata lain, mengurangi akurasi model dan karenanya efektivitas sistem. Dengan cara yang sama, mekanisme data yang kompleks memungkinkan pemrosesan data tanpa mengungkapkan informasi spesifik. Namun, teknik-teknik ini melibatkan kompleksitas komputasi dan komunikasi yang lebih tinggi yang menyebabkan waktu pemrosesan lebih lambat dan menyulitkan untuk mengakomodasi data lapangan yang terus bertambah ukurannya. Kompromi antara privasi dan skala semacam itu merupakan tantangan yang sulit bagi organisasi yang merancang sistem AI. Diperlukan pertimbangan terus-menerus terhadap teknik-teknik yang menjaga privasi yang akan menghasilkan overhead rendah dalam hal kinerja sistem, yang merupakan kontradiksi dengan kebutuhan akan privasi.

3. Hasil dan Pembahasan

3.1. Hasil

Teknik Pelestarian Privasi untuk AI. Karena munculnya kekhawatiran keamanan terkait privasi data, metode-metode berikut telah diajukan untuk mengurangi faktor risiko tanpa mengorbankan efektivitas AI. Metode-metode ini dimaksudkan untuk menjaga keamanan data pengguna di seluruh siklus hidup data AI, termasuk akuisisi data, pengembangan model, dan penggunaan. Privasi diferensial untuk data guna memenuhi privasi individu, pembelajaran federasi melatih pada beberapa perangkat tanpa berbagi data mentah [11]. Enkripsi homomorfik memungkinkan komputasi pada data yang dienkripsi, dan komputasi multi-pihak yang aman sangat berguna ketika pihak-pihak yang berbeda perlu menganalisis data bersama-sama, tanpa mengungkapkannya satu sama lain. Semua pendekatan ini membantu mencapai ukuran privasi yang tepat sekaligus meningkatkan efisiensi AI yang dibutuhkan.

Privasi diferensial adalah bentuk privasi statistik di mana privasi setiap catatan tidak dapat dikompromikan dari catatan lain. Pendekatan ini digunakan secara universal dalam sistem AI di mana perlu untuk mengevaluasi hasil yang diperoleh tanpa mengungkapkan input individu. Misalnya, Google menggunakan konsep yang dikenal sebagai privasi diferensial, dalam sistem penambangan datanya, perusahaan mengumpulkan data agregat dari perangkat pengguna tanpa melanggar hak privasi setiap pengguna. Privasi diferensial juga digunakan dalam bidang yang dikenal sebagai pembelajaran federasi di mana beberapa perangkat memungkinkan pelatihan model AI tanpa perlu berbagi data mentah dengan yang lain. Meskipun demikian, estimasi kepadatan menggunakan mekanisme menawarkan privasi pada titik data individual [12].

Pembelajaran federasi adalah salah satu solusi unik yang memungkinkan pelatihan model AI di seluruh perangkat terdesentralisasi sambil menyimpan data penting hanya di perangkat [13]. Dibandingkan dengan metode tradisional berbagi data mentah untuk pemrosesan lebih lanjut di server pusat, pembelajaran federasi hanya berbagi gradien, sehingga memastikan bahwa model dilatih tanpa mengakses informasi pribadi orang [14]. Oleh karena itu, pembelajaran federasi paling tepat diterapkan dalam situasi di mana privasi data merupakan pertimbangan penting, di perangkat seluler, perawatan kesehatan, dan keuangan. Misalnya, di sektor perawatan kesehatan, rumah sakit dapat secara kolektif melatih model pembelajaran mesin untuk prognosis hasil pasien tetapi bukan data pasien. Terlepas dari semua rencana pembelajaran federasi ini yang dapat mengurangi risiko privasi, namun masih menunjukkan beberapa tantangan heterogenitas data, efisiensi komunikasi, dan konvergensi model.

Enkripsi homomorfik adalah jenis kriptografi yang memungkinkan komputasi dilakukan pada data yang telah dienkripsi. Hal ini memungkinkan model AI untuk menelusuri dan menganalisis data sensitif karena data tersebut diamankan dari paparan di balik bentuk terenkripsi saat diproses. Teknik ini berpotensi mengubah status privasi dalam AI dan khususnya di bidang keuangan di mana data keuangan dapat diproses tanpa pengungkapan lebih lanjut [15]. Meskipun enkripsi homomorfik sangat menjamin keamanan, komputasi dengan data yang dienkripsi secara homomorfik mahal. Melakukan operasi pada data terenkripsi jauh lebih menuntut secara komputasi daripada operasi pada data teks biasa, yang menyulitkan sistem AI untuk berkembang [16]. Namun, saat ini sedang dilakukan upaya untuk mempercepat enkripsi homomorfik, dan penelitian telah menunjukkan bahwa dengan mempertimbangkan beban komputasi dan waktu, kemajuan telah dicapai untuk mengurangi aspek-aspek enkripsi homomorfik tersebut [17].

Komputasi multi-pihak yang aman. Komputasi multi-pihak yang aman adalah metode kriptografi di mana sejumlah pihak dapat mengerjakan suatu fungsi untuk menghasilkan hasil yang serupa dengan bantuan data masing-masing tanpa pihak lain dapat melihat data mereka. Pada dasarnya, dalam penerapan AI, melalui Komputasi multi-pihak yang aman, berbagai organisasi atau entitas dapat bekerja sama dan melatih model pembelajaran mesin tanpa benar-benar berbagi kumpulan data mereka sendiri. Misalnya, dalam penelitian medis, beberapa rumah sakit dapat berbagi data melalui model AI terfederasi untuk memperkirakan hasil penyakit yang tidak dibagikan antar rumah sakit yang berbeda [18]. Komputasi multi-pihak yang aman juga memastikan bahwa informasi setiap organisasi yang diterima tetap aman, tetapi pada saat yang sama dapat dibagikan antar lembaga. Meskipun demikian, metode ini memiliki beberapa

kekurangan; termasuk biaya komputasi yang tinggi dan mungkin mengalami beberapa masalah dalam komunikasi dan sinkronisasi antar pihak.

3.2. Pembahasan

Menyeimbangkan Privasi dan Akurasi Model. Dampak Kebisingan pada Akurasi. Pendekatan seperti privasi diferensial bertindak secara berharga dengan menambahkan kebisingan pada data atau model yang bertujuan untuk menjaga privasi individu [19]. Namun demikian, kebisingan ini cenderung mengurangi efektivitas model AI, terutama ketika presisi merupakan penentu utama dalam penerapan model seperti di bidang perawatan kesehatan atau keuangan. Misalnya, ketika menyangkut diagnosis penyakit di rumah sakit, kesalahan kecil dalam keluaran model akan menyebabkan efek negatif pada klien. Oleh karena itu, efektivitas teknik peningkatan privasi tetap menjadi masalah kritis; perlu untuk menjamin bahwa teknik tersebut tidak akan secara signifikan menurunkan kinerja model [20]. Sambil mencoba meminimalkan efek kebisingan, ada upaya berkelanjutan untuk membangun tingkat jaminan privasi yang lebih tinggi. Salah satu pendekatan tersebut adalah penerapan skema amplifikasi privasi yang memungkinkan untuk mengurangi tingkat kebisingan yang dibutuhkan tanpa menurunkan kinerja model.

Beban Komputasi Tambahan. Teknologi peningkatan privasi seperti enkripsi homomorfik penuh dan komputasi multipihak yang aman misalnya, dapat menimbulkan beban tambahan yang besar pada sistem yang sedang dibangun untuk AI. Mempertimbangkan waktu, sumber daya, dan peningkatan kebutuhan komputasi untuk enkripsi data atau untuk melakukan komputasi yang aman menghambat kemampuan model AI dan mempertajam inefisiensinya, terutama dalam penggunaan waktu nyata yang membutuhkan respons invasif yang cepat [21]. Untuk mengatasi tantangan ini, upaya penelitian terbaru telah berupaya untuk menyempurnakan metode pelestarian privasi ini untuk meningkatkan kinerjanya. Misalnya, telah disarankan bahwa mengintegrasikan pembelajaran terawasi/tidak terawasi klasik dengan metode yang meningkatkan privasi dapat mengurangi beban komputasi dengan tetap menjaga privasi.

Kompleksitas dan Interpretasi Model. Pada kenyataannya, berbagai teknik pelestarian privasi, terutama pembelajaran federasi dan komputasi Multipihak Aman, memaksakan tingkat peningkatan lain pada model AI. Model-model ini mungkin melibatkan integrasi lebih dari satu perangkat atau institusi yang membuatnya sulit untuk dikelola atau diinterpretasikan. Selain itu, karena meningkatnya kerumitan sistem, menjadi sulit untuk memverifikasi bahwa model berjalan sesuai program, yang pada gilirannya dapat memperbesar kemungkinan kesalahan atau pengambilan keputusan yang tepat dalam sistem

[22]. Meningkatkan interpretasi subjek model dan membuat metode peningkatan privasi lebih mudah diterapkan tanpa mengurangi keamanan masih menjadi masalah terbuka di bidang ini. Dengan demikian, hanya sistem kecerdasan buatan yang aman dan sekaligus dapat dijelaskan yang dapat dipercaya oleh pengguna dan disetujui oleh regulator.

Praktik Terbaik untuk Memastikan Privasi dan Keamanan Data dalam AI. Minimisasi dan Anonimisasi Data. Privasi melalui minimisasi data adalah konsep di mana sistem AI hanya mengumpulkan data yang diperlukan untuk menjalankan fungsinya. Membatasi jenis dan jumlah data yang dikumpulkan meminimalkan risiko perusahaan untuk merugikan privasi individu dan memastikan bahwa perusahaan mematuhi undang-undang privasi. Ini melibatkan langkah-langkah anonimisasi yang mengacu pada contoh detail pribadi dengan pengidentifikasi yang disamarkan atau diganti dengan yang palsu. Metode-metode ini tidak hanya mengurangi ancaman privasi tetapi juga meningkatkan keamanan informasi pribadi yang memungkinkan organisasi untuk menggunakan data tersebut untuk analisis sambil tetap menjaga anonimitas pengguna [23].

Transparansi dan Kontrol Pengguna. Karena meningkatnya kesadaran pengguna akan potensi risiko, transparansi memainkan peran penting dalam membangun kepercayaan pada organisasi. Untuk tujuan ini, organisasi harus dapat menjelaskan data apa yang dikumpulkan, bagaimana data tersebut akan digunakan, dan langkah-langkah apa yang diambil untuk memastikan privasi informasi yang dikumpulkan. Memberikan informasi ini membantu pengguna untuk membuat keputusan yang tepat saat berinteraksi dengan sistem cerdas. Lebih lanjut, sangat penting bahwa pengguna sebagai individu memiliki hak-hak tertentu atas kontrol tersebut mereka memiliki hak untuk melihat, menghapus, atau menolak pengambilan data. Langkah-langkah ini tidak hanya memenuhi persyaratan peraturan privasi tetapi juga meningkatkan kepercayaan pengguna dan karenanya meningkatkan interaksi mereka dengan produk.

Kepatuhan terhadap Regulasi. Kecerdasan buatan (AI) diharuskan untuk mengikuti hukum dan kebijakan regulasi perlindungan data di seluruh dunia, seperti Peraturan Perlindungan Data Umum, Peraturan Perlindungan Data Pribadi, dan sebagainya. Organisasi harus mengintegrasikan privasi dalam segala hal untuk mengembangkan dan menggunakan sistem AI mereka, yang disebut sebagai privasi berdasarkan desain (privacy by design). Penilaian Dampak Perlindungan Data (Data Protection Impact Assessment) sangat penting untuk memeriksa dampak dan risiko yang mungkin terjadi pada privasi dan mengendalikannya dengan benar. Lebih lanjut, organisasi perlu dapat menunjukkan bahwa aktivitas pemrosesan dan tujuan

pemrosesan mereka jelas, mudah diidentifikasi, didokumentasikan dengan benar, dan dapat dibenarkan berdasarkan persyaratan hukum. Perusahaan dapat menjamin kepatuhan terhadap peraturan, berupaya melindungi hak-hak pengguna yang berbeda, dan bertanggung jawab atas pemrosesan data pribadi.

Audit Keamanan dan Pengujian Penetrasi Berkala. Risiko pelanggaran keamanan harus ditangani dengan tinjauan keamanan berkala dan pengujian penetrasi sebagai cara untuk mengungkap kelemahan struktural dalam sistem AI dan memberikan perlindungan komprehensif terhadap ancaman siber. Tidak seperti pengujian normal, pengujian penetrasi melibatkan simulasi serangan dunia nyata yang menunjukkan bagaimana sistem akan merespons ancaman keamanan ini dan area mana yang paling rentan terhadap serangan tersebut. Pendekatan-pendekatan ini bermanfaat untuk melindungi organisasi dari berbagai risiko dengan upaya mencegah risiko tersebut dimanfaatkan. Selain itu, sistem AI harus ditingkatkan dari waktu ke waktu terkait dengan ancaman keamanan yang lebih baru. Manajemen kerentanan dan perbaikan kerentanan merupakan langkah penting dalam memastikan bahwa sistem yang menggunakan AI tidak dieksploitasi seiring berjalannya waktu.

4. Kesimpulan

Karena penggunaan teknologi AI terus berkembang dan diadopsi di berbagai bidang, perlindungan data pengguna menjadi isu yang sangat penting. Masalah privasi dan perlindungan data dalam sistem AI disebabkan oleh meningkatnya volume informasi pribadi dan sensitif yang diproses oleh sistem AI, yang merupakan target menarik bagi ancaman siber. Untuk mengatasi tantangan ini, semakin banyak teknik seperti privasi diferensial, pembelajaran federasi, enkripsi homomorfik, dan komputasi multipihak yang aman telah muncul. Teknik-teknik ini memungkinkan data untuk dianalisis atau diproses tetapi dengan invasi minimal terhadap privasi pemiliknya, sehingga mengurangi kerentanan privasi. Namun, metode-metode ini memberikan manfaat privasi yang besar, tetapi harus diimbangi dengan tujuan untuk mendapatkan akurasi, skalabilitas, dan efisiensi sistem yang tinggi. Mencapai keseimbangan ini tidak mudah, desain dan adopsi praktik terbaik dalam pengembangan AI sangat penting. Dalam studi ini, anggota organisasi harus mengikuti minimalisasi data, yang berarti mereka hanya dapat menggunakan data yang relevan dengan tugas-tugas tertentu, di antara hal-hal lain yang harus dikomunikasikan kepada pengguna data tersebut. Namun, ada juga kebutuhan untuk memenuhi standar berbagai peraturan termasuk Peraturan Perlindungan Data dan Undang-Undang Privasi Konsumen, serta melakukan pemeriksaan keamanan biasa, termasuk pengujian penetrasi. Secara keseluruhan, dimungkinkan untuk memfasilitasi konsep AI etis yang menghormati

privasi pengguna dan mengganggu kinerja sebagai fungsi penting dengan menggunakan kerangka kerja yang diusulkan.

Daftar Rujukan

- [1] Z. Munawar, "Research developments in the field neurocomputing," in *Proceedings of 2016 4th International Conference on Cyber and IT Service Management, CITSM 2016*, 2016, no. 59, pp. 1–6.
- [2] Z. Munawar, "Research developments in the field neurocomputing," in *2016 4th International Conference on Cyber and IT Service Management*, 2016, pp. 1–6.
- [3] W. Dong, C. Lin, X. He, X. Huang, and S. Xu, "Privacy-Preserving Federated Learning via Homomorphic Adversarial Networks," in *ICLR 2025 Conference*, 2024, pp. 1–20.
- [4] Z. Munawar, "Manfaat Teknologi Informasi di Masa Pandemi Covid-19," *J-SIKA|Jurnal Sist. Inf. Karya Anak Bangsa*, vol. 3, no. 2, pp. 53–63, Dec. 2021.
- [5] Hartatik, M. E. Koibur, A. W. Murdiyanto, and Z. Munawar, *Sains data : strategi, teknik, dan model analisis data*, First. Bandung: Kaizen Media Publishing, 2023.
- [6] Z. Munawar, S. Sutjiningtyas, N. I. Putri, R. Komalasari, and H. Soerjono, "Manfaat Kecerdasan Buatan pada Proses Belajar Mengajar di Pendidikan Tinggi," *J. Teknol. Inf. Dan Komun.*, vol. 11, no. 2, pp. 213–224, 2024.
- [7] J. Shen, Y. Zhao, S. Huang, and Y. Ren, "Secure and Flexible Privacy-Preserving Federated Learning Based on Multi-Key Fully Homomorphic Encryption," *Electronics*, vol. 13, no. 22, pp. 1–23, 2024.
- [8] V. V. Vegesna, "Privacy-Preserving Techniques in AI-Powered Cyber Security: Challenges and Opportunities," *Int. J. Mach. Learn. Sustain. Dev.*, vol. 5, no. 4, p. 8, 2023.
- [9] Z. Munawar and N. I. Putri, "Keamanan Jaringan Komputer Pada Era Big Data," *J-SIKA|Jurnal Sist. Inf. Karya Anak Bangsa*, vol. 2, no. 01 SE-Articles, pp. 14–20, Jul. 2020.
- [10] Z. Munawar and N. I. Putri, "Keamanan Jaringan Komputer Pada Era Big Data," *J-SIKA| J. Sist. Inf. Karya Anak Bangsa*, vol. 02, no. 1, pp. 14–20, 2020.
- [11] R. R. Palle and K. C. R. Kathala, "Privacy in the Age of Innovation," in *AI Solutions for Information Security*, Berkeley,: Apress, 2024, pp. 47–61.
- [12] C. Dwork and A. Roth, *The Algorithmic Foundations of Differential Privacy*. 2014.
- [13] V. V. Vegesna, "Research on differential privacy protection algorithm for federated learning based on user privacy requirements," *Int. J. Mach. Learn. Sustain. Dev.*, vol. 5, no. 4, pp. 1–8, 2023.
- [14] K. Bonawitz, H. Eichner, and D. Huba, "Towards Federated Learning at Scale: System Design," in *Proceedings of the 2nd System Machine Learning*, 2019, vol. 4, no. 2, p. 15.
- [15] Z. Munawar, "Penerapan Metode Soft Computing Dalam Menyelesaikan Permasalahan Di Bidang Teknik," *Temat. J. Teknol. Inf. Komun.*, vol. 3, no. 2, pp. 1–13, Dec. 2016.
- [16] Z. Munawar, "Machine Learning Approach for Analysis of Social Media," *ADRI Int. Journal. Information. Technol.*, vol. 1, no. 1, pp. 5–8, 2017.
- [17] D. Xin, J. Ji, F. Jing, and M. Gao, "Efficient Fully homomorphic encryption scheme using Ring-LWE," *J. Phys. Conf. Ser.*, vol. 1738, no. 1, p. 9, 2021.
- [18] M. Song et al., "In-situ AI: Towards Autonomous and Incremental Deep Learning for IoT Systems," in *2018 IEEE International Symposium on High Performance Computer Architecture (HPCA)*, 2018, pp. 92–103.
- [19] Z. Munawar, Y. Herdiana, Y. Suharya, and N. I. Putri, "Pemanfaatan Teknologi Digital Di Masa Pandemi Covid-19," *Temat. J. Teknol. Inf. Komun.*, vol. 8, no. 2, pp. 160–175, Dec. 2021.
- [20] M. Abadi and A. Chu, "Deep Learning with Differential Privacy," in *CM SIGSAC Conference on Computer and Communications Security*, 2016, pp. 308–318.
- [21] O. Gozes, M. Frid, H. Greenspan, and D. Patrick, "Rapid AI Development Cycle for the Coronavirus (COVID-19) Pandemic: Initial Results for Automated Detection & Patient Monitoring using Deep Learning CT Image Analysis Article Type: Authors: Summary Statement: Key Results: List of abbreviati," in *arXiv:2003.05037*, 2020.
- [22] T. Li, A. K. Sahu, A. Talwalkar, and V. Smth, "Federated Learning: Challenges, Methods, and Future Directions," *IEEE Signal Process. Mag.*, vol. 37, no. 3, pp. 50–60, 2020.
- [23] DFKI, "Human Centric AI – Intelligent Solutions for the Knowledge Society," *The German Research Center for Artificial Intelligence*, 2015. [Online]. Available: <https://www.dfki.de/en/web/about-us/dfki-at-a-glance/company-profile>.