

Integrasi UTAUT dan TTAT dalam Kerangka IS Dual Factor pada Niat Penggunaan Teknologi AI di Jabodetabek

Manda Raihana Laksita^{1*}, Diena Noviarini², Adnan Kasofi³

Bisnis Digital, Fakultas Ekonomi dan Bisnis, Universitas Negeri Jakarta, Jakarta, Indonesia^{1,2,3}

*Email Korespondensi: mandaraihanal04@gmail.com

Diterima: 08-08-2026 | Disetujui: 13-08-2026 | Diterbitkan: 15-08-2026

ABSTRACT

This study examines the influence of variables within the Unified Theory of Acceptance and Use of Technology (UTAUT) and the Technology Threat Avoidance Theory (TTAT), integrated through an IS Dual Factor approach, on Generative AI usage intention among users in the Jabodetabek region, addressing a gap in prior research that has tended to examine enabling and inhibiting factors along separate trajectories despite users experiencing these dual pulls simultaneously in practice, and thereby demonstrating the relevance of the IS Dual Factor framework, which treats the two sets of determinants as distinct yet concurrently operating constructs. A quantitative survey method was employed involving 200 respondents domiciled in Jabodetabek who possess knowledge of and/or experience using Generative AI technology, with data collected through an online questionnaire from September 2025 to April 2026 and analyzed using Covariance-Based Structural Equation Modeling (CB-SEM) via IBM AMOS, encompassing Confirmatory Factor Analysis and structural model testing. The results indicate that all six hypotheses were supported, with Performance Expectancy ($\beta = 0.564$) and Facilitating Conditions ($\beta = 0.566$) emerging as the strongest predictors, followed by Effort Expectancy ($\beta = 0.379$) and Social Influence ($\beta = 0.243$), which also exerted significant positive effects, while on the inhibiting side Perceived Threats ($\beta = -0.277$) significantly reduced usage intention and Perceived Avoidability ($\beta = 0.287$) functioned as a buffer attenuating this negative effect, collectively demonstrating that integrating UTAUT and TTAT within an IS Dual Factor framework yields a more comprehensive predictive model than unidimensional approaches for explaining Generative AI usage intention among Indonesia's urban population.

Keywords: Generative AI; UTAUT; TTAT; IS Dual Factor; Usage Intention; Technology Adoption; Jabodetabek; CB-SEM

ABSTRAK

Penelitian ini bertujuan menguji pengaruh variabel-variabel dalam kerangka Unified Theory of Acceptance and Use of Technology (UTAUT) dan Technology Threat Avoidance Theory (TTAT) yang diintegrasikan melalui pendekatan IS Dual Factor terhadap niat penggunaan Generative AI pada pengguna di wilayah Jabodetabek, mengingat penelitian terdahulu mengenai adopsi teknologi AI umumnya mengkaji faktor pendorong dan faktor penghambat secara terpisah padahal dalam praktiknya pengguna mengalami kedua tarikan tersebut secara bersamaan, sehingga penelitian ini hadir untuk mengisi celah tersebut dengan menguji relevansi kerangka IS Dual Factor yang memandang kedua kelompok determinan sebagai entitas berbeda namun bekerja secara simultan. Pendekatan kuantitatif digunakan melalui metode survei terhadap 200 responden berdomisili di Jabodetabek yang memiliki pengetahuan dan/atau pengalaman menggunakan Generative AI, dengan data dikumpulkan melalui kuesioner daring pada periode September 2025 hingga April 2026, kemudian dianalisis menggunakan Covariance-Based Structural Equation Modeling (CB-SEM) dengan bantuan IBM AMOS, meliputi tahap Confirmatory Factor Analysis dan pengujian model struktural. Hasil penelitian menunjukkan bahwa keenam hipotesis diterima, dengan

Performance Expectancy ($\beta = 0,564$) dan Facilitating Conditions ($\beta = 0,566$) sebagai prediktor terkuat, diikuti Effort Expectancy ($\beta = 0,379$) dan Social Influence ($\beta = 0,243$) yang berpengaruh positif signifikan, sementara pada sisi penghambat Perceived Threats ($\beta = -0,277$) terbukti menurunkan niat penggunaan dan Perceived Avoidability ($\beta = 0,287$) berfungsi melemahkan efek negatif tersebut, sehingga secara keseluruhan penelitian ini membuktikan bahwa integrasi UTAUT dan TTAT dalam kerangka IS Dual Factor menghasilkan model prediktif yang lebih komprehensif dibandingkan pendekatan unidimensional dan relevan dalam menjelaskan dinamika niat penggunaan Generative AI pada populasi urban Indonesia.

Katakunci: Generative AI; UTAUT; TTAT; IS Dual Factor; Usage Intention; Adopsi Teknologi; Jabodetabek; CB-SEM

Bagaimana Cara Sitasi Artikel ini:

Laksita, M. R., Noviarini, D. ., & Kasofi, A. (2026). Integrasi UTAUT dan TTAT dalam Kerangka IS Dual Factor pada Niat Penggunaan Teknologi AI di Jabodetabek. Jurnal Ilmu Sosial Dan Humaniora, 2(3), 3381-3401. <https://doi.org/10.63822/t6fvsw15>

PENDAHULUAN

Perkembangan kecerdasan artifisial, khususnya *generative AI*, menunjukkan akselerasi adopsi menuju arus utama dan mendorong perubahan proses kerja, pembelajaran, serta pengambilan keputusan di berbagai sektor. Berdasarkan data Resourcera (2025), pada tahun 2025 terdapat sekitar 378,8 juta pengguna aktif alat AI di seluruh dunia, setara dengan 3,9% populasi global, dan diproyeksikan meningkat 92,4% hingga mencapai 729,1 juta pada tahun 2030. Tren ini menandai pergeseran dari fase eksplorasi menuju pemanfaatan yang semakin terstruktur, baik pada tingkat individu maupun organisasi. Di Indonesia, dorongan mengadopsi AI kian menguat seiring dengan peningkatan ketersediaan alat berbahasa Indonesia, ekosistem pelatihan, dan penguatan tata kelola di sektor publik maupun privat. Namun, kebutuhan untuk memahami bagaimana dan mengapa pengguna berniat serta benar-benar menggunakan (*actual use*) AI generatif tetap mendesak, mengingat teknologi ini membawa manfaat produktivitas sekaligus risiko seperti kebocoran data, halusinasi, bias, dan tantangan kepatuhan.

Secara teoritis, *Unified Theory of Acceptance and Use of Technology* (UTAUT) merupakan kerangka paling dominan untuk menjelaskan penerimaan teknologi. Empat konstruk utamanya—*performance expectancy* (PE), *effort expectancy* (EE), *social influence* (SI), dan *facilitating conditions* (FC)—menjabarkan pendorong (*enablers*) yang memengaruhi niat dan perilaku penggunaan aktual. Dalam konteks AI generatif, PE mencerminkan persepsi bahwa AI dapat meningkatkan kualitas dan kecepatan kerja; EE berkaitan dengan kemudahan mempelajari antarmuka berbasis *prompt*; SI terlihat dari norma dan ekspektasi kolega atau atasan; sedangkan FC merujuk pada ketersediaan kebijakan, pelatihan, dan infrastruktur TI. Keempatnya membentuk jalur *approach/adopt* yang secara empiris terbukti relevan. Namun, mengandalkan UTAUT semata sering kali menghasilkan pemahaman yang bias ke sisi positif, sehingga kurang menangkap kekuatan negatif yang mendorong sikap menolak atau menghindari AI. *Technology Threat Avoidance Theory* (TTAT) melengkapi celah tersebut dengan menempatkan *perceived threats* (persepsi keparahan dan kerentanan terhadap ancaman seperti kebocoran informasi, pelanggaran hak cipta, dan halusinasi) serta *perceived avoidability* (keyakinan bahwa ancaman dapat dihindari melalui kontrol teknis dan kebijakan) sebagai determinan motivasi penghindaran. Pengguna yang menilai ancaman tinggi dan kemampuan mitigasi rendah cenderung menghindari, menurunkan niat dan penggunaan aktual; sebaliknya, ketika *guardrails* seperti pembatasan *prompt* sensitif dan audit kualitas dinilai efektif, motivasi menghindari menurun dan peluang adopsi terbuka lebih lebar.

Literatur terkini memperlihatkan dua arus besar yang berjalan berdampingan. Pada sisi penerimaan, eksperimen acak terkontrol menunjukkan pekerja penulisan profesional yang dibantu *generative AI* meningkatkan produktivitas dan mutu hasil kerja (Noy & Zhang, 2023), sementara studi lapangan berskala besar melaporkan lonjakan produktivitas agen layanan pelanggan sekitar 14–15% (Brynjolfsson, Li, & Raymond, 2023). Kajian berbasis UTAUT pada pendidikan tinggi juga menemukan bahwa ekspektasi kinerja dan kemudahan pakai mendorong niat dan penggunaan aktual *generative AI* pada mahasiswa (Wu, Tian, & Liu, 2025). Di sisi lain, banyak individu tetap membatasi atau menunda penggunaan AI karena kekhawatiran privasi, ketidakakuratan, pelanggaran akademik, dan persepsi ancaman. Temuan pada pengguna ChatGPT memperlihatkan bahwa *privacy concern* dan persepsi risiko menekan kepercayaan serta kesiapan menggunakan AI (Heliyon, 2024), sementara di lingkungan pendidikan tinggi, resistensi pendidik banyak dipicu oleh isu etika, kualitas pembelajaran, dan kesiapan infrastruktur (Kalmus & Nikiforova, 2024). Potret ganda ini menegaskan bahwa pendorong dan penghambat bukanlah sekadar pro dan kontra di satu garis, melainkan dua gugus faktor berbeda yang dapat bekerja secara bersamaan.

Tabel 1. Penelitian Terdahulu

Penulis (Tahun)	Teori / Metode	Variabel Independen	Produk / Layanan	Topik
Noy & Zhang (2023)	Eksperimen acak terkontrol (RCT)	Treatment akses ChatGPT	Generative AI asisten penulisan	Dampak AI pada produktivitas & kualitas penulisan
Brynjolfsson, Li, & Raymond (2025)	Kuasi-eksperimental; Difference-in-Differences	Akses asisten AI; pengalaman agen	Asisten AI untuk customer support	Kenaikan produktivitas agen $\pm 15\%$
Wu, Tian, & Liu (2025)	UTAUT3; SEM (AMOS)	PE, EE, SI, Hedonic Motivation, Price Value	ChatGPT (pendidikan tinggi)	Determinan intensi & perilaku penggunaan mahasiswa
Tehreem et al. (2025)	UTAUT; PLS-SEM & MGA	PE, EE, SI, Learning Value	ChatGPT (pendidikan tinggi)	Pendorong niat adopsi mahasiswa
Kalmus & Nikiforova (2024)	IRT + TOE; kuantitatif & kualitatif	Hambatan organisasional & individu	Generative AI di pendidikan	Resistensi pendidik terhadap adopsi AI

(Sumber: Diolah oleh peneliti, 2025)

Berdasarkan penelusuran tersebut, tampak kesenjangan riset yang relevan. Banyak studi berhenti pada niat menggunakan dan jarang menguji perilaku penggunaan aktual secara komprehensif, terlebih lagi yang memadukan pendorong dan penghambat dalam satu model empiris masih terbatas. Padahal dalam praktiknya, pengguna kerap mengalami tarikan ganda: dorongan memanfaatkan manfaat AI sekaligus dorongan menghindari risiko. Kerangka IS *Dual Factor* menegaskan bahwa pendorong dan penghambat adalah dua set determinan yang berbeda, bukan sekadar kebalikan, dan keduanya memengaruhi perilaku penggunaan.

Sebuah pra-survei yang melibatkan responden dari berbagai latar belakang dilakukan untuk memperdalam pemahaman tentang dinamika adopsi *generative AI*. Hasilnya mengungkapkan persepsi positif yang kuat terhadap manfaat AI, seperti kemudahan pencarian referensi, pembuatan *outline*, analisis data, penyusunan laporan, dan pembuatan konten promosi. Sebagian besar responden mengakui bahwa AI menghemat waktu pengerjaan secara substansial, sehingga meningkatkan efisiensi dan produktivitas kerja. Terkait kemudahan penggunaan, mayoritas responden menilai AI cukup mudah digunakan, terutama bagi mereka yang sudah terbiasa dengan teknologi digital. Pengaruh sosial juga signifikan, terutama dari rekan kerja, dosen, dan media sosial yang telah lebih dulu menggunakan AI dan merekomendasikannya.

Di sisi lain, pra-survei juga mengungkap berbagai kekhawatiran. Ketakutan akan jawaban yang tidak akurat menjadi keraguan utama, diikuti oleh kekhawatiran plagiarisme dan keamanan data, terutama terkait informasi klien dan data perusahaan yang sensitif. Untuk meminimalkan risiko, responden menyarankan pengecekan ulang jawaban AI dengan merujuk pada sumber resmi, menghindari penggunaan data sensitif, serta menerapkan kebijakan internal yang membatasi data yang dibagikan. Terkait motivasi keberlanjutan, penghematan waktu menjadi faktor utama yang membuat responden tertarik untuk terus menggunakan AI, diikuti oleh peningkatan efisiensi dan efektivitas kerja.

Menjawab celah tersebut, penelitian ini mengintegrasikan dua lensa teoretik yang saling melengkapi: UTAUT sebagai kerangka pendorong dan TTAT sebagai kerangka penghambat, yang dijembatani oleh IS *Dual-Factor* sehingga *enablers* dan *inhibitors* dapat diuji bersamaan terhadap intensi dan *actual use*. Pendekatan *dual-factor* sendiri terbukti relevan pada teknologi AI dan layak diadaptasi pada konteks *generative AI* (Dwivedi et al., 2021).

Berdasarkan uraian latar belakang tersebut, penelitian ini mengajukan enam rumusan masalah, yaitu apakah *performance expectancy*, *effort expectancy*, *social influence*, *facilitating conditions*, dan *perceived avoidability* berpengaruh positif terhadap *generative AI usage intention*, serta apakah *perceived threats* berpengaruh negatif terhadap *generative AI usage intention*. Adapun tujuan penelitian ini adalah untuk menganalisis pengaruh masing-masing variabel tersebut terhadap *generative AI usage intention* pada pengguna di kawasan Jabodetabek.

METODE PENELITIAN

Waktu dan Tempat Penelitian

Penelitian ini dilaksanakan di wilayah Jabodetabek (Jakarta, Bogor, Depok, Tangerang, dan Bekasi) selama periode September 2025 hingga April 2026. Pemilihan lokasi ini didasarkan pada pertimbangan bahwa Jabodetabek sebagai kawasan metropolitan terbesar di Indonesia merupakan tempat yang tepat untuk mengamati bagaimana masyarakat menerima dan menggunakan teknologi kecerdasan buatan (AI). Sebagai pusat ekonomi, sosial, dan teknologi, wilayah ini dihuni oleh populasi yang beragam dari berbagai latar belakang, sehingga memungkinkan peneliti untuk memperoleh gambaran yang lebih lengkap tentang pengguna AI. Jakarta, khususnya, berperan sebagai barometer adopsi teknologi baru di Indonesia karena masyarakatnya cenderung lebih cepat dalam mengadopsi inovasi teknologi, baik untuk keperluan pribadi maupun profesional. Secara praktis, Jabodetabek juga memberikan kemudahan akses dalam proses pengumpulan data melalui penyebaran kuesioner, sehingga penelitian dapat berjalan efisien tanpa mengorbankan kualitas dan representativitas data.

Desain Penelitian

Penelitian ini mengadopsi pendekatan kuantitatif dengan metode survei sebagai strategi utamanya. Rachman (2024) menjelaskan bahwa penelitian kuantitatif berfokus pada data berupa angka dan analisis statistik yang memungkinkan pengujian hubungan antarvariabel secara terukur dan sistematis. Melalui pendekatan ini, temuan penelitian dapat digeneralisasi ke populasi yang lebih luas, memberikan pemahaman yang lebih menyeluruh tentang fenomena yang diteliti. Desain survei kuantitatif dipilih karena sangat sesuai untuk menguji model yang mengintegrasikan dua teori besar, yaitu *Unified Theory of Acceptance and Use of Technology* (UTAUT) dan *Technology Threat Avoidance Theory* (TTAT), dalam kerangka *IS Dual Factor*. Model ini melihat adopsi teknologi dari dua sisi: faktor pendorong seperti *performance expectancy*, *effort expectancy*, *social influence*, dan *facilitating conditions*, serta faktor penghambat seperti *perceived threats* dan *perceived avoidability*. Semua faktor ini pada akhirnya membentuk niat dan perilaku penggunaan aktual teknologi *generative AI* di Indonesia.

Populasi dan Sampel

Populasi Penelitian

Populasi dalam penelitian ini adalah seluruh individu berusia 17 tahun hingga 55 tahun ke atas yang tinggal di wilayah Jabodetabek dan memiliki pengetahuan atau pengalaman menggunakan teknologi *generative AI* seperti ChatGPT, Gemini, Copilot, dan aplikasi sejenis. Pemilihan populasi ini sejalan dengan tujuan penelitian untuk memahami dinamika penerimaan dan kekhawatiran terhadap teknologi AI di lingkungan perkotaan yang memiliki akses luas terhadap internet dan layanan digital.

Teknik Pengambilan Sampel

Penelitian ini menggunakan teknik *non-probability sampling* dengan metode *purposive sampling*, yaitu responden dipilih berdasarkan kriteria tertentu yang sesuai dengan tujuan penelitian. Metode ini relevan ketika peneliti membutuhkan responden dengan karakteristik khusus, dalam hal ini pengalaman atau pengetahuan terkait *generative AI*. Kriteria responden yang ditetapkan adalah sebagai berikut: (1) berusia minimal 17 tahun hingga 55 tahun ke atas, mencakup individu yang secara hukum telah dianggap dewasa dan memiliki kapasitas untuk membuat keputusan terkait penggunaan teknologi; (2) berdomisili di wilayah Jabodetabek yang merepresentasikan kawasan urban dengan infrastruktur digital dan penetrasi teknologi yang tinggi; (3) mengetahui atau menggunakan teknologi *generative AI*, misalnya pernah menggunakan atau minimal mengetahui keberadaan aplikasi AI untuk keperluan produktivitas, pembelajaran, hiburan, atau pekerjaan.

Mengacu pada pedoman *Structural Equation Modeling* (SEM) yang membutuhkan ukuran sampel yang memadai untuk menghasilkan estimasi parameter yang stabil, jumlah sampel minimum yang digunakan dalam penelitian ini adalah 200 responden. Pendekatan ini sesuai dengan rekomendasi Kline (2016) yang menyatakan bahwa untuk model SEM, sampel yang lebih besar diperlukan, terutama jika model mengandung beberapa konstruk laten dan indikator. Hal ini juga didukung oleh Westland (2010) yang menggarisbawahi pentingnya menyesuaikan ukuran sampel dengan kompleksitas model yang digunakan. Dengan demikian, target sampel penelitian ini adalah 200 responden. Selain itu, penelitian ini juga mempertimbangkan *margin of error* sebesar 5% yang merupakan standar dalam penelitian sosial kuantitatif sebagai batas toleransi kesalahan (Puisa et al., 2023). Dengan target minimal 200 responden yang tersebar di wilayah Jabodetabek, penelitian ini diharapkan menghasilkan estimasi yang representatif terhadap populasi sasaran.

Definisi Konseptual

Definisi konseptual diperlukan agar setiap variabel dalam model penelitian terbangun secara jelas dan konsisten dengan teori yang mendasarinya. Rachman (2024) menekankan pentingnya kejelasan konstruk dan hubungan antarkonstruk dalam penelitian kuantitatif, karena hal ini akan memengaruhi kualitas instrumen dan ketepatan analisis. Dalam penelitian ini, variabel-variabel dibagi menjadi faktor pendorong (*enablers*) dan faktor penghambat (*inhibitors*) dalam kerangka IS *Dual Factor*, serta variabel niat dan perilaku penggunaan aktual.

Performance expectancy adalah tingkat keyakinan seseorang bahwa menggunakan teknologi *generative AI* akan meningkatkan kinerja mereka dalam hal belajar, bekerja, atau aktivitas sehari-hari. Konstruk ini diadaptasi dari UTAUT (Venkatesh et al., 2003) dan mencerminkan persepsi manfaat fungsional seperti efisiensi waktu, peningkatan kualitas hasil, dan kemudahan menyelesaikan tugas melalui bantuan AI.

Effort expectancy menggambarkan sejauh mana seseorang memandang teknologi *generative AI* sebagai sesuatu yang mudah dipelajari dan dioperasikan. Konstruk ini berkaitan dengan persepsi tingkat usaha dan kompleksitas penggunaan, termasuk kemudahan antarmuka, kejelasan instruksi, dan pengalaman pengguna secara keseluruhan.

Social influence adalah sejauh mana persepsi seseorang untuk menggunakan *generative AI* dipengaruhi oleh pandangan orang-orang penting di sekitarnya, seperti keluarga, teman, rekan kerja, dosen, atau atasan. Jika lingkungan sosial memberikan dorongan positif atau menganggap penggunaan AI

sebagai hal yang wajar dan bermanfaat, seseorang cenderung memiliki niat yang lebih kuat untuk mengadopsinya.

Facilitating conditions merujuk pada sejauh mana seseorang percaya bahwa infrastruktur teknis dan organisasi tersedia untuk mendukung penggunaan *generative AI*. Konstruk ini mencakup ketersediaan sumber daya, pengetahuan, dan dukungan teknis yang diperlukan untuk menggunakan teknologi secara efektif.

Perceived threats adalah persepsi seseorang terhadap potensi risiko dan ancaman yang timbul dari penggunaan *generative AI*, seperti risiko privasi, keamanan data, kesalahan *output*, bias algoritmik, atau kekhawatiran terhadap penggantian peran manusia. Konstruk ini diadaptasi dari TTAT yang memandang ancaman teknologi sebagai salah satu faktor utama yang mendorong perilaku penghindaran.

Perceived avoidability adalah keyakinan seseorang mengenai sejauh mana ancaman yang terkait dengan penggunaan *generative AI* dapat dihindari atau dikendalikan, misalnya melalui pengaturan privasi, verifikasi manual, penggunaan sumber resmi, atau pembatasan jenis tugas yang diserahkan kepada AI. Dalam TTAT, persepsi bahwa ancaman dapat dihindari dapat memengaruhi apakah seseorang memilih untuk tetap menggunakan atau menghindari teknologi tersebut.

Generative AI usage intention mengacu pada kecenderungan dan komitmen psikologis seseorang untuk menggunakan teknologi *generative AI* di masa mendatang. Niat ini tercermin dalam keinginan untuk terus menggunakan, meningkatkan frekuensi penggunaan, dan merekomendasikan *generative AI* kepada orang lain.

Generative AI actual usage adalah tingkat penggunaan nyata teknologi *generative AI* oleh seseorang dalam kurun waktu tertentu. Penggunaan aktual dapat diukur melalui frekuensi, durasi, ragam aktivitas (misalnya penulisan, desain, analisis data), serta konteks penggunaan (pendidikan, pekerjaan, usaha, hiburan).

Pengembangan Instrumen

Instrumen penelitian disusun dalam bentuk kuesioner terstruktur yang memuat sejumlah pernyataan untuk mengukur masing-masing variabel. Penyusunan instrumen mengacu pada langkah-langkah standar pengembangan alat ukur kuantitatif yang menekankan pentingnya validitas dan reliabilitas agar data yang dihasilkan akurat dan dapat dipercaya (Rachman, 2024).

Tabel 2. Instrumen Indikator

Variabel	Sumber	Item Asli	Item Adaptasi
Performance Expectancy	Venkatesh et al. (2003)	I would find the system useful in my job	Saya merasa teknologi AI bermanfaat dalam pekerjaan/kegiatan saya
		Using the system enables me to accomplish tasks more quickly	Menggunakan teknologi AI memungkinkan saya menyelesaikan tugas dengan lebih cepat
		Using the system increases my productivity	Menggunakan teknologi AI meningkatkan produktivitas saya
		If I use the system, I will increase my chances of getting a raise	Jika saya menggunakan teknologi AI, saya akan meningkatkan peluang untuk mendapatkan kenaikan gaji/penghargaan
Effort Expectancy	Venkatesh et al. (2003)	My interaction with the system would be clear and understandable	Interaksi saya dengan teknologi AI jelas dan mudah dipahami

Social Influence	Venkatesh et al. (2003)	et	It would be easy for me to become skillful at using the system	Menjadi terampil dalam menggunakan teknologi AI mudah bagi saya
			I would find the system easy to use	Saya merasa teknologi AI mudah digunakan
			Learning to operate the system is easy for me	Mempelajari cara mengoperasikan teknologi AI mudah bagi saya
			People who influence my behavior think that I should use the system	Orang-orang yang memengaruhi perilaku saya berpikir bahwa saya harus menggunakan teknologi AI
			People who are important to me think that I should use the system	Orang-orang yang penting bagi saya berpikir bahwa saya harus menggunakan teknologi AI
Facilitating Conditions	Venkatesh et al. (2003)	et	The senior management of this business has been helpful in the use of the system	Senior di organisasi/institusi saya mendukung penggunaan teknologi AI
			In general, the organization has supported the use of the system	Secara umum, organisasi/institusi saya mendukung penggunaan teknologi AI
			I have the resources necessary to use the system	Saya memiliki sumber daya yang diperlukan untuk menggunakan teknologi AI
			I have the knowledge necessary to use the system	Saya memiliki pengetahuan yang diperlukan untuk menggunakan teknologi AI
			The system is not compatible with other systems I use	Teknologi AI kompatibel dengan sistem lain yang saya gunakan
Perceived Threats	Liang & Xue (2009)	Xue	A specific person (or group) is available for assistance with system difficulties	Terdapat individu atau kelompok tertentu yang siap membantu ketika saya mengalami kesulitan dalam menggunakan teknologi AI
			My fear of exposure to Generative AI's risks is high	Ketakutan saya terhadap risiko penggunaan teknologi AI sangat tinggi
			The extent of my anxiety about potential loss due to Generative AI's risks is high	Tingkat kecemasan saya terhadap potensi kerugian akibat risiko teknologi AI sangat tinggi
			The extent of my worry about Generative AI's risks due to misuse is high	Tingkat kekhawatiran saya terhadap risiko teknologi AI akibat penyalahgunaan sangat tinggi
			Taking everything into consideration, the threat of generative AI could be prevented	Dengan mempertimbangkan segala aspek, ancaman teknologi AI dapat dicegah
Perceived Avoidability	Liang & Xue (2009)	Xue	Taking everything into consideration, I could protect myself from the threat of generative AI	Dengan mempertimbangkan segala aspek, saya dapat melindungi diri dari ancaman teknologi AI
			I intend to use Generative AI in the future	Saya berniat menggunakan teknologi AI di masa depan
			I plan to use Generative AI in the future	Saya berencana menggunakan teknologi AI di masa depan
Generative AI Usage Intention	Venkatesh et al. (2012)	et		

(Sumber: Diolah oleh peneliti, 2025)

Teknik Pengumpulan Data

Data yang digunakan dalam penelitian ini adalah data primer, yaitu data yang dikumpulkan langsung oleh peneliti untuk menjawab pertanyaan penelitian secara spesifik. Cerar et al. (2021) menjelaskan bahwa data primer memungkinkan peneliti memperoleh informasi yang relevan dengan tujuan studi dan dapat dikontrol sesuai kebutuhan analisis.

Kuesioner disusun menggunakan Google Form dan dibagi menjadi dua bagian utama: bagian pertama memuat pertanyaan demografis (usia, jenis kelamin, tingkat pendidikan, pekerjaan, dan frekuensi penggunaan *generative AI*), sedangkan bagian kedua berisi pernyataan untuk mengukur delapan variabel penelitian (PE, EE, SI, FC, PT, PA, UI, AU). Seluruh butir pernyataan diukur menggunakan skala Likert enam poin dengan rentang skor 1 (Sangat Tidak Setuju) hingga 6 (Sangat Setuju). Penggunaan skala Likert enam poin dipilih karena mampu mendorong responden untuk memberikan sikap yang tegas setuju atau tidak setuju tanpa mengambil posisi netral, sehingga menghasilkan data yang lebih jelas dan mengurangi ambiguitas dalam mengukur persepsi serta niat penggunaan teknologi *generative AI*.

Kuesioner disebarakan secara daring melalui berbagai kanal digital seperti WhatsApp, Telegram, Instagram, X, dan komunitas mahasiswa untuk menjangkau responden yang relevan dengan karakteristik populasi. Pengumpulan data dilengkapi dengan pernyataan persetujuan (*informed consent*) mengenai kerahasiaan dan penggunaan data, serta pengaturan "satu respon per akun" untuk meminimalkan duplikasi. Data yang tidak lengkap atau berisi pola jawaban tidak konsisten akan dieliminasi pada tahap prapengolahan data. Pengumpulan data direncanakan berlangsung selama bulan November 2025 hingga Januari 2026 dengan target minimal 200 responden sesuai hasil perhitungan jumlah sampel pada bagian sebelumnya.

Teknik Analisis Data

Teknik analisis data yang digunakan dalam penelitian ini adalah *Covariance-Based Structural Equation Modeling* (CB-SEM) menggunakan software AMOS, yang terdiri dari dua tahap utama: *Confirmatory Factor Analysis* (CFA) dan *Structural Equation Modeling* (SEM).

Confirmatory Factor Analysis (CFA)

Langkah pertama dalam analisis adalah CFA, yang digunakan untuk menguji validitas konstruk dalam model yang diajukan. CFA bertujuan untuk memastikan bahwa indikator yang digunakan benar-benar mengukur variabel laten yang dimaksud, serta untuk mengonfirmasi struktur faktor yang diusulkan dalam model penelitian. Dalam CFA, dilakukan pengujian validitas konvergen dan diskriminan untuk memastikan bahwa indikator-indikator pada setiap konstruk saling berkorelasi dengan baik dan tidak tumpang tindih antar konstruk yang berbeda (Mia et al., 2019).

Tabel 3. Rule of Thumb Goodness of Fit

Rule of Thumb	Cut-off Value
Goodness of Fit Index (GFI)	> 0,90
Adjusted Goodness of Fit Index (AGFI)	> 0,90
Normed Fit Index (NFI)	> 0,95
Comparative Fit Index (CFI)	> 0,95
Root Mean Square Residual (RMR)	< 0,05
Root Mean Square Error of Approximation (RMSEA)	< 0,05
Chi-Square	> 0,05

(Sumber: Mia et al., 2019)

Pada tahap ini, reliabilitas konstruk juga diuji menggunakan *Composite Reliability* (CR) dan *Average Variance Extracted* (AVE). Nilai CR > 0,7 dan AVE > 0,5 menunjukkan bahwa konstruk yang diuji memiliki reliabilitas dan validitas yang memadai.

Structural Equation Modeling (SEM)

Setelah model konstruk diuji pada tahap CFA, tahap selanjutnya adalah penerapan SEM untuk menguji hubungan antar variabel laten dalam model. SEM memungkinkan pengujian hubungan yang lebih kompleks antar konstruk, termasuk pengaruh langsung dan tidak langsung, serta memungkinkan peneliti untuk menguji keseluruhan model struktural dengan mempertimbangkan interaksi antar variabel secara simultan. Pada tahap SEM, pengujian hipotesis dilakukan untuk menilai apakah hubungan antara variabel dalam model teoritis didukung oleh data empiris. Indikator *goodness of fit* yang digunakan pada tahap SEM hampir sama dengan CFA, namun pada SEM, *chi-square* dan RMSEA menjadi lebih sensitif terhadap perubahan dalam model yang lebih kompleks. Reliabilitas model SEM juga diuji dengan menggunakan CR dan AVE yang sama dengan CFA, yang harus menunjukkan nilai CR > 0,7 dan AVE > 0,5 untuk memastikan bahwa model SEM memiliki validitas yang baik. Pengujian hipotesis dalam SEM dilakukan dengan melihat nilai *Critical Ratio* (CR) dan *p-value*, di mana hipotesis diterima jika nilai CR > 1,96 dan *p-value* < 0,05.

HASIL DAN PEMBAHASAN

Deskripsi Data

Pengumpulan data dalam penelitian ini dilakukan melalui penyebaran kuesioner daring menggunakan Google Form kepada responden yang berdomisili di wilayah Jabodetabek dan memiliki pengetahuan atau pengalaman menggunakan teknologi *generative AI*. Pengumpulan data dilaksanakan pada bulan September 2025 hingga April 2026. Sebanyak 223 kuesioner disebarkan melalui berbagai kanal digital, termasuk WhatsApp, Telegram, Instagram, dan komunitas mahasiswa. Dari jumlah tersebut, sebanyak 200 kuesioner dinyatakan valid dan memenuhi syarat untuk dianalisis lebih lanjut.

Profil Responden

Profil demografis responden penelitian mencakup karakteristik jenis kelamin, usia, pendidikan terakhir, pekerjaan, dan domisili. Berdasarkan hasil analisis, responden dalam penelitian ini didominasi oleh perempuan sebanyak 113 orang (56,5%), sementara responden laki-laki berjumlah 87 orang (43,5%). Komposisi ini mencerminkan keterwakilan kedua gender yang relatif berimbang dalam sampel penelitian. Tingginya proporsi responden perempuan sejalan dengan tren terkini adopsi *generative AI*, di mana tingkat adopsi perempuan terhadap AI mengalami pertumbuhan yang pesat dan mulai menyamai laki-laki, terutama di kalangan profesional muda (Deloitte, 2025).

Dilihat dari usia, sebagian besar responden berada pada kelompok usia 26-35 tahun (44,5%) yang mengindikasikan bahwa sampel penelitian ini didominasi oleh generasi muda yang secara umum lebih akrab dengan teknologi digital, termasuk *generative AI*. Temuan ini konsisten dengan laporan Carolan et al. (2025) yang menemukan bahwa *Millennials* (usia 26-44 tahun) merupakan *power user generative AI* dengan tingkat penggunaan harian tertinggi dibandingkan generasi lainnya.

Dari segi pendidikan terakhir, mayoritas responden berpendidikan Sarjana/S1 (45%). Tingginya proporsi responden berpendidikan sarjana mencerminkan populasi target penelitian yang merupakan pengguna teknologi dengan latar belakang pendidikan memadai untuk memahami dan memanfaatkan *generative AI*. Dari segi pekerjaan, responden terbanyak adalah karyawan swasta (39,5%), konsisten dengan temuan Bick et al. (2025) bahwa karyawan sektor swasta menunjukkan tingkat adopsi AI yang lebih tinggi dibandingkan sektor publik karena tekanan kompetitif dan tuntutan produktivitas yang lebih tinggi. Dari segi domisili, responden tersebar di seluruh wilayah Jabodetabek dengan Jakarta sebagai daerah asal terbanyak (26,5%).

Statistik Deskriptif Variabel Penelitian

Tabel 4. Statistik Deskriptif Variabel Penelitian

Variabel	Rata-rata Total	Kategori
Performance Expectancy (PE)	5,05	Tinggi
Effort Expectancy (EE)	4,53	Tinggi
Social Influence (SI)	4,45	Tinggi
Facilitating Conditions (FC)	4,66	Tinggi
Perceived Threats (PT)	3,88	Cukup
Perceived Avoidability (PA)	4,01	Cukup-Tinggi
Usage Intention (UI)	4,55	Tinggi

(Sumber: Diolah oleh peneliti, 2026)

Berdasarkan Tabel 4, seluruh konstruk UTAUT (*Performance Expectancy*, *Effort Expectancy*, *Social Influence*, dan *Facilitating Conditions*) berada pada kategori tinggi dengan rata-rata total di atas 4,45. *Performance Expectancy* mencatatkan nilai tertinggi ($M = 5,05$), yang mengindikasikan bahwa responden memiliki ekspektasi kinerja yang sangat tinggi terhadap penggunaan teknologi AI. Hal ini sejalan dengan temuan Noy & Zhang (2023) bahwa persepsi manfaat produktivitas merupakan pendorong dominan adopsi AI.

Nilai rata-rata total *Perceived Threats* sebesar 3,88 berada pada kategori cukup, mengindikasikan tingkat kekhawatiran yang moderat terhadap risiko penggunaan AI. Sementara itu, *Perceived Avoidability* dengan rata-rata 4,01 berada pada kategori cukup hingga tinggi, mengonfirmasi peran konstruk ini sebagai penyeimbang (*buffer*) terhadap *Perceived Threats* sebagaimana diprediksikan oleh TTAT. *Usage Intention* dengan rata-rata 4,55 berada pada kategori tinggi, yang mengindikasikan bahwa niat penggunaan *generative AI* di kalangan responden Jabodetabek tergolong kuat.

Hasil Uji First Order - Confirmatory Factor Analysis (CFA)

Tahap pertama analisis data menggunakan pendekatan *First Order - Confirmatory Factor Analysis* (CFA) yang bertujuan untuk menguji validitas konstruk (*construct validity*) dan reliabilitas konstruk (*construct reliability*) dari setiap variabel laten dalam model penelitian. FO-CFA dilakukan menggunakan perangkat lunak IBM AMOS versi 26 dengan estimasi parameter menggunakan metode *Maximum Likelihood Estimation* (MLE).

Evaluasi Goodness of Fit Model CFA

Tabel 5. Hasil Uji Goodness of Fit Model CFA

Indikator	Nama Indikator	Nilai	Cut-off	Status
CMIN/DF	Chi-Square/Derajat Kebebasan	1,014	$\leq 3,0$	Good Fit
GFI	Goodness of Fit Index	0,914	$> 0,90$	Good Fit

Integrasi UTAUT dan TTAT dalam Kerangka IS Dual Factor pada Niat Penggunaan Teknologi AI di Jabodetabek
(Laksita, et al.)

AGFI	Adjusted GFI	0,888	> 0,90	Marginal
RMR	Root Mean Square Residual	0,033	< 0,05	Good Fit
RMSEA	Root Mean Square Error of Approximation	0,008	< 0,05	Good Fit
PCLOSE	P-value of Close Fit	1,000	> 0,05	Good Fit
NFI	Normed Fit Index	0,936	> 0,95	Good Fit
RFI	Relative Fit Index	0,924	> 0,90	Good Fit
IFI	Incremental Fit Index	0,999	> 0,90	Good Fit
TLI	Tucker-Lewis Index	0,999	> 0,95	Good Fit
CFI	Comparative Fit Index	0,999	> 0,95	Good Fit
AIC	Akaike Information Criterion	372,203	< Saturated (600,000)	Good Fit

(Sumber: Diolah oleh peneliti, 2026)

Tabel 5 menunjukkan bahwa hampir seluruh indikator *goodness of fit* memenuhi kriteria yang ditetapkan. Nilai CMIN/DF = 1,014 berada jauh di bawah batas 3,0. GFI = 0,914 dan RMR = 0,033 memenuhi *cut-off*. Indikator *incremental fit* sangat memuaskan: CFI = 0,999, TLI = 0,999, dan IFI = 0,999 seluruhnya melampaui ambang ideal $\geq 0,95$. Nilai RMSEA = 0,008 mengindikasikan *close fit*, dikonfirmasi oleh PCLOSE = 1,000. Nilai AGFI = 0,888 sedikit di bawah *cut-off* ideal 0,90, namun masih dapat dikategorikan sebagai *marginal fit* yang *acceptable*. Baumgartner & Homburg (1996) menegaskan bahwa nilai AGFI di atas 0,80 masih memenuhi syarat kelayakan dalam penelitian aplikatif. Secara keseluruhan, model FO-CFA dinyatakan *fit*.

Uji Validitas Konvergen

Tabel 6. Hasil Uji Validitas Konvergen

Konstruk	Jumlah Item	AVE	CR	$\sqrt{AVE}$	Keterangan
Performance Expectancy (PE)	4	0,766	0,929	0,875	Valid & Reliabel
Effort Expectancy (EE)	4	0,740	0,919	0,860	Valid & Reliabel
Social Influence (SI)	4	0,744	0,921	0,863	Valid & Reliabel
Facilitating Conditions (FC)	4	0,745	0,921	0,863	Valid & Reliabel
Perceived Threats (PT)	3	0,712	0,881	0,844	Valid & Reliabel
Perceived Avoidability (PA)	3	0,740	0,895	0,860	Valid & Reliabel
Usage Intention (UI)	2	0,784	0,879	0,885	Valid & Reliabel

(Sumber: Diolah oleh peneliti, 2026)

Seluruh konstruk memenuhi kriteria validitas konvergen. Nilai AVE berkisar antara 0,712 (PT) hingga 0,784 (UI), seluruhnya $\geq 0,50$. Nilai CR berkisar antara 0,879 (UI) hingga 0,929 (PE), seluruhnya $\geq 0,70$. Semua konstruk terbukti memiliki *internal consistency* dan validitas konvergen yang kuat (Fornell & Larcker, 1981). Seluruh 24 indikator menunjukkan nilai *standardized factor loading* (λ) jauh di atas batas minimum 0,60 dengan nilai terendah FC4 = 0,825 dan tertinggi UI2 = 0,910, sehingga tidak ada indikator yang perlu dieliminasi.

Hasil Uji Structural Equation Modeling (SEM)

Setelah model pengukuran FO-CFA dinyatakan valid dan reliabel, dilakukan pengujian model struktural menggunakan SEM yang menguji pengaruh langsung variabel eksogen (PE, EE, SI, FC, PT, PA) terhadap variabel endogen *Usage Intention* secara simultan.

Evaluasi Goodness of Fit Model SEM

Tabel 7. Hasil Uji Goodness of Fit Model SEM

Indikator	Nama Indikator	Nilai	Cut-off	Status
CMIN/DF	Chi-Square/Derajat Kebebasan	1,625	$\leq 3,0$	Good Fit
GFI	Goodness of Fit Index	0,861	$> 0,90$	Marginal
AGFI	Adjusted GFI	0,832	$> 0,90$	Marginal
RMR	Root Mean Square Residual	0,115	$< 0,05$	Poor Fit
RMSEA	Root Mean Square Error of Approximation	0,056	$< 0,05$	Acceptable
PCLOSE	P-value of Close Fit	0,157	$> 0,05$	Good Fit
NFI	Normed Fit Index	0,891	$> 0,95$	Marginal
RFI	Relative Fit Index	0,878	$> 0,90$	Marginal
IFI	Incremental Fit Index	0,955	$> 0,90$	Good Fit
TLI	Tucker-Lewis Index	0,949	$> 0,95$	Acceptable
CFI	Comparative Fit Index	1,000	$> 0,95$	Good Fit
AIC	Akaike Information Criterion	507,434	$< \text{Saturated (600,000)}$	Good Fit

(Sumber: Diolah oleh peneliti, 2026)

Tabel 7 menunjukkan bahwa model SEM secara keseluruhan *fit*. Nilai CMIN/DF = 1,625 memenuhi kriteria $\leq 3,0$. CFI = 1,000 dan IFI = 0,955 keduanya memenuhi kriteria $\geq 0,95$. Nilai RMSEA = 0,056 sedikit melampaui batas ideal $< 0,05$, namun masih *acceptable* ($< 0,08$) berdasarkan Browne & Cudeck (1992), dan dikonfirmasi oleh PCLOSE = 0,157 ($> 0,05$) yang menunjukkan model tidak berbeda signifikan dari populasi. Nilai GFI = 0,861 dan AGFI = 0,832 berada di bawah *cut-off* ideal 0,90 dan nilai RMR = 0,115 melampaui batas $< 0,05$. Hal ini lazim pada model SEM yang lebih kompleks karena penambahan enam jalur struktural. Hair et al. (2010) menegaskan bahwa pada model yang kompleks, kombinasi CFI, RMSEA, dan CMIN/DF merupakan indikator utama yang diutamakan; ketiganya memenuhi kriteria sehingga model SEM dinyatakan layak untuk pengujian hipotesis.

Hasil Uji Hipotesis

Pengujian hipotesis dilakukan dengan melihat nilai *Critical Ratio* (CR) dan *p-value* pada tabel *Regression Weights* output AMOS. Sesuai ketentuan, hipotesis diterima apabila CR $> 1,96$ dan *p-value* $< 0,05$.

Tabel 8. Hasil Uji Hipotesis

Hipotesis	Path Hubungan	Arah	Estimate	S.E.	C.R.	p-value	Std. β	Keputusan
H1	PE $\rightarrow$ Usage Intention	(+)	0,309	0,042	7,335	0,000	0,564	Diterima
H2	EE $\rightarrow$ Usage Intention	(+)	0,210	0,040	5,227	0,000	0,379	Diterima
H3	SI $\rightarrow$ Usage Intention	(+)	0,123	0,035	3,537	0,000	0,243	Diterima
H4	FC $\rightarrow$ Usage Intention	(+)	0,303	0,041	7,407	0,000	0,566	Diterima
H5	PT $\rightarrow$ Usage Intention	(-)	-0,145	0,037	-3,882	0,000	-0,277	Diterima
H6	PA $\rightarrow$ Usage Intention	(+)	0,150	0,037	4,082	0,000	0,287	Diterima

(Sumber: Diolah oleh peneliti, 2026)

Berdasarkan Tabel 8, seluruh enam hipotesis penelitian diterima. Seluruh nilai CR melampaui nilai kritis 1,96 dengan tingkat signifikansi $p < 0,001$. Berdasarkan besarnya pengaruh (Std. β), urutan konstruk prediktor adalah: FC ($\beta = 0,566$) $>$ PE ($\beta = 0,564$) $>$ EE ($\beta = 0,379$) $>$ PA ($\beta = 0,287$) $>$ PT ($\beta = -0,277$) $>$ SI ($\beta = 0,243$). *Facilitating Conditions* dan *Performance Expectancy* secara bersama-sama merupakan dua prediktor terkuat *generative AI usage intention*.

Pembahasan

Pengaruh Performance Expectancy terhadap Generative AI Usage Intention

Hasil pengujian H1 menunjukkan bahwa *Performance Expectancy* (PE) berpengaruh positif dan signifikan terhadap *Generative AI Usage Intention* (Estimate = 0,309; CR = 7,335; $p < 0,001$; Std. $\beta = 0,564$). Artinya, setiap kenaikan satu poin persepsi responden terhadap *Performance Expectancy* akan meningkatkan *Usage Intention* sebesar 0,309 poin. Dengan nilai $\beta = 0,564$, PE merupakan salah satu dari dua prediktor terkuat dalam model, konsisten dengan proposisi utama UTAUT (Venkatesh et al., 2003) bahwa keyakinan terhadap manfaat kinerja merupakan pendorong dominan niat penggunaan teknologi. Temuan ini sejalan dengan Cabero-Almenara et al. (2024) dan Kim & Blazquez (2024). Kuatnya pengaruh PE juga sejalan dengan pola adopsi *generative AI* secara global yang bergerak cepat justru karena persepsi manfaat produktivitas yang langsung dirasakan pengguna sejak interaksi pertama, sebagaimana dicatat oleh Bick et al. (2025). Dalam konteks Jabodetabek, dominasi PE dapat dimaknai bahwa pengguna cenderung mengevaluasi *generative AI* secara instrumental, yaitu sejauh mana teknologi ini mempercepat pekerjaan dan meningkatkan kualitas *output* sebelum mempertimbangkan aspek lain seperti kemudahan penggunaan atau pengaruh sosial.

Pengaruh Effort Expectancy terhadap Generative AI Usage Intention

Hasil pengujian H2 menunjukkan bahwa *Effort Expectancy* (EE) berpengaruh positif dan signifikan terhadap *Generative AI Usage Intention* (Estimate = 0,210; CR = 5,227; $p < 0,001$; Std. $\beta = 0,379$). Setiap kenaikan satu poin persepsi kemudahan penggunaan akan meningkatkan *Usage Intention* sebesar 0,210 poin. Antarmuka *generative AI* yang berbasis percakapan alami terbukti menurunkan hambatan kognitif penggunaannya, konsisten dengan Mishra et al. (2023) dan Kim & Blazquez (2024). Meski signifikan, kontribusi EE ($\beta = 0,379$) berada di bawah PE dan FC, yang mengindikasikan bahwa bagi pengguna di Jabodetabek, kemudahan penggunaan berperan sebagai faktor pendukung, bukan penentu utama niat penggunaan. Pola ini juga ditemukan Venkatesh et al. (2012) dalam pengembangan UTAUT, di mana *effort expectancy* cenderung lebih dominan pada tahap awal pengenalan teknologi dan melemah seiring meningkatnya pengalaman pengguna. Mengingat mayoritas responden penelitian ini telah memiliki pengetahuan atau pengalaman menggunakan *generative AI*, wajar apabila persepsi kemudahan tidak lagi menjadi pertimbangan dominan dibandingkan ekspektasi manfaat kinerjanya.

Pengaruh Social Influence terhadap Generative AI Usage Intention

Hasil pengujian H3 menunjukkan bahwa *Social Influence* (SI) berpengaruh positif dan signifikan terhadap *Generative AI Usage Intention* (Estimate = 0,123; CR = 3,537; $p < 0,001$; Std. $\beta = 0,243$). Setiap kenaikan satu poin pengaruh sosial yang dirasakan responden akan meningkatkan *Usage Intention* sebesar 0,123 poin. Meski merupakan prediktor dengan pengaruh terkecil, norma sosial tetap bermakna secara statistik dan konsisten dengan Wang et al. (2024) dan Zhai et al. (2021). Kecilnya pengaruh SI dibandingkan PE dan FC mengindikasikan bahwa keputusan menggunakan *generative AI* di kalangan responden lebih bersifat individual dan berbasis evaluasi manfaat pribadi ketimbang dorongan lingkungan sosial atau rekan kerja. Hal ini berbeda dengan konteks pendidikan yang diteliti Wah & Hashim (2021), di mana norma sosial dari sesama pengajar cenderung berperan lebih besar. Perbedaan ini masuk akal mengingat *generative AI* saat ini sudah menjadi teknologi yang mudah diakses secara mandiri, sehingga tekanan atau ajakan dari lingkungan sosial bukan lagi prasyarat utama untuk mencoba dan mengadopsinya.

Pengaruh Facilitating Conditions terhadap Generative AI Usage Intention

Hasil pengujian H4 menunjukkan bahwa *Facilitating Conditions* (FC) berpengaruh positif dan signifikan terhadap *Generative AI Usage Intention* (Estimate = 0,303; CR = 7,407; $p < 0,001$; Std. $\beta = 0,566$). Setiap kenaikan satu poin persepsi ketersediaan fasilitas dan dukungan akan meningkatkan *Usage Intention* sebesar 0,303 poin. Dengan $\beta = 0,566$ dan CR tertinggi di antara seluruh jalur, FC merupakan prediktor terkuat dalam model ini, konsisten dengan Yuniawan et al. (2025). Dominannya FC menunjukkan bahwa ketersediaan infrastruktur pendukung seperti akses internet stabil, perangkat memadai, dan dukungan organisasi atau institusi menjadi prasyarat mendasar sebelum faktor psikologis lain seperti PE atau EE dapat berperan optimal, sejalan dengan temuan Qiao et al. (2025) pada adopsi AI oleh tenaga pengajar klinis yang juga menempatkan *facilitating conditions* sebagai penentu utama. Pola ini relevan dengan karakteristik wilayah Jabodetabek yang meskipun memiliki penetrasi internet tinggi, tetap menunjukkan kesenjangan akses fasilitas dan literasi digital antarkelompok pengguna.

Pengaruh Perceived Threats terhadap Generative AI Usage Intention

Hasil pengujian H5 menunjukkan bahwa *Perceived Threats* (PT) berpengaruh negatif dan signifikan terhadap *Generative AI Usage Intention* (Estimate = -0,145; CR = -3,882; $p < 0,001$; Std. $\beta = -0,277$). Setiap kenaikan satu poin persepsi ancaman responden terhadap *generative AI* akan menurunkan *Usage Intention* sebesar 0,145 poin. Koefisien negatif ini mengonfirmasi prediksi TTAT (Liang & Xue, 2009) bahwa ancaman yang dirasakan secara signifikan menghambat niat penggunaan teknologi, konsisten dengan Chen & Li (2025) dan Shi et al. (2025). Signifikannya PT mengindikasikan bahwa kekhawatiran terhadap risiko *generative AI* seperti kebocoran data, ketergantungan berlebihan, atau kesalahan *output* tetap menjadi pertimbangan nyata meskipun manfaat kinerjanya diakui secara luas, sejalan dengan temuan Alsharida & Al-sharafi (2025) yang menunjukkan bahwa persepsi ancaman keamanan digital secara konsisten menurunkan niat penggunaan teknologi berbasis AI. Implikasinya, upaya mendorong adopsi *generative AI* tidak cukup hanya menonjolkan manfaat, tetapi juga perlu disertai transparansi kebijakan privasi dan edukasi mitigasi risiko agar persepsi ancaman pengguna dapat ditekan.

Pengaruh Perceived Avoidability terhadap Generative AI Usage Intention

Hasil pengujian H6 menunjukkan bahwa *Perceived Avoidability* (PA) berpengaruh positif dan signifikan terhadap *Generative AI Usage Intention* (Estimate = 0,150; CR = 4,082; $p < 0,001$; Std. $\beta = 0,287$). Setiap kenaikan satu poin keyakinan responden bahwa risiko AI dapat dikelola akan meningkatkan *Usage Intention* sebesar 0,150 poin. Temuan ini memperluas proposisi TTAT ke konteks *generative AI* dan menegaskan peran PA sebagai *buffer* yang menetralkan efek negatif PT, konsisten dengan Wang et al. (2024). Diterimanya H6 memperkuat argumen bahwa dalam kerangka TTAT, persepsi kemampuan untuk menghindari atau mengelola risiko (*avoidability*) merupakan mekanisme *coping* yang berdiri sendiri, bukan sekadar kebalikan dari persepsi ancaman, sejalan dengan temuan Liu et al. (2025) pada perluasan *Acceptance-Avoidance Framework* di konteks pengguna *generative AI* di Tiongkok. Artinya, pengguna yang merasa memiliki kendali atas risiko, misalnya melalui verifikasi manual hasil AI atau pembatasan data sensitif yang dibagikan, tetap berniat menggunakan teknologi ini meski menyadari adanya ancaman. Temuan ini memberi implikasi praktis bahwa penguatan efikasi mitigasi risiko pengguna (*coping efficacy*), bukan hanya penurunan ancaman itu

sendiri, menjadi strategi kunci untuk meningkatkan niat penggunaan *generative AI*.

Diskusi Integratif: IS Dual Factor dalam Konteks Generative AI

Secara keseluruhan, temuan penelitian ini memberikan bukti empiris yang kuat bahwa kerangka IS *Dual Factor* (Cenfetelli, 2004; Dwivedi et al., 2023) relevan dan dapat diterapkan dalam konteks adopsi *generative AI* di wilayah Jabodetabek. Model penelitian yang mengintegrasikan UTAUT sebagai jalur *enablers* dan TTAT sebagai jalur *inhibitors* terbukti mampu menjelaskan *generative AI usage intention* secara komprehensif dari kedua perspektif sekaligus, suatu pendekatan yang dinyatakan Dwivedi et al. (2023) sebagai lebih akurat dibandingkan model unidimensional yang hanya melihat satu sisi.

Dari sisi *enablers*, keempat konstruk UTAUT seluruhnya berpengaruh positif signifikan dengan FC ($\beta = 0,566$) dan PE ($\beta = 0,564$) sebagai dua prediktor terkuat, diikuti EE ($\beta = 0,379$) dan SI ($\beta = 0,243$). Pola ini mencerminkan bahwa di konteks Jabodetabek, infrastruktur dan persepsi manfaat kinerja merupakan dua faktor paling determinan. Temuan ini sejalan dengan Yuniawan et al. (2025) yang menemukan FC dominan pada adopsi AI di perbankan Indonesia, serta memperluas temuan Kim & Blazquez (2024) yang membuktikan peran PE di lingkungan korporat Korea ke konteks Indonesia. Dari sisi *inhibitors*, PT ($\beta = -0,277$) terbukti berperan sebagai penghambat yang nyata, mengonfirmasi relevansi TTAT dalam konteks *generative AI*. Menariknya, PA ($\beta = 0,287$) berfungsi sebagai *buffer* yang melemahkan efek negatif persepsi ancaman. Interaksi PT-PA ini selaras dengan temuan Shi et al. (2025) pada pengguna ChatGPT di industri pariwisata yang juga menemukan *perceived threats* menurunkan niat penggunaan sementara kemampuan mitigasi meningkatkannya. Temuan ini secara keseluruhan menegaskan bahwa strategi adopsi *generative AI* yang efektif harus bersifat *dual*: memperkuat infrastruktur dan persepsi manfaat sekaligus menurunkan persepsi ancaman dan meningkatkan *coping efficacy* pengguna, persis seperti yang diprediksikan oleh kerangka IS *Dual Factor*.

KESIMPULAN

Penelitian ini bertujuan untuk menguji pengaruh variabel-variabel dalam kerangka *Unified Theory of Acceptance and Use of Technology* (UTAUT) dan *Technology Threat Avoidance Theory* (TTAT) yang diintegrasikan melalui pendekatan IS *Dual Factor* terhadap *Generative AI Usage Intention* pada pengguna di wilayah Jabodetabek. Penelitian dilakukan terhadap 200 responden yang berdomisili di Jabodetabek dan memiliki pengetahuan dan/atau pengalaman menggunakan teknologi *generative AI*, dengan pengumpulan data melalui kuesioner daring pada periode September 2025 hingga April 2026. Analisis data menggunakan pendekatan *Covariance-Based Structural Equation Modeling* (CB-SEM) melalui perangkat lunak IBM AMOS, mencakup tahap *Confirmatory Factor Analysis* (CFA) dan pengujian model struktural. Berdasarkan hasil analisis yang telah diuraikan pada Bab IV, berikut adalah kesimpulan yang dapat ditarik.

Pengujian hipotesis pertama mengonfirmasi bahwa *Performance Expectancy* (PE) berpengaruh positif dan signifikan terhadap *Generative AI Usage Intention* pada pengguna teknologi *generative AI* di wilayah Jabodetabek, dengan nilai koefisien jalur terstandarisasi sebesar 0,564 dan *Critical Ratio* sebesar 7,335 ($p < 0,001$). Temuan ini mengindikasikan bahwa semakin kuat keyakinan pengguna bahwa *generative AI* mampu meningkatkan kecepatan kerja, produktivitas, dan peluang pengembangan karier mereka, semakin besar pula niat mereka untuk menggunakan teknologi tersebut secara berkelanjutan. Dengan demikian, hipotesis pertama diterima.

Hasil pengujian hipotesis kedua menunjukkan bahwa *Effort Expectancy* (EE) berpengaruh positif dan signifikan terhadap *Generative AI Usage Intention*, dengan nilai koefisien jalur terstandarisasi sebesar 0,379 dan *Critical Ratio* sebesar 5,227 ($p < 0,001$). Temuan ini mencerminkan bahwa kemudahan berinteraksi dan mengoperasikan *generative AI* berkontribusi dalam membentuk niat pengguna, khususnya karena antarmuka berbasis percakapan alami menurunkan hambatan kognitif dalam mempelajari teknologi baru. Dengan demikian, hipotesis kedua diterima.

Pengujian atas hipotesis ketiga membuktikan bahwa *Social Influence* (SI) berpengaruh positif dan signifikan terhadap *Generative AI Usage Intention*, dengan nilai koefisien jalur terstandarisasi sebesar 0,243 dan *Critical Ratio* sebesar 3,537 ($p < 0,001$). Temuan ini mengindikasikan bahwa dorongan dari rekan kerja, dosen, maupun pimpinan institusi turut membentuk niat pengguna untuk mengadopsi *generative AI*, meskipun kontribusinya relatif lebih kecil dibandingkan konstruk UTAUT lainnya. Dengan demikian, hipotesis ketiga diterima.

Hasil pengujian hipotesis keempat menunjukkan bahwa *Facilitating Conditions* (FC) berpengaruh positif dan signifikan terhadap *Generative AI Usage Intention*, dengan nilai koefisien jalur terstandarisasi sebesar 0,566 dan *Critical Ratio* sebesar 7,407 ($p < 0,001$), menjadikannya prediktor dengan pengaruh terbesar dalam model penelitian ini. Temuan ini mengindikasikan bahwa ketersediaan infrastruktur, akses internet, pengetahuan, serta dukungan teknis berkontribusi besar dalam membentuk kesiapan dan niat pengguna untuk memanfaatkan *generative AI* di wilayah Jabodetabek. Dengan demikian, hipotesis keempat diterima.

Pengujian hipotesis kelima mengonfirmasi bahwa *Perceived Threats* (PT) berpengaruh negatif dan signifikan terhadap *Generative AI Usage Intention*, dengan nilai koefisien jalur terstandarisasi sebesar -0,277 dan *Critical Ratio* sebesar -3,882 ($p < 0,001$). Temuan ini mengindikasikan bahwa kekhawatiran pengguna terhadap risiko keamanan data, privasi, dan akurasi *output generative AI* berkontribusi dalam menurunkan niat mereka untuk menggunakan teknologi tersebut, sejalan dengan proposisi *Technology Threat Avoidance Theory*. Dengan demikian, hipotesis kelima diterima.

Hasil pengujian hipotesis keenam menunjukkan bahwa *Perceived Avoidability* (PA) berpengaruh positif dan signifikan terhadap *Generative AI Usage Intention*, dengan nilai koefisien jalur terstandarisasi sebesar 0,287 dan *Critical Ratio* sebesar 4,082 ($p < 0,001$). Temuan ini mengindikasikan bahwa keyakinan pengguna akan kemampuannya mengelola atau menghindari risiko *generative AI*, misalnya melalui verifikasi *output* dan pembatasan data sensitif, berkontribusi dalam memperkuat niat penggunaan sekaligus meredam dampak negatif dari persepsi ancaman. Dengan demikian, hipotesis keenam diterima.

Secara keseluruhan, temuan penelitian ini memberikan bukti empiris yang kuat bahwa seluruh enam hipotesis diterima dan kerangka *IS Dual Factor* yang mengintegrasikan UTAUT sebagai jalur *enablers* dan TTAT sebagai jalur *inhibitors* terbukti relevan dalam menjelaskan niat penggunaan *generative AI* di populasi urban Indonesia. *Facilitating Conditions* dan *Performance Expectancy* secara bersama-sama merupakan dua pendorong terkuat, sementara *Perceived Threats* berperan sebagai penghambat yang signifikan dan *Perceived Avoidability* berfungsi sebagai penyeimbang yang meningkatkan niat penggunaan meskipun ancaman dirasakan.

DAFTAR PUSTAKA

Alelyani, M., Alamri, S., Alqahtani, M. S., Musa, A., Almater, H., Alqahtani, N., Alshahrani, F., &

- Alelyani, S. (2021). Radiology community attitude in Saudi Arabia about the applications of artificial intelligence in radiology. *Journal of Radiation Research and Applied Sciences*, 14(1), 1-10.
- Alsharida, R. A., & Al-sharafi, M. A. (2025). Predicting cybersecurity behaviors in the metaverse through the lenses of TTAT and TPB: A hybrid SEM-ANN approach. *Online Information Review*, 49(4), 729-750. <https://doi.org/10.1108/OIR-08-2023-0425>
- Arslankara, V. B. (2024). Generative artificial intelligence as a lifelong learning self efficacy: Usage and competence scale. *Journal of Educational Technology*, 6(2), 288-302.
- Baumgartner, H., & Homburg, C. (1996). Applications of structural equation modeling in marketing and consumer research: A review. *International Journal of Research in Marketing*, 13(2), 139-161.
- Beglar, D., & Nemoto, T. (2014). Developing Likert-scale questionnaires. *JALT2013 Conference Proceedings*, 1-8.
- Bick, A., Blandin, A., & Deming, D. J. (2025). *The rapid adoption of generative AI*. National Bureau of Economic Research.
- Breitwieser, M., Zirknitzer, S., Poslusny, K., Freude, T., Scholsching, J., Bodenschatz, K., Wagner, A., Hergan, K., Schaffert, M., Metzger, R., & Marko, P. (2025). AI in fracture detection: A cross-disciplinary analysis of physician acceptance using the UTAUT model. *Journal of Medical Systems*, 49(1), 1-16.
- Brown, R., Sillence, E., & Branley-bell, D. (2025). AcademAI: Investigating AI usage, attitudes, and literacy in higher education and research. *Journal of Educational Computing Research*, 63(2), 1-28. <https://doi.org/10.1177/00472395251347304>
- Browne, M. W., & Cudeck, R. (1992). Alternative ways of assessing model fit. *Sociological Methods & Research*, 21(2), 230-258. <https://doi.org/10.1177/0049124192021002005>
- Brynjolfsson, E., Li, D., & Raymond, L. R. (2023). Generative AI at work. *Quarterly Journal of Economics*, 140(1), 1-45.
- Cabero-Almenara, J., Palacios-Rodríguez, A., Loaiza-Aguirre, M. I., & Andrade-Abarca, P. S. (2024). The impact of pedagogical beliefs on the adoption of generative AI in higher education: Predictive model from UTAUT2. *Frontiers in Artificial Intelligence*, 7, 1-11. <https://doi.org/10.3389/frai.2024.1497705>
- Carolan, S., Martin, A. W., Gong, C. C., & Borja, S. (2025). 2025: *The state of consumer AI*. AI's Consumer Tipping Point Has Arrived. 1-44.
- Cenfetelli, R. T. (2004). Inhibitors and enablers as dual factor concepts in technology usage. *Journal of the Association for Information Systems*, 5(11), 16. <https://doi.org/10.17705/1jais.00059>
- Cerar, J., Nell, P. C., & Reiche, B. S. (2021). The declining share of primary data and the neglect of the individual level in international business research. *Journal of International Business Studies*, 52(7), 1365-1374. <https://doi.org/10.1057/s41267-021-00451-0>
- Chen, H., & Li, W. (2025). Mobile device users' privacy security assurance behavior: A technology threat avoidance perspective. *Information & Computer Security*, 33(3), 330-344. <https://doi.org/10.1108/ICS-04-2016-0027>
- Choi, H., & Park, S. (2024). Enhancing participatory security culture in public institutions: An analysis of organizational employees' security threat recognition processes. *IEEE Access*, 12, 47543-47558. <https://doi.org/10.1109/ACCESS.2024.3383311>
- Collins, C., Dennehy, D., Conboy, K., & Mikalef, P. (2021). Artificial intelligence in information systems research: A systematic literature review and research agenda. *International Journal of Information*

- Management, 60, 102383. <https://doi.org/10.1016/j.ijinfomgt.2021.102383>
- Cranney, K., Delecourt, S., & Koning, R. (2026). *Global evidence on gender gaps and generative AI over time*. Harvard Business School Working Paper.
- Deloitte. (2025). Women and generative AI: The adoption gap is closing fast, but a trust gap persists. *Deloitte Insights*, 1-10.
- Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., Baabdullah, A. M., Koohang, A., Raghavan, V., Ahuja, M., Albanna, H., Albashrawi, M. A., Al-Busaidi, A. S., Balakrishnan, J., Barlette, Y., Basu, S., Bose, I., Brooks, L., Buhalis, D., ... Wright, R. (2023). Opinion paper: "So what if ChatGPT wrote it?" Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- Fornell, C., & Larcker, D. F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18(1), 39-50.
- Gupta, V. (2024). *An empirical evaluation of a generative artificial intelligence technology adoption model from entrepreneurs' perspectives*. Doctoral dissertation.
- Hair, J. F., Black, W. C., Babin, B. J., & Anderson, R. E. (2010). *Multivariate data analysis* (7th ed.). Pearson Prentice Hall.
- Joshi, A., Kale, S., Chandel, S., & Pal, D. K. (2015). Likert scale: Explored and explained. *British Journal of Applied Science & Technology*, 7(4), 396-403. <https://doi.org/10.9734/BJAST/2015/14975>
- Kalmus, K., & Nikiforova, A. (2024). Resistance to generative AI in higher education. *arXiv Preprint*, 1-15.
- Keerativutisest, V., Chaiyasonthorn, W., Khalid, B., Ślusarczyk, B., & Chaveesuk, S. (2024). Determinants of AI-based applications adoption in the agricultural sector: Multi-group analysis. *Journal of Agricultural Science*, 72(1), 750-764.
- Kim, M. J., Lee, C.-K., Petrick, J. F., & Kim, Y. S. (2020). The influence of perceived risk and intervention on international tourists' behavior during the Hong Kong protest: Application of an extended model of goal-directed behavior. *Journal of Hospitality and Tourism Management*, 45, 622-632. <https://doi.org/10.1016/j.jhtm.2020.11.003>
- Kim, Y., & Blazquez, V. (2024). Determinants of generative AI system adoption and usage behavior in Korean companies: Applying the UTAUT model. *Journal of Business Research*, 180, 1-15.
- Kline, R. B. (2016). *Principles and practice of structural equation modeling* (4th ed.). Guilford Press.
- Lee, J., Park, J., & Mostafa, N. A. (2025). Artificial intelligence dimensions and design features for the knowledge-based work of the future: A socio-technical perspective. *Data Science and Management*, 8(2), 112-128.
- Liang, H., & Xue, Y. (2009). Avoidance of information technology threats: A theoretical perspective. *MIS Quarterly*, 33(1), 71-90. <https://doi.org/10.2307/20650279>
- Liu, C., Yang, L., Dong, X., & Li, X. (2025). Factors influencing generative AI usage intention in China: Extending the acceptance-avoidance framework with perceived AI literacy. *Computers in Human Behavior*, 158, 1-27.
- Mia, M., Majri, Y., Ibrahim, P., & Abdul, K. (2019). Covariance based-structural equation modeling (CB-SEM) using AMOS in management research. *Journal of Business and Management*, 21(1), 56-61. <https://doi.org/10.9790/487X-2101025661>

- Mishra, S., Mishra, A., Dubey, A., & Dwivedi, Y. K. (2023). Virtual reality in retailing: A meta-analysis to determine the purchase and non-purchase behavioural intention of consumers. *Journal of Retailing and Consumer Services*, 72, 1-15.
- Noy, S., & Zhang, W. (2023). Experimental evidence on the productivity effects of generative artificial intelligence. *Science*, 381(6654), 187-192.
- Puisa, R., Montewka, J., & Krata, P. (2023). A framework estimating the minimum sample size and margin of error for maritime quantitative risk analysis. *Reliability Engineering and System Safety*, 235, 109221. <https://doi.org/10.1016/j.ress.2023.109221>
- Qiao, J., Li, K., Han, W., & Qi, J. (2025). Investigating clinical teachers' adoption intention and use behavior of AI-assisted teaching: An extended UTAUT2 model approach. *Medical Education Online*, 30(1), 1-18.
- Rachman, A. (2024). *Metode penelitian kuantitatif, kualitatif dan R&D*. RajaGrafindo Persada.
- Razaque, A., Ajlan, A. Al, Melaoune, N., Alotaibi, M., Alotaibi, B., Dias, I., Oad, A., Hariri, S., & Zhao, C. (2021). Avoidance of cybersecurity threats with the deployment of a web-based blockchain-enabled cybersecurity awareness system. *Applied Sciences*, 11(1), 1-21.
- Resourcera. (2025). Global AI usage statistics 2025. *Resourcera Insights*.
- Shi, J., Lee, M., Girish, V. G., Xiao, G., & Lee, C.-K. (2025). Embracing the ChatGPT revolution: Unlocking new horizons for tourism. *Journal of Hospitality and Tourism Technology*, 15(3), 433-448. <https://doi.org/10.1108/JHTT-07-2023-0203>
- Tehreem, S., Ahmad, M., & Khan, S. (2025). Determinants of ChatGPT adoption intention among university students: An extension of UTAUT. *Asian Journal of University Education (AJUE)*, 21(1), 45-62.
- Upadhyay, N. (2022). Theorizing artificial intelligence acceptance and digital entrepreneurship model. *International Journal of Entrepreneurial Behavior & Research*, 28(5), 1138-1166. <https://doi.org/10.1108/IJEER-01-2021-0052>
- Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS Quarterly*, 27(3), 425-478. <https://doi.org/10.2307/30036540>
- Venkatesh, V., Thong, J. Y. L., & Xu, X. (2012). Consumer acceptance and use of information technology: Extending the unified theory of acceptance and use of technology. *MIS Quarterly*, 36(1), 157-178. <https://doi.org/10.2307/41410412>
- Wah, L. L., & Hashim, H. (2021). Determining pre-service teachers' intention of using technology for teaching English as a second language (ESL). *Journal of Technology and Education*, 15(2), 78-95.
- Wang, K., Ruan, Q., Zhang, X., & Fu, C. (2024). Pre-service teachers' GenAI anxiety, technology self-efficacy, and TPACK: Their structural relations with behavioral intention to design GenAI-assisted teaching. *Computers & Education*, 215, 1-18.
- Westland, J. C. (2010). Lower bounds on sample size in structural equation modeling. *Electronic Commerce Research and Applications*, 9(6), 476-487. <https://doi.org/10.1016/j.elerap.2010.07.003>
- Wichmann, J., Gesk, T. S., & Leyer, M. (2024). Acceptance of AI in health care for short- and long-term treatments: Pilot development study of an integrated theoretical model. *JMIR Human Factors*, 11(1), e48600. <https://doi.org/10.2196/48600>
- Wu, C., Gong, X., Luo, L., Zhao, Q., Hu, S., & Mou, Y. (2021). Applying control-value theory and unified theory of acceptance and use of technology to explore pre-service teachers' academic emotions and

- learning satisfaction. *Frontiers in Psychology*, 12, 738959. <https://doi.org/10.3389/fpsyg.2021.738959>
- Wu, D., Tian, L., & Liu, Q. (2025). Understanding university students' adoption of ChatGPT: A UTAUT3 perspective. *Current Psychology*, 44(3), 1567-1582.
- Yasmin, N., Zakaria, K., & Hashim, H. (2024). Shaping the future of education: Conceptualising pre-service teachers' perspectives on artificial intelligence (AI) integration. *International Journal of Academic Research in Business and Social Sciences*, 14(5), 643-653. <https://doi.org/10.6007/IJARBSS/v14-i5/21584>
- Young, D., Carpenter, D., & Mcleod, A. (2016). Malware avoidance motivations and behaviors: A technology threat avoidance replication. *Journal of Information Systems*, 30(2), 1-17.
- Yuniawan, A., Hersugondo, H., Mas, F., & Ali, M. (2025). Beyond the hype: Understanding barriers to AI adoption through lens. *Asian Journal of Information Management*, 19(3), 45-62. <https://doi.org/10.1108/AJIM-08-2024-0619>
- Zhai, H., Yang, X., Xue, J., Lavender, C., & Ye, T. (2021). Radiation oncologists' perceptions of adopting an artificial intelligence-assisted contouring technology: Model development and questionnaire study. *JMIR Medical Informatics*, 9(5), e27122. <https://doi.org/10.2196/27122>