



PENGARUH ATTITUDE TOWARD AI, PERCEIVED AI ACCURACY, AI ANXIETY DAN TRUST IN AI TERHADAP WILLINGNESS TO ACCEPT AI UNTUK Pengerjaan TUGAS KULIAH MAHASISWA SURABAYA

Sandrak Nababan¹⁾, Anggraeni Widya Purwita²⁾

¹⁾Sistem Informasi, Fakultas Teknik, Universitas Negeri Surabaya, Surabaya, Indonesia
Email: sandrak.22141@mhs.unesa.ac.id

²⁾Sistem Informasi, Fakultas Teknik, Universitas Negeri Surabaya, Surabaya, Indonesia
Email: anggraenipurwita@unesa.ac.id

Abstract

Generative artificial intelligence (AI) is now widely used in university learning. However, frequent use does not always mean that students understand the technology or trust it appropriately. This study examines the effects of Attitude Toward AI, Perceived AI Accuracy, and AI Anxiety on Trust in AI. It also examines the effect of Trust in AI on students' Willingness to Accept AI for academic assignments. This quantitative cross-sectional study involved 204 active students from public and private universities in Surabaya. All respondents had used generative AI at least three times in the previous month. Data were collected through a 20-item questionnaire using a seven-point Likert scale and analyzed using PLS-SEM in SmartPLS 4. The measurement model met the criteria for convergent validity, discriminant validity, and reliability. Attitude Toward AI ($\beta = 0.438$), Perceived AI Accuracy ($\beta = 0.394$), and AI Anxiety ($\beta = -0.474$) significantly explained Trust in AI ($R^2 = 0.584$). Trust in AI had a strong positive effect on Willingness to Accept AI ($\beta = 0.749$; $R^2 = 0.561$). All four paths were significant at $p = 0.001$. These results show that trust has an important role in students' acceptance of generative AI. Positive attitudes and perceived accuracy increase trust, while AI Anxiety reduces it.

Keywords: Generative AI; Trust in AI; AI Anxiety; Perceived AI Accuracy; Willingness to Accept AI.

Abstrak

Kecerdasan buatan generatif kini digunakan secara luas dalam pembelajaran di perguruan tinggi. Namun, penggunaan yang sering tidak selalu berarti bahwa mahasiswa memahami teknologi tersebut atau memercayainya secara tepat. Penelitian ini menguji pengaruh Attitude Toward AI, Perceived AI Accuracy, dan AI Anxiety terhadap Trust in AI, serta pengaruh Trust in AI terhadap Willingness to Accept AI untuk menyelesaikan tugas akademik. Penelitian kuantitatif dengan desain potong lintang ini melibatkan 204 mahasiswa aktif dari perguruan tinggi negeri dan swasta di Surabaya. Seluruh responden telah menggunakan AI generatif setidaknya tiga kali dalam satu bulan terakhir. Data dikumpulkan melalui kuesioner 20 butir dengan skala Likert tujuh poin dan dianalisis menggunakan PLS-SEM pada SmartPLS 4. Model pengukuran memenuhi kriteria validitas konvergen, validitas diskriminan, dan reliabilitas. Attitude Toward AI ($\beta = 0,438$), Perceived AI Accuracy ($\beta = 0,394$), dan AI Anxiety ($\beta = -0,474$) secara signifikan menjelaskan Trust in AI ($R^2 = 0,584$). Trust in AI berpengaruh positif kuat terhadap Willingness to Accept AI ($\beta = 0,749$; $R^2 = 0,561$). Seluruh jalur signifikan pada $p = 0,001$. Temuan ini menunjukkan bahwa kepercayaan berperan penting dalam penerimaan AI generatif oleh mahasiswa. Sikap positif dan persepsi atas akurasi meningkatkan kepercayaan, sedangkan kecemasan terhadap AI menurunkannya.

Kata Kunci: AI Generatif; Kepercayaan terhadap AI; Kecemasan terhadap AI; Persepsi Akurasi AI; Kesiapan Menerima AI.

PENDAHULUAN

Generative AI seperti ChatGPT, Gemini, dan Microsoft Copilot telah mengubah secara cepat cara mahasiswa mencari informasi, mengembangkan gagasan, dan menyelesaikan tugas akademik. Digital Education Council (2024) melaporkan bahwa 86% mahasiswa telah menggunakan AI untuk pembelajaran, lebih dari 54% menggunakannya setiap minggu, dan rata-rata mahasiswa menggunakan lebih dari dua perangkat AI. Survei HEPI-Kortext (2025) selanjutnya menunjukkan bahwa penggunaan AI oleh mahasiswa meningkat dari 66% pada 2024 menjadi sekitar 92% pada 2025. Angka tersebut menunjukkan bahwa AI generatif bukan lagi sekadar alat bantu sesekali, melainkan telah menjadi bagian dari praktik sehari-hari di perguruan tinggi.

Namun, tingginya frekuensi penggunaan tidak serta-merta menunjukkan bahwa mahasiswa memahami cara menggunakan AI secara efektif. Zhao et al. (2025) menemukan bahwa 84,46% mahasiswa menggunakan AI generatif untuk kegiatan akademik, tetapi hanya sekitar 24% yang merasa mampu menjelaskan pertanyaan secara efektif dan lebih dari 83% mengalami kesulitan menyusun prompt yang tepat. Kesenjangan ini penting karena AI generatif dapat menghasilkan informasi yang tampak meyakinkan, tetapi keliru, tidak logis, atau menyesatkan (Kim et al., 2025). Mahasiswa dapat menghargai kemudahan yang ditawarkan AI, tetapi tetap meragukan akurasi, keandalan, privasi, integritas akademik, serta dampaknya terhadap kemampuan berpikir mandiri.

Persepsi mahasiswa membantu menjelaskan mengapa intensitas penggunaan yang tinggi dapat berjalan beriringan dengan penerimaan yang tetap berhati-hati. Sikap positif terhadap AI (Attitude Toward AI) terbentuk ketika teknologi membuat pembelajaran lebih efisien, mendukung pencarian informasi, dan membantu penyelesaian tugas (Chan & Hu, 2023). Sikap tersebut dapat memperkuat niat adopsi (Belanche et al., 2019). Persepsi bahwa keluaran AI akurat dan relevan (Perceived AI Accuracy) juga dapat mendorong penggunaan berkelanjutan (Elena et al., 2025). Sebaliknya, kecemasan terhadap AI (AI Anxiety) dapat muncul akibat keluaran yang tidak akurat, ketergantungan, masalah privasi, dan risiko integritas akademik, serta dapat memengaruhi motivasi dan perilaku belajar (Wang et al., 2022).

Kepercayaan terhadap AI (Trust in AI) sangat penting dalam konteks akademik karena mahasiswa menggunakan materi yang dihasilkan AI untuk tugas yang menuntut bukti kredibel dan penalaran yang tepat. Sikap positif dapat meningkatkan keyakinan terhadap teknologi (Daly et al., 2025), sedangkan persepsi akurasi menjadi dasar praktis untuk menilai sistem yang proses internalnya tidak transparan bagi pengguna (Kleine et al., 2026). Kekhawatiran terhadap jawaban yang keliru,

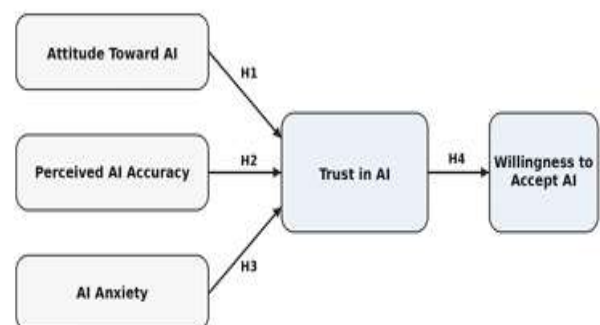
ketergantungan, dan persoalan etika dapat menurunkan kepercayaan (Zhou et al., 2025). Ketika kepercayaan telah terbentuk, pengguna cenderung lebih bersedia menerima AI untuk aktivitas yang memiliki konsekuensi penting (Svestkova et al., 2025).

Model penelitian ini mengadaptasi kerangka berbasis kepercayaan dari Pal et al. (2023) yang dikembangkan untuk sistem rekomendasi AI dalam investasi saham, serta Ananth et al. (2025) yang menguji konten hasil AI yang sensitif terhadap budaya. Bukti empiris dalam konteks pendidikan tinggi masih terbatas, khususnya pada mahasiswa di Surabaya yang menggunakan AI generatif secara berulang untuk kegiatan perkuliahan. Oleh karena itu, penelitian ini menguji pengaruh sikap terhadap AI (AT), persepsi akurasi AI (PA), dan kecemasan terhadap AI (AA) terhadap kepercayaan terhadap AI (TI), kemudian menguji pengaruh TI terhadap kesediaan menerima AI (WA). Penelitian ini memberikan bukti kontekstual mengenai proses terbentuknya kepercayaan dan perubahannya menjadi penerimaan dalam penggunaan akademik.

METODE PELAKSANAAN

3.1 Model Penelitian dan Hipotesis

Penelitian ini menggunakan desain kuantitatif potong lintang untuk menguji lima konstruk laten. Sikap terhadap AI (Attitude Toward AI/AT) mencerminkan penilaian positif atau negatif mahasiswa terhadap penggunaan AI. Persepsi akurasi AI (Perceived AI Accuracy/PA) mencerminkan penilaian mahasiswa mengenai ketepatan dan konsistensi keluaran AI. Kecemasan terhadap AI (AI Anxiety/AA) mencerminkan kekhawatiran ketika mengandalkan AI, termasuk informasi yang tidak akurat, ketergantungan, dan persoalan integritas akademik. Kepercayaan terhadap AI (Trust in AI/TI) merupakan keyakinan bahwa AI dapat diandalkan, aman, dan tepercaya. Kesediaan menerima AI (Willingness to Accept AI/WA) merupakan kesiapan untuk menggunakan serta menerima informasi atau rekomendasi dari AI generatif. Hubungan yang diajukan disajikan pada Gambar 1.



Gambar 1. Model penelitian dan hubungan hipotesis

Sumber: Pal et al. (2023)



H1:

H₀₁: Attitude Toward AI tidak berpengaruh terhadap Trust in AI.

H_{a1}: Attitude Toward AI berpengaruh terhadap Trust in AI.

H2:

H₀₂: Perceived AI Accuracy tidak berpengaruh terhadap Trust in AI.

H_{a2}: Perceived AI Accuracy berpengaruh terhadap Trust in AI.

H3:

H₀₃: AI Anxiety tidak berpengaruh terhadap Trust in AI.

H_{a3}: AI Anxiety berpengaruh terhadap Trust in AI.

H4:

H₀₄: Trust in AI tidak berpengaruh terhadap Willingness to Accept Generative AI Tools.

H_{a4}: Trust in AI berpengaruh terhadap Willingness to Accept Generative AI Tools

3.2 Populasi, Sampel, dan Pengumpulan Data

Populasi sasaran penelitian adalah mahasiswa aktif pada perguruan tinggi negeri atau swasta di Surabaya yang menggunakan AI generatif untuk mendukung penyelesaian tugas akademik. Teknik purposive non-probability sampling digunakan karena responden harus memenuhi kriteria pengalaman tertentu (Sugiyono, 2013). Responden yang memenuhi syarat adalah mahasiswa yang telah menggunakan AI generatif sekurang-kurangnya tiga kali selama satu bulan sebelum mengisi kuesioner. Kriteria tersebut memastikan bahwa responden memiliki pengalaman penggunaan berulang yang memadai untuk menilai teknologi.

Pedoman lima hingga sepuluh kali jumlah indikator (Hair et al., 2014) digunakan sebagai acuan perencanaan, sehingga model dengan 20 indikator membutuhkan sekitar 100-200 respons. Data dikumpulkan pada Maret-Juli 2026 melalui kuesioner Google Forms yang disebarakan melalui grup mahasiswa, surel, dan media sosial. Partisipasi dilakukan secara sukarela. Jawaban responden diukur menggunakan skala Likert tujuh poin, mulai dari 1 (sangat tidak setuju) hingga 7 (sangat setuju).

3.3 Instrumen Pengukuran dan Uji Coba

Setiap konstruk diukur menggunakan empat butir. Sumber butir dan makna operasional konstruk dirangkum pada Tabel 1. Redaksi butir diterjemahkan dan disesuaikan dengan konteks penggunaan AI generatif oleh mahasiswa dalam tugas akademik tanpa mengubah makna konseptual ukuran asli. Uji coba terhadap 40 responden yang memiliki karakteristik serupa dengan populasi sasaran dianalisis menggunakan SPSS 26. Seluruh 20 butir memiliki nilai di

atas *r* tabel Pearson sebesar 0,312, dengan korelasi butir berkisar antara 0,745-0,966. Nilai Cronbach's alpha pada tingkat konstruk berkisar antara 0,778 0,882, sehingga instrumen dinilai layak digunakan dalam survei utama.

Tabel 1. Operasionalisasi konstruk dan sumber

Konstruk	Kode	Makna operasional	Sumber
Attitude Toward AI	AT	Penilaian positif atau negatif terhadap penggunaan AI untuk tugas.	Belanche et al. (2019); Bhattacharjee (2000)
Perceived AI Accuracy	PA	Penilaian bahwa keluaran AI akurat, konsisten, dan dapat diterapkan.	Gursoy et al. (2019)
AI Anxiety	AA	Kekhawatiran terhadap kesalahan, ketergantungan, dan risiko integritas akademik.	Cobelli et al. (2021)
Trust in AI	TI	Keyakinan bahwa AI dapat diandalkan, aman, dan tepercaya.	Jamaludin & Ahmad (2013)
Willingness to Accept AI	WA	Kesediaan menggunakan serta menerima informasi atau rekomendasi dari AI.	Gursoy et al. (2019); Zhang et al. (2020)

Sumber: Instrumen diadaptasi oleh penulis.

3.4 Prapemrosesan data

Sebanyak 216 respons kuesioner berhasil dihimpun. Empat respons dengan pola jawaban lurus (straight-lining) dikeluarkan. Delapan respons lainnya tidak memenuhi kriteria kelayakan; dua dari delapan respons tersebut juga mengandung data hilang (missing data). Karena kedua kasus data hilang telah termasuk dalam kelompok yang tidak memenuhi kriteria, keduanya tidak dikurangkan kembali. Dengan demikian, sampel akhir untuk analisis adalah $216 - 4 - 8 = 204$ responden valid. Logika penyaringan ini mempertahankan penyebut yang tepat dan mencegah penghitungan ganda pada data yang dikeluarkan.

Arah seluruh butir AA diperiksa sebelum analisis. Pengodean balik (reverse coding) hanya diperlukan apabila persetujuan terhadap suatu butir menunjukkan tingkat konstruk yang berlawanan, bukan semata-mata karena butir tersebut menggunakan redaksi negatif (Dueber et al., 2021; Weijters et al., 2013). Persetujuan terhadap setiap butir AA secara langsung menunjukkan kecemasan yang lebih tinggi terhadap keluaran keliru, ketergantungan, atau persoalan etika akademik. Oleh sebab itu, indikator AA tidak melalui pengodean balik; koefisien negatif pada jalur AA terhadap TI mencerminkan hubungan substantif, bukan artefak penskoran.

3.5 Analisis PLS-SEM

Analisis PLS-SEM dilakukan menggunakan SmartPLS 4 karena metode ini mendukung analisis model variabel laten yang berorientasi pada prediksi tanpa mensyaratkan normalitas multivariat (Rahadi, 2021). Model pengukuran dievaluasi berdasarkan outer loading ($\geq 0,70$), average variance extracted (AVE $\geq 0,50$), Cronbach's



alpha dan composite reliability ($\geq 0,70$), HTMT ($< 0,90$), kriteria Fornell-Larcker, serta cross-loading (Kusmantini et al., 2024; Hair et al., 2022). Model struktural dievaluasi menggunakan R^2 dan f^2 . Signifikansi jalur diestimasi melalui 5.000 bootstrap resampling; hubungan dinyatakan signifikan apabila $t > 1,96$ dan $p < 0,05$.

HASIL DAN PEMBAHASAN

4.1 Profil responden dan penggunaan AI generatif

Sampel akhir terdiri atas 204 mahasiswa. Responden laki laki sedikit lebih banyak daripada responden perempuan, dan sebagian besar peserta berusia 21-23 tahun. Intensitas penggunaan AI generatif tergolong tinggi: 94,1% responden menggunakannya sedikitnya enam kali per bulan. Hampir seluruh responden menggunakan ChatGPT, sedangkan banyak responden juga menggunakan Gemini dan Claude untuk membandingkan keluaran serta memilih informasi yang sesuai dengan kebutuhan tugas. Karakteristik responden secara rinci disajikan pada Tabel 2.

Tabel 2. Karakteristik responden

Karakteristik	n	%
Jenis kelamin - Laki-laki	115	56,4
Jenis kelamin - Perempuan	89	43,6
Usia - 20 tahun	12	5,9
Usia - 21 tahun	45	22,1
Usia - 22 tahun	78	38,2
Usia - 23 tahun	50	24,5
Usia - 24 tahun	17	8,3
Usia - 25 tahun	2	1,0
Penggunaan AI - 3-5 kali/bulan	12	5,9
Penggunaan AI - 6-10 kali/bulan	100	49,0
Penggunaan AI - >10 kali/bulan	92	45,1
Platform - ChatGPT	202	99,0
Platform - Gemini	149	73,0
Platform - Claude	106	52,0
Platform - DeepSeek	62	30,4
Platform - Microsoft Copilot	31	15,2
Platform - Perplexity AI	19	9,3
Platform - Monica	1	0,5

Sumber: Data survei primer, 2026. Persentase platform memungkinkan jawaban ganda.

Responden terutama menggunakan AI untuk riset serta merangkum materi perkuliahan atau artikel jurnal (94,0%), melakukan parafrase (91,2%), dan mencari artikel ilmiah (90,8%). Penggunaan lainnya meliputi brainstorming (58,5%), pembuatan gambar atau desain (47,0%), dan bantuan pemrograman (37,3%). Alasan utama penggunaan adalah mempercepat penyelesaian tugas (93,5%), diikuti oleh mengatasi kesulitan menulis (76,5%)

dan memahami materi perkuliahan yang sulit (76,0%). Tingkat penggunaan langsung tersebut memberikan dasar pengalaman yang memadai untuk menilai akurasi, kecemasan, kepercayaan, dan penerimaan.

4.2 Model pengukuran

Outer loading berada pada rentang 0,765-0,856 ($>0,70$), AVE 0,633-0,695, Cronbach's alpha 0,807-0,853, dan composite reliability 0,873-0,901. Seluruh nilai tersebut memenuhi kriteria reliabilitas indikator, validitas konvergen, dan konsistensi internal (Tabel 3).

Tabel 3. Evaluasi model pengukuran

Konstruk	Rentang loading	AVE	α	rho_A	rho_C
AA	0.783-0.849	0.659	0.828	0.830	0.886
AT	0.765-0.856	0.656	0.825	0.828	0.884
PA	0.768-0.824	0.633	0.807	0.807	0.873
TI	0.807-0.852	0.695	0.853	0.855	0.901
WA	0.784-0.841	0.665	0.832	0.833	0.888

Validitas diskriminan dikonfirmasi melalui tiga pemeriksaan. Seluruh nilai HTMT berada di bawah 0,90; nilai 4 tertinggi adalah 0,888 pada hubungan TI dan WA (Tabel 4). Berdasarkan kriteria Fornell-Larcker, akar kuadrat AVE setiap konstruk lebih besar daripada korelasinya dengan konstruk lain (Tabel 5). Selain itu, seluruh indikator memiliki loading tertinggi pada konstruk yang diukur. Secara keseluruhan, hasil tersebut menunjukkan bahwa kelima konstruk dapat dibedakan secara empiris.

Tabel 4. Rasio heterotrait-monotrait (HTMT)

Konstruk	AA	AT	PA	TI	WA
AA	-				
AT	0.125	-			
PA	0.056	0.197	-		
TI	0.510	0.537	0.558	-	
WA	0.328	0.422	0.517	0.888	-

Tabel 5. Kriteria Fornell-Larcker

Konstruk	AA	AT	PA	TI	WA
AA	0.812				
AT	0.106	0.810			
PA	0.000	0.163	0.796		
TI	-0.428	0.452	0.465	0.834	
WA	-0.270	0.351	0.424	0.749	0.815

Nilai diagonal merupakan akar kuadrat AVE.



4.3 Model struktural dan pengujian hipotesis

AT, PA, dan AA secara bersama-sama menjelaskan 58,4% varians TI (R^2 terkoreksi = 0,578), sedangkan TI menjelaskan 56,1% varians WA (R^2 terkoreksi = 0,559). Kedua nilai tersebut menunjukkan daya jelas yang moderat dalam konteks pendidikan (Tabel 6).

Tabel 6. Koefisien determinasi

Konstruk endogen	R^2	R^2 terkoreksi	Interpretasi
Kepercayaan terhadap AI (TI)	0.584	0.578	Moderat
Kesediaan Menerima AI (WA)	0.561	0.559	Moderat

Seluruh hubungan memiliki ukuran efek yang besar. Efek langsung terbesar terdapat pada jalur TI terhadap WA ($F^2 = 1,277$), sedangkan AA, AT, dan PA juga memberikan efek besar terhadap TI (Tabel 7).

Tabel 7. Ukuran efek

Hubungan	F^2	Interpretasi
AA $\rightarrow$ TI	0.535	Besar
AT $\rightarrow$ TI	0.444	Besar
PA $\rightarrow$ TI	0.363	Besar
TI $\rightarrow$ WA	1.277	Besar

Hasil bootstrapping mendukung seluruh hipotesis (Tabel 8). AT dan PA berpengaruh positif dan signifikan terhadap TI; AA berpengaruh negatif dan signifikan terhadap TI; sedangkan TI berpengaruh positif dan signifikan terhadap WA. Nilai koefisien, galat baku, t, dan p dilaporkan tanpa perubahan dari kumpulan data jurnal yang dianalisis.

Tabel 8. Pengujian hipotesis

Hip.	Jalur	β	SE	t	p	Keputusan
H2	PA $\rightarrow$ TI	0.394	0.049	8.073	0.001	Didukung
H3	AA $\rightarrow$ TI	-0.474	0.045	10.485	0.001	Didukung
H4	TI $\rightarrow$ WA	0.749	0.035	21.697	0.001	Didukung

4.4 Sikap terhadap AI dan kepercayaan terhadap AI

H1 didukung ($\beta = 0,438$; $t = 8,828$; $p = 0,001$). Mahasiswa dengan sikap yang lebih positif terhadap AI cenderung memiliki kepercayaan yang lebih tinggi terhadap AI. Mahasiswa yang menilai AI generatif bermanfaat, membantu, dan sesuai untuk kegiatan akademik lebih mungkin memercayai teknologi tersebut. Interpretasi ini selaras dengan pola penggunaan responden: banyak mahasiswa mengandalkan AI untuk mempercepat penyelesaian tugas, menghasilkan gagasan, dan memahami materi yang sulit. Temuan ini konsisten dengan Pal et al. (2023) dan Ananth et al. (2025) yang mengaitkan penilaian

positif terhadap AI dengan keyakinan yang lebih tinggi pada sistem berbasis AI.

4.5 Persepsi akurasi AI dan kepercayaan terhadap AI

H2 didukung ($\beta = 0,394$; $t = 8,073$; $p = 0,001$). Mahasiswa yang menilai keluaran AI akurat, relevan, dan sesuai dengan kebutuhan perkuliahan cenderung memandang teknologi tersebut dapat diandalkan. Akurasi menjadi sangat penting karena riset, peringkasan, pencarian artikel ilmiah, dan parafrase merupakan beberapa penggunaan yang paling umum. Ketika proses internal sistem tidak transparan, kualitas keluaran menjadi sinyal praktis dalam membentuk kepercayaan (Kleine et al., 2026). Hasil ini juga selaras dengan bukti bahwa keluaran yang relevan dan akurat memperkuat kepercayaan pengguna di perguruan tinggi (Elena et al., 2025).

4.6 Kecemasan terhadap AI dan kepercayaan terhadap AI

H3 didukung ($\beta = -0,474$; $t = 10,485$; $p = 0,001$). Kecemasan terhadap AI yang lebih tinggi berkaitan dengan kepercayaan terhadap AI yang lebih rendah. Mahasiswa mengkhawatirkan informasi yang keliru atau salah ditafsirkan, ketergantungan terhadap AI, serta persoalan etika akademik seperti plagiarisme atau ketidakpatuhan terhadap instruksi dosen. Jalur negatif tersebut bukan disebabkan oleh reverse coding karena seluruh indikator AA secara langsung mengukur kecemasan. Hasil ini mendukung Cobelli et al. (2021) yang memandang kekhawatiran terhadap teknologi sebagai biaya psikologis, serta Zhou et al. (2025) yang mengaitkan kecemasan terhadap AI dengan penurunan kepercayaan.

4.7 Kepercayaan terhadap AI dan kesediaan menerima AI

H4 didukung ($\beta = 0,749$; $t = 21,697$; $p = 0,001$) dan merupakan hubungan terkuat dalam model ($F^2 = 1,277$). Mahasiswa yang meyakini bahwa AI generatif aman, dapat diandalkan, dan mampu memberikan informasi yang bermanfaat menunjukkan kesediaan yang lebih tinggi untuk menggunakannya dalam kegiatan akademik. Dengan demikian, kepercayaan menjadi mekanisme utama yang menerjemahkan persepsi terhadap AI menjadi penerimaan. Hasil ini konsisten dengan bukti terdahulu bahwa kesediaan menerima teknologi meningkat ketika pengguna memercayai kemampuannya (Jamaludin & Ahmad, 2013; Svestkova et al., 2025).

KESIMPULAN

Penelitian ini menunjukkan bahwa sikap terhadap AI dan persepsi akurasi AI secara signifikan meningkatkan kepercayaan terhadap AI, sedangkan kecemasan terhadap AI secara signifikan menurunkannya. Kepercayaan



terhadap AI selanjutnya berpengaruh positif kuat terhadap kesediaan menerima AI dan menjelaskan 56,1% variansnya. Mahasiswa lebih cenderung memercayai dan menerima AI generatif ketika teknologi tersebut membantu menyelesaikan tugas, menyediakan informasi yang dinilai akurat dan relevan, serta dirasakan aman untuk digunakan. Sebaliknya, kepercayaan menurun ketika mahasiswa mengkhawatirkan keluaran yang keliru, ketergantungan, atau persoalan integritas akademik.

Temuan penelitian ini memperluas penerapan model penerimaan AI berbasis kepercayaan dari konteks rekomendasi investasi dan konten budaya hasil AI ke konteks pendidikan tinggi. Bagi perguruan tinggi, penyediaan akses dan petunjuk operasional saja belum memadai. Literasi AI perlu melatih mahasiswa untuk memverifikasi informasi yang dihasilkan AI dengan sumber kredibel, mengenali risiko privasi dan integritas akademik, serta menggunakan AI tanpa menggantikan penalaran mandiri. Praktik tersebut dapat membentuk kepercayaan yang terkalibrasi, bukan ketergantungan tanpa sikap kritis. Penelitian selanjutnya perlu menguji model ini pada wilayah, institusi, disiplin ilmu, dan desain pembelajaran yang berbeda untuk mengetahui apakah konteks mengubah hubungan yang ditemukan.

DAFTAR PUSTAKA

- Ananth, A. S., Hussain, M., Dabic, M., & Vrontis, D. (2025). Cultural sensitivity and AI adoption: Exploring the willingness to embrace AI generated content. *Journal of Enterprise Information Management*. <https://doi.org/10.1142/S0219649225501308>
- Belanche, D., Casaló, L. V., & Flavián, C. (2019). Artificial intelligence in FinTech: Understanding robo-advisors adoption among customers. *Industrial Management & Data Systems*, 119(7), 1411-1430. <https://doi.org/10.1108/IMDS-08-2018-0368>
- Bhattacharjee, A. (2000). Acceptance of e-commerce services: The case of electronic brokerages. *IEEE Transactions on Systems, Man, and Cybernetics - Part A*, 30(4), 411-420.
- Chan, C. K. Y., & Hu, W. (2023). Students' voices on generative AI: Perceptions, benefits, and challenges in higher education. *International Journal of Educational Technology in Higher Education*, 20(1). <https://doi.org/10.1186/s41239-023-00411-8>
- Cobelli, N., Cassia, F., & Burro, R. (2021). Factors affecting the choices of adoption/non-adoption of future technologies during coronavirus pandemic. *Technological Forecasting and Social Change*, 169, 120814. <https://doi.org/10.1016/j.techfore.2021.120814>
- Daly, S. J., Wiewiora, A., & Hearn, G. (2025). Shifting attitudes and trust in AI: Influences on organizational AI adoption. *Technological Forecasting and Social Change*, <https://doi.org/10.1016/j.techfore.2025.124108> 215, 124108
- Digital Education Council. (2024). AI or not AI: What students want. <https://s.id/Digital-Education-Council>
- Dueber, D. M., et al. (2021). To reverse item orientation or not to reverse item orientation, that is the question. <https://doi.org/10.1177/10731911211017635> Assessment.
- Elena, D., Domagoj, F., & Marin, M. (2025). Trust in generative AI tools: A comparative study of higher education students, teachers, and researchers. *Information*, <https://doi.org/10.3390/info16070622> 16(7), 622.
- Gursoy, D., Chi, O. H., Lu, L., & Nunkoo, R. (2019). Consumers acceptance of artificially intelligent (AI) device use in service delivery. *International Journal of Information Management*, 49, 157-169. <https://doi.org/10.1016/j.ijinfomgt.2019.03.008>
- Hair, J. F., Hult, G. T. M., Ringle, C. M., Sarstedt, M., Lux, G., & Troiville, J. (2022). An introduction to structural equation modeling.
- Hair, J. F., Sarstedt, M., Hopkins, L., & Kuppelwieser, V. G. (2014). Partial least squares structural equation modeling (PLS-SEM): An emerging tool in business research. *European Business Review*, 26(2), 106-121. <https://doi.org/10.1108/EBR-10-2013-0128>
- HEPI-Kortext. (2025). Student generative AI survey 2025. <https://www.hepi.ac.uk/wp-content/uploads/2025/02/HEPI-Kortext-Student-Generative-AI-Survey-2025.pdf>
- Jamaludin, A., & Ahmad, F. (2013). Investigating the relationship between trust and intention to purchase online. *Business and Management Horizons*, 1(1), 1. <https://doi.org/10.5296/bmh.v1i1.3253>
- Kim, J., Yu, S., Detrick, R., & Li, N. (2025). Exploring students' perspectives on generative AI-assisted academic writing. *Education and Information Technologies*, 30(1). <https://doi.org/10.1007/s10639-024-12878-7>
- Kleine, A., Schaffernak, I., & Lerner, E. (2026). When do researchers trust AI? Trust as a mediator and experience as a boundary condition in researchers' intentions to use generative AI. *Computers in Human Behavior: Artificial Humans*, <https://doi.org/10.1016/j.chbah.2026.100310> 8, 100310.
- Kusmantini, T., Mubarak, A. A., & Anggraini, S. P. (2024). Metode analisis SmartPLS.



- Pal, S., Chua, A. Y. K., & Banerjee, S. (2023). Examining trust and willingness to accept AI recommendation systems. *Journal of Information Science*.
- Rahadi, D. R. (2021). Partial least squares structural equation modeling. In *Handbook of Market Research* https://doi.org/10.1007/978-3-319-57413-4_15 (pp. 587-632).
- Sugiyono. (2013). *Metode penelitian kuantitatif, kualitatif dan R&D*. Alfabeta.
- Svestkova, A., Huang, Y., & Smahel, D. (2025). Factors that influence trust and willingness to use generative AI for health information: A cross sectional study. *Digital Health*.
<https://doi.org/10.1177/20552076251360973>
- Wang, Y., et al. (2022). What drives students' AI learning behavior: A perspective of AI anxiety. *Interactive Learning Environments*, 1-17.
<https://doi.org/10.1080/10494820.2022.2153147>
- Weijters, B., Baumgartner, H., & Schillewaert, N. (2013). Reversed item bias: An integrative model. *Psychological Methods*.
- Zhang, Z., Genc, Y., Xing, A., Wang, D., Fan, X., & Citardi, D. (2020). Lay individuals' perceptions of artificial intelligence (AI)-empowered healthcare systems. *Proceedings of the Association for Information Science and Technology*, 57(1), 1-9.
<https://doi.org/10.1002/pra2.326>
- Zhao, Y., et al. (2025). The current status, knowledge, attitudes, and challenges of generative artificial intelligence use among undergraduate nursing students: A single-center cross-sectional survey of western China. *Frontiers in Public Health*, 13.
<https://doi.org/10.3389/fpubh.2025.1648416>
- Zhou, Q., et al. (2025). The mediation of trust on artificial intelligence anxiety and continuous adoption of artificial intelligence technology among primary nurses: A cross-sectional study.