

Analisis Deskriptif Komparatif Pemanfaatan ChatGPT, Kualitas Pemahaman, dan Efisiensi Tugas Berdasarkan Status Kerja Mahasiswa

Jonathan Wijaya^{*}, R Widya Henisaputri

Fakultas Komputer, Program Studi Sistem Informasi, Universitas Universal, Batam

Jl. Pasir Putih, Sadai, Bengkong, Kota Batam, Kepulauan Riau, Indonesia

Email: ^{1,*}jonathanwjyskom@gmail.com, ²widya@uvers.ac.id

Email Penulis Korespondensi: jonathanwjyskom@gmail.com

Submitted: 10/06/2026; Accepted: 04/07/2026; Published: 05/07/2026

Abstrak—Pemanfaatan ChatGPT dalam kegiatan akademik dapat memberikan penilaian yang berbeda pada aspek pemahaman dan efisiensi tugas. Perbedaan tuntutan antara mahasiswa bekerja dan tidak bekerja juga memungkinkan munculnya pola penggunaan yang tidak sama. Penelitian ini bertujuan menggambarkan Pemanfaatan ChatGPT, Kualitas Pemahaman, dan Efisiensi Tugas, membandingkan ketiga konstruk berdasarkan status kerja, serta membandingkan Kualitas Pemahaman dan Efisiensi Tugas pada responden yang sama. Penelitian menggunakan pendekatan kuantitatif dengan desain deskriptif-komparatif. Data diperoleh dari 101 mahasiswa Universitas Universal yang dipilih melalui purposive sampling menggunakan kuesioner skala Likert empat poin. Evaluasi instrumen menghasilkan 26 item final yang terbagi ke dalam tiga konstruk dengan nilai Cronbach's Alpha sebesar 0,787–0,913. Hasil Mann–Whitney U menunjukkan tidak terdapat perbedaan signifikan antara mahasiswa bekerja dan tidak bekerja pada Pemanfaatan ChatGPT ($p=0,923$), Kualitas Pemahaman ($p=0,244$), dan Efisiensi Tugas ($p=0,079$). Hasil Wilcoxon Signed-Rank Test menunjukkan bahwa mean per item Efisiensi Tugas sebesar 3,260 lebih tinggi daripada Kualitas Pemahaman sebesar 3,000 ($Z=-5,779$; $p<0,001$; $r=0,593$). Berdasarkan persepsi responden, penggunaan ChatGPT lebih menonjol pada kepraktisan dan efisiensi pengerjaan tugas dibandingkan pada kualitas pemahaman, sedangkan status kerja tidak menunjukkan perbedaan yang berarti pada ketiga konstruk. Kontribusi penelitian ini adalah memberikan bukti empiris yang memisahkan aspek pemanfaatan, pemahaman, dan efisiensi serta menyediakan dasar bagi perguruan tinggi dalam menyusun penggunaan ChatGPT yang tetap menekankan verifikasi informasi dan proses pemahaman mahasiswa.

Kata Kunci: ChatGPT; Kecerdasan Buatan Generatif; Kualitas Pemahaman; Efisiensi Tugas; Status Kerja Mahasiswa

Abstract—The use of ChatGPT in academic activities may be perceived differently in relation to students' understanding and task efficiency. Differences in academic demands between working and non-working students may also lead to different patterns of use. This study aimed to describe ChatGPT Utilization, Quality of Understanding, and Task Efficiency, compare the three constructs according to student employment status, and compare Quality of Understanding and Task Efficiency within the same respondents. A quantitative descriptive-comparative design was employed. Data were collected from 101 Universitas Universal students selected through purposive sampling using a four-point Likert-scale questionnaire. Instrument evaluation resulted in 26 final items across three constructs, with Cronbach's Alpha values ranging from 0.787 to 0.913. The Mann–Whitney U test indicated no significant differences between working and non-working students in ChatGPT Utilization ($p=0.923$), Quality of Understanding ($p=0.244$), or Task Efficiency ($p=0.079$). The Wilcoxon Signed-Rank Test showed that the mean item score for Task Efficiency was higher than that for Quality of Understanding, at 3.260 and 3.000, respectively ($Z=-5.779$; $p<0.001$; $r=0.593$). Based on respondents' perceptions, ChatGPT use was more prominent in supporting practical and efficient task completion than in the quality of understanding, while student employment status did not produce meaningful differences across the three constructs. This study contributes empirical evidence by distinguishing utilization, understanding, and efficiency as separate aspects and provides a practical basis for higher education institutions to guide ChatGPT use while maintaining information verification and students' understanding processes.

Keywords: ChatGPT; Generative Artificial Intelligence; Quality Of Understanding; Task Efficiency; Student Employment Status

1. PENDAHULUAN

Perkembangan AI generatif mulai mengubah cara mahasiswa memperoleh dukungan belajar di perguruan tinggi. ChatGPT menjadi salah satu aplikasi yang sering digunakan karena mampu merespons pertanyaan berbasis teks, memberi gambaran awal, membantu penafsiran instruksi, dan menyediakan alternatif ide untuk tugas akademik [1]. Sejumlah kajian juga menunjukkan bahwa integrasi AI dalam pembelajaran berpotensi mendukung pengalaman belajar, produktivitas, dan hasil akademik mahasiswa [2], [3], [4]. Meskipun demikian, manfaat tersebut perlu dikaji lebih lanjut karena penggunaan ChatGPT tidak hanya berkaitan dengan kemudahan akses, tetapi juga dengan kualitas pemahaman dan efisiensi proses pengerjaan tugas.

Pada konteks akademik, ChatGPT dapat berfungsi sebagai pendamping belajar awal, misalnya untuk menjelaskan materi, membantu penyusunan kerangka tugas, dan mempercepat pencarian informasi pendukung [1], [5], [6]. Namun, penggunaan yang tidak kritis dapat menimbulkan masalah seperti informasi yang kurang akurat, ketergantungan, melemahnya proses berpikir mandiri, serta potensi pelanggaran integritas akademik [7], [8], [9], [10]. Oleh sebab itu, peran ChatGPT perlu ditempatkan sebagai alat pendukung, sementara validasi, analisis, dan pengambilan keputusan akademik tetap menjadi tanggung jawab mahasiswa.

Penelitian ini memisahkan dua bentuk keluaran penggunaan ChatGPT. Keluaran pertama adalah kualitas pemahaman akademik, yang berkaitan dengan kemampuan mahasiswa menangkap instruksi, memahami penjelasan, mengaitkan konsep, dan menafsirkan topik perkuliahan. Keluaran kedua adalah efisiensi penyelesaian tugas, yang berhubungan dengan penghematan waktu, keterarahan proses pengerjaan, dan kepraktisan dalam menyusun tugas. Pemisahan ini penting karena bantuan yang mempercepat pengerjaan belum tentu selalu menghasilkan pemahaman yang lebih mendalam.

Status kerja digunakan sebagai dasar perbandingan karena mahasiswa bekerja dan tidak bekerja menghadapi susunan tuntutan serta sumber daya belajar yang berbeda. Dalam Study Demands–Resources Theory, tuntutan akademik memerlukan waktu, perhatian, dan energi, sedangkan sumber daya belajar membantu mahasiswa mengelola tuntutan tersebut [11]. Mahasiswa yang bekerja juga perlu membagi sumber daya yang dimiliki antara peran pekerjaan dan peran akademik. Pengelolaan batas waktu antara kedua peran tersebut berkaitan dengan keseimbangan kerja–studi dan kepuasan akademik [12]. Dalam kondisi ini, ChatGPT dapat dipandang sebagai salah satu sumber daya pendukung yang mungkin digunakan untuk membantu pengelolaan informasi dan penyelesaian tugas. Oleh karena itu, status kerja relevan digunakan sebagai pembeda kelompok untuk melihat apakah perbedaan tuntutan tersebut diikuti oleh perbedaan Pemanfaatan ChatGPT, Kualitas Pemahaman, dan Efisiensi Tugas.

Sejumlah penelitian terdahulu telah membahas penggunaan ChatGPT dalam kegiatan akademik dari sudut pandang yang berbeda. Simamora et al. (2025) menilai peran ChatGPT sebagai alat bantu penyelesaian tugas mahasiswa, sedangkan Hakiki et al. (2023) menguji penggunaan ChatGPT dalam kaitannya dengan hasil belajar pada pembelajaran teknologi. Deng et al. (2025) merangkum bukti eksperimental mengenai manfaat ChatGPT terhadap pembelajaran mahasiswa melalui tinjauan sistematis dan meta-analisis. Sementara itu, Tompunu et al. (2025) secara khusus mengkaji pemanfaatan ChatGPT untuk membantu pencarian referensi dalam penyusunan tugas akhir. Keempat penelitian tersebut memberikan gambaran mengenai manfaat ChatGPT, tetapi lebih banyak berfokus pada efektivitas umum, hasil belajar, atau bentuk penggunaan tertentu [1], [3], [4], [5].

Berdasarkan perbandingan tersebut, masih terdapat ruang penelitian dalam memisahkan Pemanfaatan ChatGPT, Kualitas Pemahaman, dan Efisiensi Tugas sebagai tiga konstruk yang berbeda. Penelitian terdahulu yang ditelaah juga belum membandingkan ketiga konstruk tersebut berdasarkan status kerja mahasiswa serta belum menilai secara langsung apakah Kualitas Pemahaman atau Efisiensi Tugas memperoleh penilaian lebih tinggi pada responden yang sama. Penelitian ini mengisi kesenjangan tersebut melalui pendekatan deskriptif-komparatif pada mahasiswa bekerja dan tidak bekerja dalam satu lingkungan perguruan tinggi.

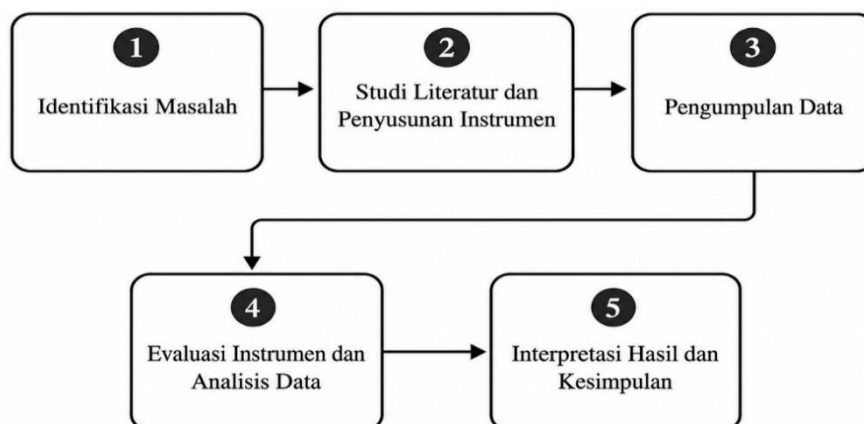
Penelitian ini bertujuan menggambarkan pemanfaatan ChatGPT, kualitas pemahaman, dan efisiensi tugas mahasiswa Universitas Universal, kemudian membandingkan ketiga konstruk tersebut berdasarkan status kerja mahasiswa. Penelitian ini juga membandingkan kualitas pemahaman dan efisiensi tugas pada responden yang sama untuk mengetahui aspek yang memperoleh penilaian lebih tinggi. Pembeda penelitian ini terletak pada pengukuran ketiga konstruk secara terpisah serta penggunaan status kerja sebagai dasar perbandingan dalam satu lingkungan perguruan tinggi. Kontribusi penelitian ini terletak pada penyajian bukti empiris yang memisahkan Pemanfaatan ChatGPT, Kualitas Pemahaman, dan Efisiensi Tugas sebagai tiga konstruk, membandingkannya berdasarkan status kerja, serta menilai perbedaan Kualitas Pemahaman dan Efisiensi Tugas pada responden yang sama. Temuan penelitian diharapkan dapat menjadi dasar praktis bagi perguruan tinggi dan dosen dalam menyusun pedoman penggunaan ChatGPT yang tidak hanya menekankan kemudahan pengerjaan tugas, tetapi juga verifikasi informasi dan proses pemahaman mahasiswa.

2. METODOLOGI PENELITIAN

2.1 Tahapan Penelitian

Penelitian dilaksanakan melalui beberapa tahapan yang saling berkaitan agar proses pengumpulan dan analisis data dapat dilakukan secara terarah. Tahap awal dilakukan dengan merumuskan masalah mengenai pemanfaatan ChatGPT, kualitas pemahaman, dan efisiensi tugas mahasiswa berdasarkan status kerja. Setelah itu, kajian literatur digunakan sebagai dasar dalam menentukan konstruk penelitian, menyusun indikator, dan merumuskan butir kuesioner yang sesuai dengan fokus penelitian. Instrumen kemudian disebar kepada mahasiswa Universitas Universal yang pernah menggunakan ChatGPT dalam kegiatan akademik. Data yang terkumpul diperiksa terlebih dahulu untuk memastikan kelengkapan dan konsistensi jawaban, kemudian disusun dan diolah menggunakan Microsoft Excel serta SPSS versi 26. Kelayakan instrumen dievaluasi melalui Exploratory Factor Analysis dan pengujian reliabilitas hingga diperoleh 26 item final yang digunakan dalam penghitungan skor setiap konstruk. Analisis selanjutnya meliputi statistik deskriptif, pemeriksaan ceiling effect, Harman's single-factor test, uji normalitas, Mann–Whitney U, Wilcoxon Signed-Rank Test, serta pemeriksaan perbedaan antara responden awal dan akhir. Setiap teknik analisis dipilih sesuai dengan tujuan pengujian, karakteristik data, serta bentuk perbandingan yang dilakukan dalam penelitian. Urutan analisis disusun mulai dari pemeriksaan kualitas instrumen dan distribusi data hingga pengujian perbedaan antarkelompok dan perbandingan antarkonstruk. Hasil dari setiap pengujian kemudian dirangkum untuk melihat pola jawaban responden, perbedaan berdasarkan status kerja, dan perbandingan antara kualitas pemahaman dan efisiensi tugas. Tahap terakhir dilakukan dengan menafsirkan hasil

secara deskriptif dan komparatif, menghubungkannya dengan tujuan penelitian, serta menyusun kesimpulan dan rekomendasi berdasarkan temuan yang diperoleh.



Gambar 1. Tahapan Penelitian

2.2 Desain Penelitian

Penelitian ini menggunakan pendekatan kuantitatif dengan desain deskriptif-komparatif. Analisis deskriptif digunakan untuk menggambarkan pemanfaatan ChatGPT, kualitas pemahaman, dan efisiensi tugas berdasarkan jawaban responden. Analisis komparatif diterapkan untuk membandingkan ketiga konstruk tersebut antara mahasiswa bekerja dan mahasiswa tidak bekerja [13]. Penelitian ini juga membandingkan kualitas pemahaman dan efisiensi tugas pada responden yang sama guna mengetahui aspek yang memperoleh penilaian lebih tinggi. Karena data diperoleh melalui survei pada satu periode pengukuran, hasil penelitian ditafsirkan sebagai gambaran persepsi dan perbedaan kelompok, bukan sebagai bukti hubungan sebab-akibat [14].

2.3 Populasi, Sampel, dan Teknik Pengumpulan Data

Populasi penelitian mencakup seluruh mahasiswa aktif Universitas Universal sebanyak 1.307 orang, perhitungan menggunakan rumus Slovin dengan tingkat kesalahan 10% menghasilkan kebutuhan minimum sekitar 93 responden, sedangkan data yang berhasil dianalisis mencakup 101 responden [15]. Pemilihan responden dilakukan melalui purposive sampling dengan kriteria mahasiswa yang pernah atau sedang menggunakan ChatGPT dalam kegiatan akademik yang dapat memberikan data sesuai fokus penelitian. Meskipun jumlah sampel telah melampaui kebutuhan minimum, penggunaan teknik nonprobabilitas membatasi generalisasi hasil di luar karakteristik responden dan lingkungan penelitian ini [16].

Data dikumpulkan melalui survei daring menggunakan Google Form pada 28 April–30 Mei 2026. Survei dipilih karena sesuai untuk memperoleh jawaban terstruktur dari responden dalam jumlah tertentu [17]. Untuk menjaga kualitas data, setiap akun hanya dapat mengisi satu kali dan data masuk diperiksa sebelum dianalisis. Skala yang digunakan adalah Likert empat poin, mulai dari sangat tidak setuju hingga sangat setuju, agar responden memberikan kecenderungan sikap yang lebih jelas tanpa pilihan netral [15]. Namun, keputusan menghilangkan pilihan netral tetap dipertimbangkan sebagai keterbatasan karena keberadaan atau ketiadaan kategori tengah dapat memengaruhi pola respons dan karakteristik pengukuran [18].

2.4 Variabel dan Instrumen Penelitian

Instrumen awal disusun untuk mengukur tiga konstruk penelitian, yaitu pemanfaatan ChatGPT, kualitas pemahaman, dan efisiensi tugas. Setiap konstruk pada awalnya terdiri atas sepuluh pernyataan, sehingga jumlah keseluruhan instrumen mencapai 30 item. Penyusunan butir mempertimbangkan bentuk penggunaan ChatGPT dalam kegiatan akademik, pemahaman terhadap materi dan tugas, serta kemudahan dan kecepatan proses penyelesaian tugas [4], [9]. Sebelum kuesioner disebarkan, redaksi dan kesesuaian isi setiap pernyataan dikonsultasikan kepada dosen pembimbing.

Setelah data terkumpul, struktur instrumen diperiksa melalui Exploratory Factor Analysis dan konsistensi internalnya dinilai menggunakan Cronbach's Alpha. Hasil evaluasi menghasilkan 26 item final, yang terdiri atas delapan item Pemanfaatan ChatGPT (X), delapan item Kualitas Pemahaman (Y1), dan sepuluh item Efisiensi Tugas (Y2). Item X1.8, X1.9, Y1.6, dan Y1.7 tidak disertakan dalam penghitungan skor final karena kurang mendukung pemisahan konstruk berdasarkan hasil analisis faktor. Status kerja mahasiswa dikelompokkan menjadi bekerja dan tidak bekerja serta digunakan sebagai dasar perbandingan antarkelompok.

Tabel 1. Variabel Pemanfaatan ChatGPT

Variabel	Kode	Pertanyaan
Pemanfaatan	X1.1	ChatGPT membantu saya dalam proses menyelesaikan tugas akademik

Variabel	Kode	Pertanyaan
Pemanfaatan	X1.2	ChatGPT mudah digunakan dalam mendukung penyelesaian tugas akademik saya
Pemanfaatan	X1.3	ChatGPT membantu saya memahami instruksi tugas dengan lebih jelas
Pemanfaatan	X1.4	ChatGPT memberikan jawaban yang relevan dengan tugas akademik saya
Pemanfaatan	X1.5	ChatGPT membantu saya memperoleh ide dalam penyelesaian tugas akademik
Pemanfaatan	X1.6	ChatGPT membantu memperoleh informasi awal yang berkaitan dengan tugas
Pemanfaatan	X1.7	ChatGPT membantu saya menyelesaikan tugas akademik secara lebih terarah
Pemanfaatan	X1.10	ChatGPT membantu menjelaskan kembali materi dengan bahasa yang lebih sederhana
Pemanfaatan	X1	Total Variabel X

Tabel 2. Variabel Kualitas Pemahaman

Variabel	Kode	Pertanyaan
Kualitas	Y1.1	Informasi yang diberikan ChatGPT akurat dan sesuai dengan sumber yang relevan
Kualitas	Y1.2	ChatGPT membantu saya memahami konsep materi dengan lebih mudah
Kualitas	Y1.3	ChatGPT mampu menganalisis pertanyaan atau masalah dengan baik
Kualitas	Y1.4	Jawaban dari ChatGPT sesuai dengan konteks tugas yang saya kerjakan
Kualitas	Y1.5	Penjelasan dari ChatGPT dapat diterapkan / diaplikasikan dalam tugas saya
Kualitas	Y1.8	ChatGPT mampu memberikan ide atau solusi dalam mengerjakan tugas
Kualitas	Y1.9	ChatGPT mampu menyusun materi dan jawaban tugas akademik secara sistematis
Kualitas	Y1.10	ChatGPT mampu menghubungkan berbagai konsep dalam materi
Kualitas	Y1	Total Variabel Y1

Tabel 3. Variabel Efisiensi Tugas

Variabel	Kode	Pertanyaan
Efisiensi	Y2.1	Penggunaan ChatGPT menghemat waktu dalam menyelesaikan tugas akademik
Efisiensi	Y2.2	Menyelesaikan lebih banyak tugas dalam waktu yang lebih singkat dengan ChatGPT
Efisiensi	Y2.3	ChatGPT meningkatkan produktivitas saya dalam belajar dan mengerjakan tugas
Efisiensi	Y2.4	ChatGPT membantu saya menemukan jawaban dengan cepat
Efisiensi	Y2.5	ChatGPT mengurangi waktu yang saya butuhkan untuk mencari referensi
Efisiensi	Y2.6	ChatGPT mempermudah proses penyusunan tugas akademik
Efisiensi	Y2.7	ChatGPT membuat proses pengerjaan tugas menjadi lebih praktis
Efisiensi	Y2.8	Saya merasa tugas menjadi lebih ringan dengan bantuan ChatGPT
Efisiensi	Y2.9	ChatGPT membantu menyederhanakan tugas yang kompleks
Efisiensi	Y2.10	Penggunaan ChatGPT mengurangi tingkat kesulitan dalam menyelesaikan tugas
Efisiensi	Y2	Total Variabel Y2

2.5 Teknik Analisis Data

Data penelitian diolah menggunakan Microsoft Excel dan SPSS versi 26. Analisis deskriptif dilakukan dengan menghitung jumlah data, mean, median, simpangan baku, nilai minimum, nilai maksimum, dan skewness pada setiap konstruk. Pemeriksaan ceiling effect dilakukan dengan menghitung proporsi responden yang memperoleh skor maksimum teoretis. Indikasi ceiling effect dinilai kuat apabila lebih dari 15% responden mencapai skor maksimum [19].

Kelayakan struktur instrumen diperiksa menggunakan Exploratory Factor Analysis dengan metode ekstraksi Principal Axis Factoring dan rotasi Direct Oblimin. Jumlah faktor ditetapkan sebanyak tiga sesuai dengan konstruk Pemanfaatan ChatGPT, Kualitas Pemahaman, dan Efisiensi Tugas. Kecukupan data dievaluasi melalui Kaiser-Meyer-Olkin dan Bartlett's Test of Sphericity, sedangkan keputusan mengenai item mempertimbangkan factor loading, communality, kesesuaian konstruk, dan konsistensi internal [20]. Reliabilitas setiap konstruk kemudian dihitung menggunakan Cronbach's Alpha. Kemungkinan common method bias diperiksa melalui Harman's single-factor test dengan mengekstraksi satu faktor tanpa rotasi. Pengujian tersebut digunakan sebagai pemeriksaan awal dan tidak diperlakukan sebagai bukti bahwa common method bias sepenuhnya tidak ada [21].

Distribusi skor diperiksa menggunakan Shapiro-Wilk pada taraf signifikansi 0,05. Karena hasil pengujian menunjukkan bahwa seluruh konstruk tidak secara seragam memenuhi asumsi normalitas, analisis komparatif menggunakan metode nonparametrik. Mann-Whitney U digunakan untuk membandingkan Pemanfaatan ChatGPT, Kualitas Pemahaman, dan Efisiensi Tugas antara mahasiswa bekerja dan tidak bekerja. Wilcoxon Signed-Rank Test digunakan untuk membandingkan skor rata-rata Kualitas Pemahaman dan Efisiensi Tugas pada responden yang sama. Pemilihan kedua uji tersebut disesuaikan dengan karakteristik kelompok independen dan data berpasangan serta kondisi distribusi data [13]. Besarnya perbedaan dilengkapi dengan effect size yang dihitung menggunakan rumus.

$$r = \frac{|Z|}{\sqrt{N}} \quad (1)$$

Keterangan: r = effect size, $|Z|$ = nilai absolut statistic Z, N = Jumlah observasi yang digunakan dalam perhitungan

Potensi non-response bias diperiksa melalui pendekatan early-late response dengan membandingkan 30 responden pertama dan 30 responden terakhir berdasarkan urutan waktu pengisian kuesioner. Perbedaan skor ketiga konstruk diuji menggunakan Mann-Whitney U. Komposisi status kerja pada kelompok awal dan akhir juga diperiksa menggunakan tabulasi silang dan Chi-Square agar perbedaan waktu respons dapat ditafsirkan dengan mempertimbangkan karakteristik kelompok. Hasil pemeriksaan ini diperlakukan sebagai indikasi potensi bias respons, bukan sebagai pembuktian langsung terhadap karakteristik individu yang sama sekali tidak memberikan respons [22].

3. HASIL DAN PEMBAHASAN

3.1 Hasil Penelitian

Bagian ini menyajikan hasil analisis terhadap data dari 101 mahasiswa Universitas Universal yang menggunakan ChatGPT dalam kegiatan akademik. Penyajian hasil diawali dengan karakteristik responden, dilanjutkan dengan evaluasi instrumen, pemeriksaan kemungkinan bias metode, statistik deskriptif, distribusi data, serta perbandingan berdasarkan status kerja. Analisis tambahan dilakukan untuk membandingkan kualitas pemahaman dan efisiensi tugas pada responden yang sama serta memeriksa kemungkinan perbedaan antara responden awal dan akhir.

3.1.1 Karakteristik Responden

Responden dikelompokkan berdasarkan status kerja mahasiswa yaitu bekerja dan tidak bekerja. Dari 101 responden, 56 orang atau 55,4% termasuk mahasiswa bekerja, sedangkan 45 orang atau 44,6% merupakan mahasiswa tidak bekerja. Kedua kelompok tersebut selanjutnya digunakan dalam analisis deskriptif dan pengujian Mann-Whitney U untuk memeriksa perbedaan Pemanfaatan ChatGPT, Kualitas Pemahaman, dan Efisiensi Tugas berdasarkan status kerja.

Tabel 4. Karakteristik Responden

Status Mahasiswa	Frekuensi	Persentase
Bekerja	56	55,45%
Tidak Bekerja	45	44,55%

3.1.2 Hasil Evaluasi Instrumen

Evaluasi instrumen dilakukan terhadap 30 item awal untuk memeriksa kesesuaian struktur pengukuran dengan tiga konstruk penelitian. Exploratory Factor Analysis menggunakan metode ekstraksi Principal Axis Factoring dan rotasi Direct Oblimin [20]. Hasil pengujian akhir menunjukkan nilai Kaiser-Meyer-Olkin sebesar 0,841 dan Bartlett's Test of Sphericity dengan signifikansi $p < 0,001$. Hasil tersebut menunjukkan bahwa pola korelasi antarpbutir memadai untuk dianalisis lebih lanjut melalui analisis faktor.

Proses evaluasi menghasilkan tiga faktor yang sesuai dengan konstruk Pemanfaatan ChatGPT, Kualitas Pemahaman, dan Efisiensi Tugas. Ketiga faktor tersebut menjelaskan varians kumulatif sebesar 44,93%. Berdasarkan factor loading, communality, struktur faktor, dan kesesuaian teoritis, item X1.8, X1.9, Y1.6, dan Y1.7 tidak disertakan dalam instrumen final. Dengan demikian, analisis selanjutnya menggunakan 26 item yang terdiri atas delapan item Pemanfaatan ChatGPT (X), delapan item Kualitas Pemahaman (Y1), dan sepuluh item Efisiensi Tugas (Y2).

Tabel 5. Ringkasan Hasil Exploratory Factor Analysis

Komponen Evaluasi	Hasil
Kaiser-Meyer-Olkin (KMO)	0,841
Bartlett's Test of Sphericity	$p < 0,001$
Metode ekstraksi	Principal Axis Factoring
Metode rotasi	Direct Oblimin
Jumlah faktor yang dipertahankan	3 faktor
Varians kumulatif	44,93%
Item yang dikeluarkan	X1.8, X1.9, Y1.6, dan Y1.7
Jumlah item final	26 item

Hasil rotasi menunjukkan bahwa item Efisiensi Tugas membentuk faktor pertama dengan loading antara 0,416 dan 0,804. Item Pemanfaatan ChatGPT mengelompok pada faktor kedua dengan loading antara 0,481 dan 0,658, sedangkan sebagian besar item Kualitas Pemahaman membentuk faktor ketiga dengan loading antara 0,423 dan 0,774. Pola tersebut menunjukkan bahwa ketiga konstruk dapat dibedakan berdasarkan susunan item yang membentuk setiap faktor.

Item Y1.2 tetap dipertahankan dalam konstruk Kualitas Pemahaman meskipun loading pada Pattern Matrix berada di bawah 0,40. Keputusan tersebut mempertimbangkan kesesuaian isi item dengan konsep pemahaman

materi, nilai communality sebesar 0,360, dan loading pada Structure Matrix sebesar 0,554. Dengan pertimbangan statistik dan teoritis tersebut, Y1.2 masih dinilai relevan untuk digunakan dalam instrumen final.

Tabel 6. Struktur Faktor Instrumen Final

Faktor	Konstruk	Kode Item	Rentang Loading	Jumlah Item
1	Efisiensi Tugas	Y2.1-Y2.10	0,416 - 0,804	10
2	Pemanfaatan ChatGPT	X1.1-X1.7 dan X1.10	0,481 - 0,658	8
3	Kualitas Pemahaman	Y1.1, Y1.3-Y1.5, dan Y1.8-Y1.10	0,423 - 0,774	7
3	Kualitas Pemahaman	Y1.2	<0,40	1

Nilai Cronbach's Alpha pada Pemanfaatan ChatGPT sebesar 0,787, Kualitas Pemahaman sebesar 0,851, dan Efisiensi Tugas sebesar 0,913. Seluruh nilai berada di atas 0,70 [23], sehingga masing-masing konstruk memiliki konsistensi internal yang dapat diterima dan dapat digunakan dalam analisis berikutnya.

Tabel 7. Komposisi dan Reliabilitas Instrumen Final

Konstruk	Kode Item Final	Jumlah Item	Cronbach's Alpha	Keterangan
Pemanfaatan ChatGPT	X1.1-X1.7 dan X1.10	8	0,787	Reliabel
Kualitas Pemahaman	Y1.1-Y1.5, dan Y1.8-Y1.10	8	0,851	Reliabel
Efisiensi Tugas	Y2.1-Y2.10	10	0,913	Reliabel

3.1.3 Pemeriksaan Common Method Bias

Kemungkinan common method bias diperiksa karena seluruh data konstruk diperoleh dari responden yang sama melalui satu kuesioner. Pemeriksaan dilakukan menggunakan Harman's single-factor test terhadap 26 item final dengan mengekstraksi satu faktor tanpa rotasi. Pengujian ini digunakan untuk melihat apakah satu faktor umum mendominasi variasi jawaban responden.

Hasil pengujian menunjukkan bahwa satu faktor hanya menjelaskan 32,21% dari keseluruhan varians. Nilai tersebut menunjukkan bahwa tidak terdapat satu faktor umum yang mendominasi jawaban seluruh item. Oleh karena itu, tidak ditemukan indikasi kuat common method bias pada data penelitian. Meskipun demikian, hasil Harman's single-factor test hanya diperlakukan sebagai pemeriksaan awal dan tidak dapat memastikan bahwa seluruh kemungkinan bias metode telah sepenuhnya dihilangkan [21].

3.1.4 Statistik Deskriptif dan Ceiling Effect

Statistik deskriptif digunakan untuk menggambarkan distribusi skor pada konstruk Pemanfaatan ChatGPT (X), Kualitas Pemahaman (Y1), dan Efisiensi Tugas (Y2). Pemeriksaan meliputi nilai mean, median, simpangan baku, skor minimum, skor maksimum, dan skewness. Selain itu, ceiling effect diperiksa berdasarkan jumlah responden yang mencapai skor maksimum teoretis pada masing-masing konstruk.

Pemanfaatan ChatGPT memperoleh mean sebesar 26,11 dengan median 25 dan simpangan baku 3,35. Kualitas Pemahaman memiliki mean sebesar 24,00, median 24, dan simpangan baku 3,75. Sementara itu, Efisiensi Tugas memperoleh mean sebesar 32,60 dengan median 32 dan simpangan baku 5,18. Nilai skewness ketiga konstruk berada antara -0,352 dan 0,217, yang menunjukkan bahwa distribusi skor tidak memiliki kemencengan yang sangat besar.

Sebanyak sembilan responden atau 8,91% mencapai skor maksimum pada Pemanfaatan ChatGPT, tidak terdapat responden yang mencapai skor maksimum pada Kualitas Pemahaman, dan sepuluh responden atau 9,90% mencapai skor maksimum pada Efisiensi Tugas. Seluruh persentase berada di bawah batas 15%, sehingga tidak ditemukan ceiling effect yang kuat pada ketiga konstruk. Hasil tersebut menunjukkan bahwa skor responden masih memiliki variasi yang memadai dan tidak terkonsentrasi secara berlebihan pada nilai maksimum [19].

Tabel 8. Hasil Deskriptif konstruk Penelitian

Kode Variabel	N	Mean	Median	SD	Minimum	Maksimum	Skewness	Ceiling Effect
X	101	26,11	25	3,35	20	32	0,217	9 (8,91%)
Y1	101	24,00	24	3,75	14	31	-0,231	0 (0,00%)
Y2	101	32,60	32	5,18	17	40	-0,352	10 (9,90%)

3.1.5 Hasil Uji Normalitas

Uji normalitas dilakukan menggunakan Shapiro-Wilk dengan taraf signifikansi 0,05. Data dinyatakan mengikuti distribusi normal apabila nilai signifikansi lebih besar dari 0,05, sedangkan nilai signifikansi kurang dari atau sama dengan 0,05 menunjukkan bahwa distribusi data tidak normal. Hasil pengujian terhadap ketiga konstruk disajikan pada Tabel 9.

Tabel 9. Hasil Uji Normalitas Shapiro-Wilk

Konstruk	Nilai Signifikansi	Kriteria	Keterangan
Pemanfaatan ChatGPT	$p < 0,001$	$p > 0,05$	Tidak normal
Kualitas Pemahaman	$p = 0,107$	$p > 0,05$	Normal
Efisiensi Tugas	$p = 0,001$	$p > 0,05$	Tidak normal

Hasil Shapiro-Wilk menunjukkan bahwa Pemanfaatan ChatGPT memiliki nilai $p < 0,001$ dan Efisiensi Tugas memiliki nilai $p = 0,001$, sehingga kedua konstruk tidak memenuhi asumsi normalitas. Kualitas Pemahaman memperoleh nilai $p = 0,107$ dan dinyatakan berdistribusi normal. Karena distribusi ketiga konstruk tidak seluruhnya memenuhi asumsi normalitas, analisis perbandingan selanjutnya menggunakan metode nonparametrik, yaitu Mann-Whitney U untuk kelompok independen dan Wilcoxon Signed-Rank Test untuk data berpasangan [13].

3.2 Perbandingan Berdasarkan Status Kerja

Perbandingan berdasarkan status kerja dilakukan terhadap kelompok mahasiswa bekerja sebanyak 56 responden dan mahasiswa tidak bekerja sebanyak 45 responden. Analisis diawali dengan membandingkan nilai mean setiap konstruk, kemudian dilanjutkan menggunakan Mann-Whitney U untuk mengetahui apakah terdapat perbedaan yang signifikan antara kedua kelompok.

Tabel 10. Mean Konstruk Berdasarkan Status Kerja Mahasiswa

Konstruk	(n = 56) Bekerja	(n = 45) Tidak Bekerja
Pemanfaatan ChatGPT	26,05	26,18
Kualitas Pemahaman	24,38	23,53
Efisiensi Tugas	33,48	31,51

Mean Pemanfaatan ChatGPT pada mahasiswa bekerja sebesar 26,05 dan pada mahasiswa tidak bekerja sebesar 26,18, sehingga kedua kelompok menunjukkan skor yang relatif serupa. Kualitas Pemahaman pada mahasiswa bekerja memiliki mean sebesar 24,38, lebih tinggi daripada mahasiswa tidak bekerja sebesar 23,53. Efisiensi Tugas juga menunjukkan kecenderungan lebih tinggi pada mahasiswa bekerja, yaitu 33,48 dibandingkan 31,51 pada mahasiswa tidak bekerja. Perbedaan deskriptif tersebut selanjutnya diuji menggunakan Mann-Whitney U untuk menentukan signifikansi dan besarnya perbedaan antarkelompok.

Tabel 11. Hasil Uji Mann-Whitney U Berdasarkan Status Kerja

Konstruk	U	Z	p	Effect Size (r)	Keterangan
Pemanfaatan ChatGPT	1.246	-0,096	0,923	0,010	Tidak signifikan
Kualitas Pemahaman	1.090	-1,166	0,244	0,116	Tidak signifikan
Efisiensi Tugas	1.004	-1,755	0,079	0,175	Tidak signifikan

Hasil Mann-Whitney U menunjukkan bahwa Pemanfaatan ChatGPT memperoleh nilai ($U=1.246$), ($Z=-0,096$), dan ($p=0,923$). Kualitas Pemahaman memperoleh nilai ($U=1.090$), ($Z=-1,166$), dan ($p=0,244$), sedangkan Efisiensi Tugas memperoleh nilai ($U=1.004$), ($Z=-1,755$), dan ($p=0,079$). Nilai (U) merupakan statistik pengujian yang dihitung berdasarkan peringkat skor kedua kelompok, sedangkan nilai (Z) menunjukkan bentuk standardisasi dari hasil pengujian. Nilai (p) digunakan untuk menentukan signifikansi perbedaan antarkelompok. Karena seluruh nilai (p) berada di atas 0,05, tidak terdapat perbedaan yang signifikan antara mahasiswa bekerja dan tidak bekerja pada ketiga konstruk.

Nilai effect size Pemanfaatan ChatGPT sebesar 0,010 menunjukkan ukuran perbedaan yang sangat kecil. Kualitas Pemahaman memiliki effect size sebesar 0,116 dan Efisiensi Tugas sebesar 0,175, yang keduanya termasuk dalam kategori kecil. Dengan demikian, kecenderungan mean Kualitas Pemahaman dan Efisiensi Tugas yang lebih tinggi pada mahasiswa bekerja belum menunjukkan perbedaan yang bermakna secara statistik maupun praktis.

3.3 Perbandingan Kualitas Pemahaman dan Efisiensi Tugas

Kualitas Pemahaman dan Efisiensi Tugas dibandingkan menggunakan nilai mean per item karena kedua konstruk memiliki jumlah item yang berbeda. Kualitas Pemahaman terdiri atas delapan item, sedangkan Efisiensi Tugas terdiri atas sepuluh item. Hasil deskriptif menunjukkan mean per item Kualitas Pemahaman sebesar 3,000 dan Efisiensi Tugas sebesar 3,260. Perbedaan kedua skor pada responden yang sama selanjutnya diuji menggunakan Wilcoxon Signed-Rank Test [13].

Tabel 12. Hasil Perbandingan Kualitas Pemahaman (Y1) dan Efisiensi Tugas (Y2)

Komponen	Hasil
Mean (Y1) Kualitas Pemahaman	3,000
Mean (Y2) Efisiensi Tugas	3,260

Komponen	Hasil
Negative ranks	22
Positive ranks	73
Ties	6
Z	-5,779
p	< 0,001
Effect size (r)	0,593

Keterangan: Negative ranks ($Y_2 < Y_1$), Positive ranks ($Y_2 > Y_1$), Ties ($Y_1 = Y_2$)

Hasil Wilcoxon menunjukkan bahwa 73 responden memiliki skor Efisiensi Tugas yang lebih tinggi daripada Kualitas Pemahaman, sedangkan 22 responden menunjukkan pola sebaliknya dan enam responden memiliki skor yang sama. Nilai ($Z = -5,779$) dengan ($p < 0,001$) menunjukkan bahwa terdapat perbedaan yang signifikan antara kedua konstruk. Arah peringkat dan mean per item menunjukkan bahwa Efisiensi Tugas memperoleh penilaian lebih tinggi daripada Kualitas Pemahaman.

Nilai effect size sebesar 0,593 termasuk dalam kategori besar karena nilai r yang mencapai atau melebihi 0,50 umumnya diinterpretasikan sebagai ukuran efek besar [24]. Hasil tersebut menunjukkan bahwa perbedaan antara Kualitas Pemahaman dan Efisiensi Tugas tidak hanya signifikan secara statistik, tetapi juga memiliki besaran perbedaan yang kuat. Berdasarkan persepsi responden, pemanfaatan ChatGPT lebih menonjol pada Efisiensi Tugas dibandingkan pada Kualitas Pemahaman.

3.4 Pemeriksaan Potensi Non-response Bias

Pemeriksaan potensi non-response bias dilakukan melalui pendekatan early-late response dengan membandingkan 30 responden pertama dan 30 responden terakhir berdasarkan urutan waktu pengisian kuesioner. Perbandingan skor Pemanfaatan ChatGPT, Kualitas Pemahaman, dan Efisiensi Tugas dilakukan menggunakan Mann-Whitney U. Pendekatan ini digunakan sebagai pemeriksaan awal terhadap kemungkinan perbedaan pola respons antara responden awal dan akhir, bukan sebagai pengukuran langsung terhadap individu yang tidak memberikan respons [22].

Tabel 13. Perbandingan Responden Awal dan Akhir

Konstruk	p	Keterangan
Pemanfaatan ChatGPT	0,681	Tidak berbeda signifikan
Kualitas Pemahaman	0,002	Berbeda signifikan
Efisiensi Tugas	0,002	Berbeda signifikan

Hasil pengujian menunjukkan bahwa Pemanfaatan ChatGPT tidak berbeda secara signifikan antara responden awal dan akhir karena memperoleh nilai ($p = 0,681$). Sebaliknya, Kualitas Pemahaman dan Efisiensi Tugas masing-masing memperoleh nilai ($p = 0,002$), sehingga terdapat perbedaan yang signifikan antara kedua kelompok waktu respons. Besarnya perbedaan pada Kualitas Pemahaman ditunjukkan oleh nilai ($r = 0,393$), sedangkan Efisiensi Tugas memiliki nilai ($r = 0,405$).

Tabel 14. Komposisi Status Kerja Responden Awal dan Akhir

Status Kerja	Responden Awal	Responden Akhir
Bekerja	21	9
Tidak Bekerja	9	21
Total	30	30

Komposisi status kerja menunjukkan bahwa kelompok responden awal lebih banyak terdiri atas mahasiswa bekerja, sedangkan kelompok responden akhir lebih banyak terdiri atas mahasiswa tidak bekerja. Hasil Chi-Square menunjukkan bahwa perbedaan komposisi tersebut signifikan dengan nilai ($X^2 = 9,60$) dan ($p = 0,002$). Oleh karena itu, perbedaan Kualitas Pemahaman dan Efisiensi Tugas antara responden awal dan akhir belum dapat ditafsirkan sebagai non-response bias murni karena kemungkinan berkaitan dengan ketidakseimbangan status kerja pada kedua kelompok. Hasil ini menunjukkan adanya indikasi potensi bias respons yang perlu dipertimbangkan sebagai keterbatasan penelitian [22].

3.5 Pembahasan

Evaluasi instrumen menghasilkan 26 item yang terbagi ke dalam tiga faktor, yaitu Pemanfaatan ChatGPT, Kualitas Pemahaman, dan Efisiensi Tugas. Nilai KMO sebesar 0,841 dan Bartlett's Test of Sphericity yang signifikan menunjukkan bahwa hubungan antaritem mencukupi untuk dianalisis menggunakan EFA [20]. Ketiga konstruk juga memiliki Cronbach's Alpha di atas 0,70, yaitu 0,787, 0,851, dan 0,913. Hasil tersebut menunjukkan bahwa setiap konstruk mempunyai konsistensi internal yang memadai [23]. Pemeriksaan Harman's single-factor test memperlihatkan bahwa faktor tunggal hanya menjelaskan 32,21% varians. Dengan demikian, tidak terlihat adanya satu faktor umum yang mendominasi seluruh jawaban, meskipun kemungkinan common method bias tidak dapat sepenuhnya dikesampingkan [21].

Secara deskriptif, mahasiswa menunjukkan kecenderungan positif dalam memanfaatkan ChatGPT untuk kegiatan akademik. Aplikasi tersebut digunakan untuk memperoleh gambaran awal, memahami instruksi, mencari gagasan, serta membantu penyusunan tugas. Temuan ini sejalan dengan kajian yang menjelaskan bahwa AI generatif dapat mendukung pencarian informasi, pemberian penjelasan, dan aktivitas belajar mahasiswa [4], [9]. Mean total ketiga konstruk tidak dibandingkan secara langsung karena jumlah itemnya berbeda. Selain itu, persentase responden yang mencapai skor maksimum pada setiap konstruk berada di bawah 15%, sehingga tidak ditemukan penumpukan jawaban yang kuat pada batas atas skala [19].

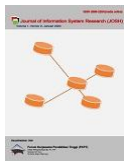
Perbandingan berdasarkan status kerja menunjukkan bahwa Pemanfaatan ChatGPT relatif serupa pada mahasiswa bekerja dan tidak bekerja. Kualitas Pemahaman dan Efisiensi Tugas memang cenderung lebih tinggi pada mahasiswa bekerja, tetapi perbedaannya tidak signifikan dan ukuran efeknya sangat kecil hingga kecil. Secara teoritis, mahasiswa bekerja menghadapi tambahan tuntutan waktu dan pembagian peran yang dapat memengaruhi cara mereka mengelola kegiatan akademik [11], [12]. Namun, hasil penelitian ini menunjukkan bahwa perbedaan tuntutan tersebut tidak diikuti oleh perbedaan yang cukup kuat pada ketiga konstruk. Temuan ini mengindikasikan bahwa ChatGPT dapat dimanfaatkan sebagai sumber daya pendukung oleh kedua kelompok, meskipun kebutuhan dan situasi penggunaannya mungkin berbeda.

Hasil Wilcoxon menunjukkan bahwa mean per item Efisiensi Tugas sebesar 3,260 lebih tinggi daripada Kualitas Pemahaman sebesar 3,000. Perbedaan tersebut signifikan dengan ($Z=-5,779$), ($p<0,001$), dan effect size sebesar 0,593. Berdasarkan persepsi responden, manfaat ChatGPT lebih menonjol dalam membantu mempercepat dan mempermudah penyelesaian tugas dibandingkan dalam mendukung pemahaman materi. Temuan ini perlu dipahami secara hati-hati karena pengerjaan yang lebih cepat belum tentu diikuti oleh pemahaman yang lebih mendalam. Kualitas penggunaan AI tetap dipengaruhi oleh kemampuan mahasiswa dalam menyusun instruksi, memeriksa informasi, dan menilai keluaran secara kritis [4], [10].

Pemeriksaan responden awal dan akhir menemukan perbedaan pada Kualitas Pemahaman dan Efisiensi Tugas, tetapi tidak pada Pemanfaatan ChatGPT. Komposisi status kerja kedua kelompok juga berbeda secara signifikan. Karena itu, hasil tersebut belum dapat langsung dianggap sebagai bukti non-response bias dan lebih tepat diperlakukan sebagai indikasi potensi bias respons [22]. Penelitian ini juga terbatas pada satu perguruan tinggi dan menggunakan purposive sampling, sehingga generalisasinya perlu dilakukan secara terbatas [16]. Selain itu, penggunaan kuesioner laporan diri memungkinkan responden memberikan jawaban yang dinilai lebih positif atau dapat diterima secara sosial, sehingga kemungkinan social desirability bias belum dapat sepenuhnya dikesampingkan [25]. Penggunaan skala Likert empat poin tanpa pilihan netral dapat pula memengaruhi kecenderungan jawaban responden [18]. Dengan demikian, hasil penelitian menggambarkan persepsi dan perbedaan kelompok, bukan membuktikan hubungan sebab-akibat.

4. KESIMPULAN

Penelitian yang melibatkan 101 mahasiswa Universitas Universal ini menggunakan 26 item final untuk mengukur Pemanfaatan ChatGPT, Kualitas Pemahaman, dan Efisiensi Tugas, dengan konsistensi internal yang memadai pada seluruh konstruk. Secara deskriptif, mahasiswa cenderung menilai ChatGPT bermanfaat untuk memperoleh informasi awal, memahami instruksi, menemukan gagasan, dan mempermudah penyelesaian tugas, serta tidak ditemukan ceiling effect yang kuat. Hasil Mann-Whitney U menunjukkan bahwa ketiga konstruk tidak berbeda secara signifikan antara mahasiswa bekerja dan tidak bekerja, meskipun Kualitas Pemahaman dan Efisiensi Tugas memiliki kecenderungan lebih tinggi pada mahasiswa bekerja dengan ukuran perbedaan yang kecil. Hasil Wilcoxon Signed-Rank Test menunjukkan bahwa Efisiensi Tugas memperoleh penilaian lebih tinggi daripada Kualitas Pemahaman dengan ($p<0,001$) dan effect size sebesar 0,593, sehingga berdasarkan persepsi responden, manfaat ChatGPT lebih menonjol pada kepraktisan dan efisiensi pengerjaan tugas daripada pada pemahaman materi. Kontribusi empiris penelitian ini terletak pada pemisahan Pemanfaatan ChatGPT, Kualitas Pemahaman, dan Efisiensi Tugas sebagai tiga konstruk yang berbeda, serta pembuktiannya melalui perbandingan berdasarkan status kerja dan pengujian pada responden yang sama. Secara praktis, temuan ini dapat menjadi dasar bagi perguruan tinggi untuk mengarahkan penggunaan ChatGPT agar tidak hanya berorientasi pada kecepatan penyelesaian tugas, tetapi juga pada verifikasi informasi dan penguatan pemahaman mahasiswa. Temuan ini perlu ditafsirkan secara terbatas karena penelitian menggunakan purposive sampling, hanya dilakukan pada satu perguruan tinggi, dan masih menunjukkan indikasi potensi bias respons. Penelitian berikutnya disarankan melibatkan sampel yang lebih luas dan menggunakan ukuran objektif, seperti mutu tugas, hasil belajar, literasi AI, intensitas penggunaan, dan kemampuan berpikir kritis. Universitas Universal juga dapat menyusun pedoman penggunaan ChatGPT yang mencakup verifikasi informasi, keterbukaan penggunaan AI, dan tanggung jawab akademik, disertai kegiatan refleksi, penjelasan lisan, atau pemeriksaan sumber oleh dosen agar kemudahan pengerjaan tugas tetap berjalan bersama proses pemahaman. Karena tidak ditemukan perbedaan berdasarkan status kerja, pedoman dan pelatihan literasi AI tersebut dapat diterapkan kepada seluruh mahasiswa dengan akses yang fleksibel bagi mahasiswa bekerja maupun tidak bekerja.



REFERENCES

- [1] Eka Finanti Simamora, Imel Simanungkalit, Prihatin Ningsih Sagala, Nurcahya Br Zandroto, Putri Br Tarigan, and Rival Ananda Gisty, “Efektivitas Peran ChatGPT Sebagai Alat Bantu Penyelesaian Tugas Akademik Mahasiswa,” *Algoritma : Jurnal Matematika, Ilmu pengetahuan Alam, Kebumian dan Angkasa*, vol. 3, no. 2, pp. 74–85, Mar. 2025, doi: 10.62383/algoritma.v3i2.445.
- [2] J. Li, K. Yin, Y. Wang, X. Jiang, and D. Chen, “Effectiveness of generative artificial intelligence-based teaching versus traditional teaching methods in medical education: a meta-analysis of randomized controlled trials,” *BMC Med. Educ.*, vol. 25, no. 1, Dec. 2025, doi: 10.1186/s12909-025-07750-2.
- [3] M. Hakiki et al., “Exploring the impact of using Chat-GPT on student learning outcomes in technology learning: The comprehensive experiment,” *Advances in Mobile Learning Educational Research*, vol. 3, no. 2, pp. 859–872, Oct. 2023, doi: 10.25082/amler.2023.02.013.
- [4] R. Deng, M. Jiang, X. Yu, Y. Lu, and S. Liu, “Does ChatGPT enhance student learning? A systematic review and meta-analysis of experimental studies,” *Comput. Educ.*, vol. 227, Apr. 2025, doi: 10.1016/j.compedu.2024.105224.
- [5] K. Angelina Tomponu et al., “ChatGPT Sebagai Alat Bantu Pencarian Referensi: Analisis Penggunaan oleh Mahasiswa Sistem Informasi UNSRI dalam Menyusun Tugas Akhir,” 2025. doi: 10.36040/jati.v9i4.14062.
- [6] A. Azizi, R. Joko Musridho, S. Wiruddin, and E. Dwi Mustika, “Pengaruh Penggunaan ChatGPT Terhadap Efektivitas Pembelajaran Mahasiswa Teknik Informatika di Universitas Pahlawan Tuanku Tambusai,” *Journal on Pustaka Cendekia Informatika*, vol. 2, no. 1, pp. 22–27, 2024, doi: 10.70292/pctif.v2i1.61.
- [7] Muhammad Tarmizi and Y. Yahfizham, “Perspektif Mahasiswa Terhadap Penggunaan Kecerdasan Buatan ChatGPT dalam Penyusunan Tugas Akhir,” *Indiktika : Jurnal Inovasi Pendidikan Matematika*, vol. 6, no. 2, pp. 151–161, Jun. 2024, doi: 10.31851/indiktika.v6i2.15425.
- [8] P. Kusumaningtyas, A. P. Arrumi, and dan S. Keren Tiurma Eunike, “Efektivitas Pemanfaatan ChatGPT dalam Tugas Esai Mahasiswa Ilmu Komunikasi Universitas Negeri Surabaya,” *Prosiding Seminar Nasional Ilmu-Ilmu Sosial (SNIIS)*, vol. 2, pp. 158–165, 2023, Accessed: Jun. 19, 2026. [Online]. Available: <https://proceeding.unesa.ac.id/index.php/sniis/article/view/794/266>
- [9] I. M. Castillo-Martínez, D. Flores-Bueno, S. M. Gómez-Puente, and V. O. Vite-León, “AI in higher education: a systematic literature review,” 2024, *Frontiers in Education*. doi: 10.3389/educ.2024.1391485.
- [10] S. Bećirović, E. Polz, and I. Tinkel, “Exploring students’ AI literacy and its effects on their AI output quality, self-efficacy, and academic performance,” *Smart Learning Environments*, vol. 12, no. 1, Dec. 2025, doi: 10.1186/s40561-025-00384-3.
- [11] A. B. Bakker and K. Mostert, “Study Demands–Resources Theory: Understanding Student Well-Being in Higher Education,” Sep. 01, 2024, Springer. doi: 10.1007/s10648-024-09940-8.
- [12] L. Eastgate, P. A. Creed, M. Hood, and A. Bialocerkowski, “It Takes Work: How University Students Manage Role Boundaries when the Future is Calling,” *Res. High. Educ.*, vol. 64, no. 7, pp. 1071–1088, Nov. 2023, doi: 10.1007/s11162-023-09736-9.
- [13] E. E. Pérez-Guerrero et al., “Methodological and Statistical Considerations for Cross-Sectional, Case–Control, and Cohort Studies,” Jul. 01, 2024, Multidisciplinary Digital Publishing Institute (MDPI). doi: 10.3390/jcm13144005.
- [14] D. Rutkowski, L. Rutkowski, G. Thompson, and Y. Canbolat, “The limits of inference: reassessing causality in international assessments,” *Large. Scale. Assess. Educ.*, vol. 12, no. 1, Dec. 2024, doi: 10.1186/s40536-024-00197-9.
- [15] L. Widiarti, Sahiruddin, and M. A. Kasri, “Efektivitas ChatGPT sebagai Asisten Virtual untuk Mendukung Produktivitas Akademik Mahasiswa Pendidikan Teknologi Informasi,” *Decode: Jurnal Pendidikan Teknologi Informasi*, vol. 5, no. 3, pp. 1218–1224, Nov. 2025, doi: 10.51454/decode.v5i3.1464.
- [16] C. Andrade, “The Inconvenient Truth About Convenience and Purposive Samples,” *Indian J. Psychol. Med.*, vol. 43, no. 1, pp. 86–88, Jan. 2021, doi: 10.1177/0253717620977000.
- [17] A. A. Syanzani, N. Azrina, and V. Fitriani, “Analisis Keefektifan ChatGPT dalam Membantu Proses Belajar pada Mahasiswa STMIK Antar Bangsa,” 2024, doi: 10.51998/jti.v10i2.580.
- [18] M. Kankaraš and S. Capecchi, “Neither agree nor disagree: use and misuse of the neutral response category in Likert-type scales,” *Metron*, vol. 83, no. 1, pp. 111–140, Apr. 2025, doi: 10.1007/s40300-024-00276-5.
- [19] C. P. Klapproth, F. Fischer, and M. Rose, “Scale agreement, ceiling and floor effects, construct validity, and relative efficiency of the PROPr and EQ-5D-3L in low back pain patients,” *Health Qual. Life Outcomes*, vol. 21, no. 1, Dec. 2023, doi: 10.1186/s12955-023-02188-w.
- [20] D. Goretzko and M. Bühner, “Robustness of factor solutions in exploratory factor analysis,” *Behaviormetrika*, vol. 49, no. 1, pp. 131–148, Jan. 2022, doi: 10.1007/s41237-021-00152-w.
- [21] Z. Zeng, G. Xie, Y. He, and S. Zhou, “Exploring psychological pathways between workplace violence and burnout among nurses in Chinese Tertiary Hospitals,” *Sci. Rep.*, vol. 15, no. 1, Dec. 2025, doi: 10.1038/s41598-025-19671-7.
- [22] B. Struminskaya and T. Gummer, “Risk of Nonresponse Bias and the Length of the Field Period in a Mixed-Mode General Population Panel,” *J. Surv. Stat. Methodol.*, vol. 10, no. 1, pp. 161–182, Feb. 2022, doi: 10.1093/jssam/smab011.
- [23] G. W. Cheung, H. D. Cooper-Thomas, R. S. Lau, and L. C. Wang, “Reporting reliability, convergent and discriminant validity with structural equation modeling: A review and best-practice recommendations,” *Asia Pacific Journal of Management*, vol. 41, no. 2, pp. 745–783, Jun. 2024, doi: 10.1007/s10490-023-09871-y.
- [24] F. Fiel Peres, “Effect sizes for nonparametric tests,” Feb. 15, 2026. doi: 10.11613/BM.2026.010101.
- [25] K. Nurumov, D. Hernández-Torrano, A. Ait Si Mhamed, and U. Ospanova, “Measuring Social Desirability in Collectivist Countries: A Psychometric Study in a Representative Sample From Kazakhstan,” *Front. Psychol.*, vol. 13, Apr. 2022, doi: 10.3389/fpsyg.2022.822931.