

EXPLORING GENERATIVE AI UTILIZATION USER PERSONA CREATION: QUALITATIVE STUDY UX PRACTITIONERS

EKSPLORASI PEMANFAATAN *GENERATIVE AI* DALAM PEMBUATAN *USER PERSONA*: STUDI KUALITATIF PADA PRAKTIKSI UX

Raihan Alghani Leksono¹, Firman Ardiansyah², and Shelve Nidya Neyman³

^{1,2,3}Computer Science Study Program, School of Data Science, Mathematics, and Informatics, IPB University, Indonesia

[†]corresponding author: [raihanalghani @apps.ipb.ac.id](mailto:raihanalghani@apps.ipb.ac.id)

Copyright © 2026 M Raihan Alghani Leksono, Firman Ardiansyah, and Shelve Nidya Neyman. This is an open-access article distributed under the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

Abstract

User persona creation represents a crucial stage in User-Centered Design (UCD) for understanding user needs; however, traditional approaches are often time-consuming and may hinder rapid iteration. The emergence of Generative AI (GenAI) powered by Large Language Models (LLMs) offers the potential to accelerate this process, although concerns remain regarding reliability, bias, and hallucination. A research gap persists due to the limited number of studies examining how UX practitioners in industry contexts utilize GenAI and critically perceive the quality of its generated artifacts. This study aims to investigate the role of GenAI in user persona creation, map the taxonomy of prompting strategies, identify methodological risks, and formulate an adaptive governance-oriented conceptual framework. An exploratory qualitative approach was employed through semi-structured in-depth interviews with eight User Experience (UX) practitioners from diverse industry sectors. Data were analyzed using Reflexive Thematic Analysis to inductively construct patterns of meaning based on meaning units. The study successfully formulated a Conceptual Framework for AI-Assisted Persona Creation integrating operational workflow and governance into four core components: input, process, output risk, and governance. Findings indicate that practitioners utilize GenAI as a cognitive assistant through contextual prompting strategies, although AI-generated outputs exhibit structural limitations, including subtle hallucinations and demographic bias. The study concludes that the validity

of AI-assisted user personas is not determined by full automation; rather, it fundamentally depends on human-in-the-loop mechanisms and the principle of insight traceability to mitigate the risk of analytical skill degradation among practitioners.

Keywords: Conceptual Framework, Generative AI, Reflexive Thematic Analysis, User Experience Research, User Persona.

Abstrak

Pembuatan *user persona* merupakan tahapan krusial dalam *User-Centered Design* (UCD) untuk memahami kebutuhan pengguna, namun pendekatan tradisional sering kali memakan waktu dan menghambat iterasi cepat. Kehadiran *Generative AI* (GenAI) berbasis *Large Language Models* (LLM) menawarkan potensi akselerasi proses tersebut, meski memicu kekhawatiran terkait reliabilitas, bias, dan halusinasi. Kesenjangan riset muncul karena belum banyak studi yang mengeksplorasi dinamika pemanfaatan GenAI oleh praktisi UX di industri serta bagaimana kualitas artefaknya dipersepsikan secara kritis. Penelitian ini bertujuan menginvestigasi peran GenAI dalam pembuatan *user persona*, memetakan taksonomi strategi *prompting*, mengidentifikasi risiko metodologis, serta merumuskan kerangka konseptual tata kelola yang adaptif. Pendekatan kualitatif eksploratif diterapkan melalui wawancara mendalam semi-terstruktur terhadap delapan praktisi *User Experience* (UX) lintas sektor industri. Data dianalisis menggunakan metode *Reflexive Thematic Analysis* untuk mengonstruksi pola makna secara induktif berbasis *meaning unit*. Hasil penelitian berhasil merumuskan *Conceptual Framework for AI-Assisted Persona Creation* yang mengintegrasikan alur operasional dan tata kelola ke dalam empat komponen utama: *input*, *process*, *output risk*, dan *governance*. Temuan menunjukkan bahwa praktisi memanfaatkan GenAI sebagai asisten kognitif melalui strategi *prompting* kontekstual, walaupun luaran AI terindikasi mengalami keterbatasan struktural berupa halusinasi halus dan bias demografi. Simpulan penelitian menegaskan bahwa validitas *user persona* berbasis AI tidak ditentukan oleh otomasi penuh, melainkan mutlak tergantung pada mekanisme *human-in-the-loop* dan prinsip keterlacakan bukti (*insight traceability*) guna menekan risiko degradasi kemampuan analitis praktisi.

Kata Kunci: Analisis Tematik *Reflexive*, *Generative AI*, Kerangka Konseptual, Riset Pengalaman Pengguna, *User Persona*.

1. Pendahuluan

Proses desain *User Experience* (UX) modern sangat bergantung pada pendekatan *User Centered Design* (UCD), di mana pembuatan *user persona* menjadi artefak fundamental untuk mewakili segmen pengguna berdasarkan data empiris (Siregar *et al.*, 2025). Persona yang efektif harus

dibangun melalui riset kualitatif yang ketat agar tim desain dapat berempati dan mengambil keputusan secara akurat (Cooper et al., 2014; Lauer et al., 2025). Namun, dalam metode tradisional, proses ekstraksi data kualitatif ini memakan waktu dan sumber daya yang besar, sehingga kerap menghambat kelancaran iterasi desain yang cepat di industri. Kemajuan *artificial intelligence*, khususnya *Generative AI* (GenAI) berbasis *Large Language Models* (LLM), menawarkan solusi atas kendala operasional tersebut. LLM mampu memproses data teks masif, mengenali konteks, dan menghasilkan luaran bahasa alami yang komprehensif (Limantara, 2025). Melalui teknik *prompting* yang adaptif, GenAI dapat diinstruksikan untuk menjalankan tugas penalaran kompleks secara instan (Reynolds dan McDonell, 2021). Meskipun menawarkan efisiensi tinggi, otomasi proses kognitif ini turut memicu tantangan metodologis baru, mulai dari kemunculan halusinasi informasi, bias demografi, hingga potensi ketergantungan praktisi terhadap sistem otomatis (Yang et al., 2023).

Studi terdahulu telah berupaya mengeksplorasi isu pemanfaatan AI ini, namun belum menyentuh realitas operasional di level industri profesional. Jiang et al., (2023) berfokus pada pengujian kapabilitas teknis model dalam mengekspresikan sifat kepribadian pada simulasi persona, sementara Strahler (2025) mengeksplorasi GenAI sebatas sebagai alat bantu simulasi profil pengguna di lingkungan pendidikan. Literatur yang ada masih berpusat pada ranah eksperimen teknis model dan belum membedah secara empiris bagaimana praktisi menavigasi risiko bias dan validitas data di lapangan. Terdapat kesenjangan literatur yang signifikan terkait dinamika interaksi praktisi UX dalam mengontrol risiko metodologis saat mengintegrasikan GenAI ke dalam alur kerja riset.

Mengisi kesenjangan literatur tersebut, penelitian ini dilaksanakan dengan tujuan utama untuk memetakan praktik dan variasi strategi *prompting* yang digunakan oleh praktisi UX profesional ke dalam taksonomi pola pemanfaatan GenAI. Selain itu, penelitian ini juga mengidentifikasi bentuk risiko metodologis (seperti bias, halusinasi, overgeneralisasi, dan ilusi kecerdasan) serta mengeksplorasi persepsi praktisi terhadap kualitas artefak persona yang dihasilkan. Berdasarkan temuan tersebut, penelitian ini merumuskan susunan *Conceptual Framework for AI-Assisted Persona Creation* yang memuat tata kelola mitigasi risiko dan mekanisme *human in the loop* dalam alur kerja praktisi.

Guna menjaga ketajaman analisis, ruang lingkup penelitian ini dibatasi pada pemanfaatan GenAI berbasis teks dalam mengekstraksi elemen inti *user persona*, yaitu tujuan (*goals*), tantangan (*pains*), kebutuhan (*needs*), latar belakang pengguna, serta perilaku penggunaan (*behaviors*). Pengumpulan data dilakukan melalui wawancara mendalam bersama praktisi UX profesional yang memiliki pengalaman nyata, kemudian dianalisis menggunakan metode *reflexive thematic analysis*. Penelitian ini secara sengaja tidak mencakup *high-fidelity prototyping*, *usability testing*,

maupun implementasi sistem agar tetap fokus pada konseptualisasi persona. Melalui batasan ini, penelitian diharapkan mampu memberikan pemahaman empiris yang mendalam serta menawarkan kerangka konseptual tata kelola yang dapat diimplementasikan langsung untuk mengoptimalkan efisiensi kerja sekaligus memitigasi risiko metodologis di lapangan.

2. Metodologi

2.1 Pendekatan Penelitian

Penelitian ini menerapkan pendekatan kualitatif eksploratif yang berlandaskan pada paradigma konstruktivis-interpretivis. Desain eksploratif digunakan karena fenomena kolaborasi praktisi UX dan GenAI belum memiliki kerangka teoritis yang mapan. Fokus penelitian telah bergeser dari sekadar merumuskan panduan operasional menjadi upaya mengonstruksi pola makna, taksonomi *prompting*, dan sintesis sebuah *Conceptual Framework* tata kelola risiko berdasarkan pengalaman empiris di lapangan.

2.2 Partisipan dan Posisionalitas Peneliti

Penelitian melibatkan delapan praktisi UX profesional yang direkrut melalui teknik *purposive sampling* dengan kriteria inklusi memiliki pengalaman minimal satu tahun dan secara aktif menggunakan *Large Language Models* (LLM) untuk sintesis artefak persona. Pengumpulan data dihentikan pada partisipan kedelapan setelah prinsip saturasi data terpenuhi (Guest *et al.*, 2006). Sebagai instrumen utama (*human instrument*), peneliti menyadari posisionalitasnya sebagai mahasiswa ilmu komputer yang memiliki ketertarikan tinggi terhadap AI. Guna mencegah bias pro-teknologi dalam mengarahkan narasi, peneliti secara sadar menerapkan sikap refleksif melalui pencarian kasus negatif (*negative cases*) serta penggunaan pertanyaan terbuka yang memberi ruang bagi partisipan untuk mengkritik teknologi secara leluasa.

2.3 Pengumpulan Data

Data primer dikumpulkan melalui wawancara mendalam semi-terstruktur secara daring menggunakan platform *zoom meeting* dengan durasi 40 hingga 60 menit per sesi. Panduan wawancara mencakup empat dimensi pokok yang meliputi praktik penggunaan, strategi *prompting*, evaluasi risiko artefak, dan mekanisme mitigasi. Seluruh sesi direkam atas persetujuan partisipan (*informed consent*) dan ditranskripsikan secara *verbatim*. Transkrip kemudian disandingkan dengan catatan lapangan (*field notes*) untuk menjaga kelengkapan konteks non-verbal.

2.4 Teknik Analisis dan Kredibilitas Data

Seluruh transkrip dianalisis menggunakan metode *Reflexive Thematic Analysis* yang berakar pada kerangka kerja Braun dan Clarke (2020). Proses analisis tidak menggunakan perangkat lunak kualitatif spesifik, melainkan dikelola secara manual menggunakan pengolah kata dan *spreadsheet* untuk menjaga kedekatan peneliti dengan data.



Gambar 1: *Thematic Analysis*.

Tahapan koding bersifat induktif dan organik (*open coding*) berbasis *meaning unit*. Ekstraksi data ini menghasilkan puluhan kode awal yang kemudian direduksi. Resolusi terhadap konflik interpretasi antarkode diselesaikan melalui penerapan *constant comparative method* (Charmaz, 2014), yakni dengan membandingkan temuan silang antarpartisipan secara konstan.

Guna menjamin keabsahan data (*trustworthiness*), serangkaian protokol ketat diterapkan. *Audit trail* dilakukan melalui pembuatan catatan analitis (*analytic memos*) untuk melacak evolusi konseptual tema. Proses pengecekan anggota pasca-analisis tidak dilakukan secara formal akibat keterbatasan waktu praktisi, namun digantikan oleh *real-time member checking* (verifikasi makna dan parafrase) yang dilakukan secara interaktif sepanjang sesi wawancara berlangsung. Pada tahap akhir, struktur tema potensial yang telah dirumuskan secara mandiri dievaluasi secara kritis melalui *peer debriefing* bersama pembimbing selaku *expert reviewer*. Masukan dari proses *review* tersebut kemudian digunakan untuk memfinalisasi tema dan menyintesisnya ke dalam bentuk kerangka konseptual tata kelola AI.

3. Hasil dan Pembahasan

3.1 Gambaran Umum Pelaksanaan Penelitian

Penelitian ini diawali dengan merumuskan instrumen wawancara yang mencakup 20 pertanyaan utama dalam empat dimensi pokok: praktik penggunaan, strategi *prompting*, evaluasi risiko artefak, dan mekanisme mitigasi. Proses rekrutmen partisipan dilakukan menggunakan teknik *purposive* dan *snowball sampling* hingga mencapai delapan orang praktisi UX profesional (Tabel 1). Pengumpulan data dihentikan pada partisipan kedelapan (P8) karena peneliti menilai telah terjadi saturasi data. Indikator kejenuhan ini dibuktikan melalui redundansi data, di mana pola temuan mengenai risiko overgeneralisasi perilaku dan prosedur anonimisasi yang

diutarakan secara mendalam oleh P7 kembali berulang secara identik pada P8 tanpa memunculkan variasi konseptual baru (Guest *et al.*, 2006).

Tabel 1: Profil 8 partisipan dalam penelitian.

ID partisipan	Peran	Sektor industri
P1	<i>Product Designer</i>	Teknologi dan informasi
P2	<i>Product Researcher</i>	Transportasi
P3	<i>UI/UX Designer</i>	<i>Outsourcing</i>
P4	<i>Product Designer</i>	Teknologi dan informasi
P5	<i>Product Designer</i>	Teknologi, informasi, dan komunikasi
P6	<i>Product Owner</i>	Perbankan
P7	<i>UX Designer</i>	<i>Fintech</i>
P8	<i>Product Designer</i>	<i>Fintech</i>

Seluruh sesi wawancara dilaksanakan secara daring dengan durasi 40 hingga 60 menit dan ditranskripsikan secara *verbatim*. Sebagai bentuk reflektivitas, peneliti secara sadar memosisikan diri sebagai pengamat netral guna mencegah bias konfirmasi akibat latar belakang keilmuan teknis AI yang dimiliki peneliti. Guna menjamin keabsahan data (*trustworthiness*), peneliti menerapkan pembentukan *audit trail* yang sistematis, di mana setiap tema yang dirumuskan memiliki rekam jejak rujukan (*traceability*) yang jelas pada kutipan langsung partisipan (Nowell *et al.*, 2017).

3.2 Pemetaan Tematis Praktik Generative AI

Data kualitatif dianalisis menggunakan *Reflexive Thematic Analysis* (Braun dan Clarke, 2020) berbasis *meaning unit* (unit makna). Draf pengodean awal yang semula bersifat operasional direduksi dan distandardisasi menjadi konsep akademis formal. Hasil ekstraksi data mengkristal menjadi enam tema final (Tabel 2) yang merepresentasikan dinamika interaksi praktisi dengan AI.

Tabel 2: Pemetaan tema hasil analisis data kualitatif.

No	Tema final (Fase 5)	Draft kode awal (Fase 2)
1	Peran AI sebagai Asisten Kognitif dalam Alur Riset	<i>Pre/post-research utilization, data synthesis automation, time-constrained operational pressure, workload scalability.</i>

2	Strategi <i>Prompting</i> yang Kontekstual dan Adaptif	<i>Iterative prompting, context-aware prompting, empirical grounding necessity, role-playing directives.</i>
3	Keterbatasan Struktural pada Artefak Persona Berbasis AI	<i>Expansive/reductive hallucination, plausible hallucination, WEIRD bias, behavioral overgeneralization.</i>
4	Ilusi Kecerdasan dan Risiko <i>Over-Trust</i> terhadap AI	<i>Cognitive offloading, deskilling risk, illusion of competence, automation bias, rejection risk.</i>
5	Batasan Peran AI dalam Aspek Empati dan Pengambilan Keputusan	<i>Empathetic limitation, contextual accountability, authenticity of evidence.</i>
6	Mekanisme Mitigasi Risiko dalam Penggunaan AI	<i>Data anonymization, selective disclosure, insight traceability, human-in-the-loop intervention.</i>

*WEIRD: *Western, Educated, Industrialized, Rich, and Democratic*.
leluasa.

3.3 Temuan Utama Pemanfaatan Generative AI

3.3.1. Manfaat Operasional dan Interaksi Kognitif

Dimensi ini mencakup Tema 1 dan Tema 2 yang menyoroti motivasi serta taktik teknis praktisi. Temuan menunjukkan bahwa pemicu utama adopsi GenAI adalah tekanan ritme kerja (*pace industri*) dan *workload scalability*. P8 secara mendalam menjelaskan ketimpangan ini: "*Skala proyeknya makin lama makin gede, dan tim risetnya nggak nambah... ini nggak sustainable.*" GenAI diposisikan sebagai asisten kognitif untuk fase pra-pengumpulan data (membangun kerangka awal) hingga pasca-riset (sintesis dan *affinity mapping*).

Dalam berinteraksi, praktisi menerapkan taksonomi *prompting* yang adaptif, mulai dari pemberian peran (*role-playing*), instruksi masif di awal (*front-loaded mega prompt*), hingga teknik memancing konteks bertahap (*reverse-prompting*). Menariknya, penelitian ini juga menemukan adanya kasus negatif (*negative case*) atau anomali strategi. Mayoritas praktisi sangat menghindari *zero-shot prompt* demi menjaga landasan empiris. Namun, P1 justru sengaja menggunakan *zero-shot prompt* di awal interaksi khusus untuk menghindari bias konfirmasi peneliti (*framing bias*). P1 menyatakan: "*Takutnya pandanganku ngebuat jawaban dia jadi seolah-olah sama kayak kesukaanku... aku butuh user persona versi dia dulu.*" Praktik ini membuktikan bahwa interaksi dengan AI tidak bersifat statis, melainkan dieksplorasi secara kritis oleh pengguna.

3.3.2. Keterbatasan Struktural dan Risiko Sosioteknis

Dimensi ini mengintegrasikan Tema 3 dan Tema 4 yang merupakan temuan paling kritis secara akademis. Keterbatasan struktural AI termanifestasi dalam bentuk halusinasi data dan *WEIRD bias* (bias kelas sosial). P2 mengilustrasikan bias tersebut: *"Dia lumayan mengambil atau menggeneralisasikan konteks persona ke ekstrem kanan, yang kaya, kaya banget, yang Jaksel, Jaksel banget... nggak merepresentasikan segmentasi yang saya punya."* AI juga kerap melakukan homogenisasi artefak (P7, P8) seolah seluruh pengguna memiliki perilaku yang seragam.

Dampak paling berbahaya dari kefasihan bahasa AI adalah munculnya fenomena ilusi kecerdasan (*illusion of competence*). Kefasihan ini memicu pelepasan beban kognitif (*cognitive offloading*) yang tidak terkontrol (Risiko dan Gilbert 2016), di mana praktisi rentan mengalami *automation bias* dan mempercayai output secara buta (*over-trust*). Hal ini secara langsung berujung pada degradasi kemampuan analitis (*deskilling*). Kekhawatiran ini diungkapkan secara personal oleh P8: *"Apakah skill gue di area ini lama-lama jadi tumpul?... Kemampuan gue buat sintesis manual sekarang mungkin tidak setajam dua tahun lalu."* Temuan ini membuktikan bahwa GenAI tidak hanya berdampak pada hasil kerja, tetapi juga mendisrupsi lanskap kognitif dan ketahanan kompetensi praktisi secara jangka panjang.

3.3.3. Tata Kelola Manusia dan Batasan Operasional

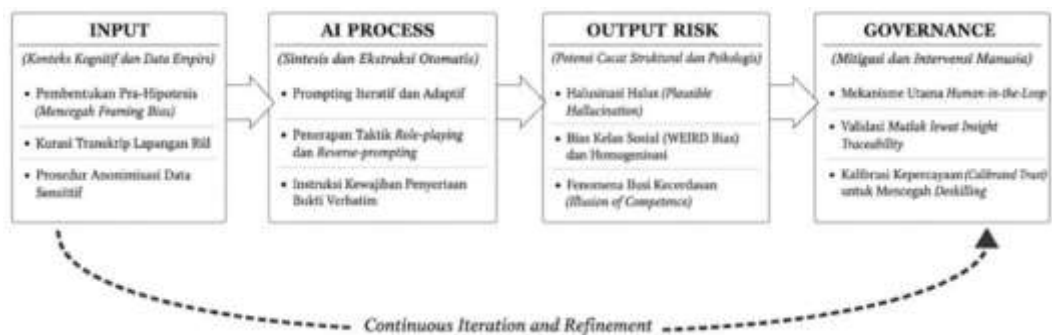
Dimensi ini mewakili tema 5 dan 6 dalam membahas strategi mitigasi yang diterapkan praktisi. Meskipun mendelegasikan tugas sintesis kepada AI, praktisi sepakat bahwa aspek empati dan pengambilan keputusan kritis merupakan batas fundamental manusia. P4 menegaskan tanggung jawab profesionalnya: *"Kita diamanahin untuk bikin produk yang bener, bukan diamanahin untuk nge-prompting AI."*

Sebagai langkah mitigasi, praktisi menerapkan kontrol ketat berupa anonimisasi data sensitif (P1, P7, P8) dan *selective data disclosure* (P6). Evaluasi mutlak bertumpu pada mekanisme *human-in-the-loop*. Untuk menjamin validitas, P7 dan P8 menerapkan prinsip *insight traceability*, di mana setiap poin pada artefak akhir harus memiliki jangkar bukti yang kuat. P8 menyatakan: *"Setiap klaim yang masuk ke persona final harus bisa gue trace balik ke minimal satu transkrip... kalau nggak bisa, mending buang."*

3.4 Kerangka Konseptual Tata Kelola AI (*Conceptual Framework*)

Berdasarkan sintesis dari ketiga dimensi tematik di atas, keluaran penelitian ini dikonstruksi menjadi *Conceptual Framework for AI-Assisted*

Persona Creation (Gambar 2). Transformasi dari temuan lapangan menjadi kerangka konseptual ini didasarkan pada pemetaan alur kausalitas sosioteknis, yaitu hubungan sebab-akibat antara manusia dan teknologi (Mumford 2006). Secara spesifik, temuan pada dimensi Manfaat Operasional diterjemahkan menjadi komponen *Input* dan *AI Process*, di mana persiapan kognitif praktisi dan strategi interaksi menjadi motor penggerak. Selanjutnya, dimensi Keterbatasan Struktural dan Risiko direpresentasikan sebagai *Output Risk*, yakni titik kritis di mana anomali seperti halusinasi dan bias berpotensi muncul. Tata Kelola Manusia, sebagai dimensi ketiga, diadopsi secara langsung menjadi komponen *Governance* sebagai lapisan pertahanan akhir yang menetapkan protokol intervensi manual terhadap risiko luaran tersebut. Kerangka ini memberikan gambaran mengenai alur dan keterhubungan setiap tahapan dalam proses penelitian, sehingga proses yang berlangsung dari persiapan data hingga validasi dapat dipahami secara lebih sistematis.



Gambar 2: Alur *Conceptual Framework*.

Berdasarkan alur pada Gambar 2, dapat dijelaskan sebagai berikut:

- Input** (Konteks Kognitif dan Data Empiris): Mewajibkan penyelarasan pemahaman konteks dengan membentuk pra-hipotesis guna mencegah *framing bias* serta penyediaan data empiris yang telah dianonimisasi untuk mencegah kebocoran data sensitif.
- AI Process** (Sintesis dan Ekstraksi): Proses *prompting* iteratif dan adaptif seperti *role-playing*, instruksi dengan mendefinisikan *constraint*, dan kewajiban penyertaan *verbatim* yang berfungsi untuk mengonversi data mentah menjadi draf wawasan struktural.
- Output Risk** (Potensi Cacat Struktural): Tahap identifikasi titik buta luaran AI yang berpotensi mengandung halusinasi halus (*plausible hallucination*), *WEIRD bias*, serta overgeneralisasi perilaku yang memicu ilusi kecerdasan.
- Governance** (Mitigasi dan Intervensi Manusia): Penerapan *human-in-the-loop* sebagai verifikasi tahap akhir. Komponen ini mewajibkan validasi menggunakan prinsip *insight traceability*, guna memastikan setiap elemen persona berbasis fakta sekaligus mencegah

terjadinya *deskilling* pada praktisi UX akibat *over-trust*.

4. Simpulan dan Saran

Penelitian ini menyimpulkan bahwa GenAI tidak diposisikan sebagai pengganti *UX researcher*, melainkan berperan sebagai asisten kognitif untuk mengakselerasi sintesis data di tengah tekanan operasional industri. Praktisi berinteraksi dengan AI melalui strategi *prompting* yang adaptif, terstruktur, dan kontekstual, di mana praktik *zero-data generation* sangat dihindari guna menjaga landasan empiris riset. Namun, pemanfaatan teknologi ini memicu risiko struktural berupa halusinasi halus dan *WEIRD bias* (bias kelas sosial), serta ancaman psikologis seperti pelepasan beban kognitif (*cognitive offloading*) dan *automation bias* yang berpotensi mendegradasi ketajaman insting analitis praktisi (*deskilling*). Sebagai kontribusi utama, penelitian ini berhasil merumuskan *Conceptual Framework for AI-Assisted Persona Creation* yang memetakan dinamika kausalitas antara komponen *Input*, *AI Process*, *Output Risk*, dan *Governance*. Kerangka konseptual ini menegaskan bahwa validitas artefak persona tidak ditentukan oleh otomasi penuh, melainkan mutlak membutuhkan intervensi *human-in-the-loop* dan penerapan prinsip keterlacakan wawasan (*insight traceability*). Elemen empati, interpretasi emosi pengguna, dan pengambilan keputusan kritis tetap menjadi domain eksklusif manusia yang tidak dapat didelegasikan kepada mesin.

Berdasarkan hasil penelitian yang telah dilakukan, terdapat beberapa saran penting yang dirumuskan untuk agenda penelitian selanjutnya. Pertama, penelitian selanjutnya dapat menggunakan pendekatan *mixed methods* atau metode kuantitatif untuk membandingkan secara langsung kualitas *user persona* yang dibuat secara manual dengan yang dibantu oleh *Generative AI*. Kedua, studi lanjutan disarankan untuk melibatkan lebih banyak partisipan dari berbagai sektor industri dan tingkat pengalaman yang lebih bervariasi agar cakupan hasil penelitian menjadi lebih luas dan beragam. Ketiga, ruang lingkup pemanfaatan *Generative AI* dapat dieksplorasi lebih jauh pada tahapan proses kerja UX lainnya, seperti *usability testing*, *customer journey mapping*, atau penyusunan skenario penggunaan. Terakhir, penelitian berikutnya perlu membahas lebih dalam mengenai aspek etika, kebijakan privasi data, serta batasan tanggung jawab profesional dalam penggunaan AI untuk riset pengguna guna memastikan integritas praktik perancangan tetap terjaga.

Daftar Pustaka

Braun V, Clarke V. 2006. Using thematic analysis in psychology. *Qualitative Research in Psychology*. [diakses 2025 Des 9]; 3(2):77–101. <https://doi.org/10.1191/1478088706qp0630a>.

Braun V, Clarke V. 2020. One size fits all? What counts as quality practice in

- (reflexive) thematic analysis?. [diakses 2026 Mei 18]; *Qualitative Research in Psychology*. 18(3):328-352. doi:10.1080/14780887.2020.1769238.
- Charmaz K. 2014. *Constructing Grounded Theory* (2nd ed.). London (UK): SAGE Publications.
- Cooper A, Reimann R, Cronin D, Noessel C. 2014. *About Face: The Essentials of Interaction Design*. Ed ke-4. Indianapolis (IN): J Wiley.
- Dellermann D, Ebel P, Söllner M, Leimeister JM. 2019. Hybrid intelligence. *Business & Information Systems Engineering*. [diakses 2026 Mei 12]; 61(5):637-643. <https://doi.org/10.1007/s12599-019-00595-2>.
- Guest G, Bunce A, Johnson L. 2006. How many interviews are enough? An experiment with data saturation and variability. *Field Methods*. [diakses 2025 Des 11]; 18(1):59–82. <https://doi.org/10.1177/1525822X05279903>.
- Hennink M, Kaiser BN. 2022. Sample sizes for saturation in qualitative research: a systematic review of empirical tests. *Soc Sci Med*. [diakses 2026 Mei 13]; 292:114523. doi:10.1016/j.socscimed.2021.114523.
- Hoff KA, Bashir M. 2015. Trust in automation: integrating empirical evidence on factors that influence trust. *Human Factors*. [diakses 2026 Mei 11]; 57(3):407-434. <https://doi.org/10.1177/0018720814547570>.
- Jiang H, Zhang X, Cao X, Kabbara J, Roy D. 2023. PersonaLLM: investigating the ability of large language models to express personality traits. *arXiv* [diakses 2025 Des 11]; 2305.02547. <https://arxiv.org/abs/2305.02547>.
- Lauer C, Storey D, Soley R. 2025. Leveraging AI toward the development of vector personas for ux research. *Journal of User Experience*. 20(4): 149-165. <https://uxpajournal.org/leveraging-ai-toward-the-development-of-vector-personas-for-ux-research/>
- Limantara L. 2025. AI sebagai alat transformasi pendidikan: mewujudkan pembelajaran kritis dan inovatif di uk petra. Di dalam: Tjiek LT, Sugiarto I, Iskandar EA, editor. *Towards an AI-Native Campus: UK Petra x Kecerdasan Artifisial*. Surabaya: Bagian Penerbit Lembaga Penelitian dan Pengabdian kepada Masyarakat, Universitas Kristen Petra. hlm 53-67.
- Mumford E. 2006. The story of socio-technical design: reflections on its successes, failures and potential. *Information Systems Journal*. 16(4):317-342. <https://doi.org/10.1111/j.1365-2575.2006.00221.x>.
- Nielsen J. 1994. *Usability Engineering*. San Diego (CA): Morgan Kaufmann.
- Nowell LS, Norris JM, White DE, Moules NJ. 2017. Thematic analysis: striving to meet the trustworthiness criteria. *Int J Qual Methods*. 16(1):1-13. doi:10.1177/1609406917733847.

- Parasuraman R, Manzey DH. 2010. Complacency and bias in human use of automation: an attentional integration. *Human Factors*. 52(3):381-410. <https://doi.org/10.1177/0018720810376055>.
- Reynolds L, McDonell K. 2021. Prompt programming for large language models: beyond the few-shot paradigm. Di dalam: Yoshifumi K, editor. CHI '21 Extended Abstracts of the 2021 CHI Conference on Human Factors in Computing Systems; 2021 Mei 8-13; Yokohama, Jepang. New York (NY): ACM. hlm 1-7. <https://doi.org/10.1145/3411763.3451760>.
- Risko EF, Gilbert SJ. 2016. Cognitive offloading. *Trends in Cognitive Sciences*. 20(9):676-688. <https://doi.org/10.1016/j.tics.2016.07.002>.
- Siregar A, Saputra E, Arsa D, Ramadhani NP. 2025. Perancangan ui/ux sistem prakerin web dengan metode ucd. *Jurnal Informatika Polinema*. 11(4): 531–540. <https://doi.org/10.33795/jip.v11i4.7625>.
- Strahler D. 2025. Generative AI (GenAI) for co-creating user personas. *Proceedings of the New York State Communication Association*. 2023(1):3. <https://docs.rwu.edu/nyscaproceedings/vol2024/iss1/3/>.
- Yang J, Jin H, Tang R, Han X, Feng Q, Jiang H, Yin B, Hu X. 2023. Harnessing the power of LLMs in practice: a survey on ChatGPT and beyond. *arXiv* [diakses 2025 Nov 21]; 2304.13712. <https://arxiv.org/abs/2304.13712>.
- Yildirim I, Paul LA. 2024. From task structures to world models: what do LLMs know?. *Trends in Cognitive Sciences*. [diakses 2026 Mei 12]; 28(5):404-415. <https://doi.org/10.1016/j.tics.2024.02.008>.