



- Hak Cipta Dilindungi Undang-undang
1. Dilarang mengutip sebagian atau seluruh karya tulis ini tanpa mencantumkan dan menyebutkan sumber :
 - a. Pengutipan hanya untuk kepentingan pendidikan, penelitian, penulisan karya ilmiah, penyusunan laporan, penulisan kritik atau tinjauan suatu masalah
 - b. Pengutipan tidak merugikan kepentingan yang wajar IPB University.
 2. Dilarang mengumumkan dan memperbanyak sebagian atau seluruh karya tulis ini dalam bentuk apapun tanpa izin IPB University.

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

M RAIHAN ALGHANI LEKSONO



PROGRAM SARJANA ILMU KOMPUTER
SEKOLAH SAINS DATA, MATEMATIKA, DAN INFORMATIKA
INSTITUT PERTANIAN BOGOR
BOGOR
2026



@Hak cipta milik IPB University

Hak Cipta Dilindungi Undang-undang

1. Dilarang mengutip sebagian atau seluruh karya tulis ini tanpa mencantumkan dan menyebutkan sumber :
 - a. Pengutipan hanya untuk kepentingan pendidikan, penelitian, penulisan karya ilmiah, penyusunan laporan, penulisan kritik atau tinjauan suatu masalah
 - b. Pengutipan tidak merugikan kepentingan yang wajar IPB University.
2. Dilarang mengumumkan dan memperbanyak sebagian atau seluruh karya tulis ini dalam bentuk apapun tanpa izin IPB University.

PERNYATAAN MENGENAI SKRIPSI DAN SUMBER INFORMASI SERTA PELIMPAHAN HAK CIPTA

Dengan ini saya menyatakan bahwa skripsi dengan judul “Eksplorasi Pemanfaatan *Generative AI* dalam Pembuatan *User Persona*: Studi Kualitatif pada Praktisi UX” adalah karya saya dengan arahan dari dosen pembimbing dan belum diajukan dalam bentuk apa pun kepada perguruan tinggi mana pun. Sumber informasi yang berasal atau dikutip dari karya yang diterbitkan maupun tidak diterbitkan dari penulis lain telah disebutkan dalam teks dan dicantumkan dalam Daftar Pustaka di bagian akhir skripsi ini.

Dengan ini saya melimpahkan hak cipta dari karya tulis saya kepada Institut Pertanian Bogor.

Bogor, Juli 2026

M Raihan Alghani Leksono
G6401221052



ABSTRAK

M RAIHAN ALGHANI LEKSONO. Eksplorasi Pemanfaatan *Generative AI* dalam Pembuatan *User Persona*: Studi Kualitatif pada Praktisi UX. Dibimbing oleh FIRMAN ARDIANSYAH dan SHELVE NIDYA NEYMAN.

Proses pembuatan *user persona* secara tradisional sering kali memakan banyak waktu dan sumber daya. Kehadiran *generative AI* (GenAI) menawarkan efisiensi tinggi untuk mengakselerasi proses tersebut, akan tetapi menyimpan risiko halusinasi dan bias yang dapat menyesatkan keputusan desain. Penelitian ini bertujuan menginvestigasi peran GenAI, memetakan taksonomi strategi instruksi (*prompting*), menganalisis risiko metodologis, serta merumuskan kerangka tata kelola operasional. Data dikumpulkan melalui wawancara mendalam terhadap delapan praktisi *User Experience* (UX) dari berbagai sektor industri dan dianalisis menggunakan metode *Reflexive Thematic Analysis*. Hasil penelitian menunjukkan bahwa GenAI secara fundamental difungsikan sebagai asisten kognitif untuk melepaskan beban pemrosesan data mentah. Meskipun efisien, pemanfaatan mesin terbukti memicu cacat bawaan berupa halusinasi dan bias demografi yang menyeragamkan profil pengguna. Sebagai respons mitigasi, praktisi menerapkan instruksi berbasis batasan empiris. Berdasarkan temuan tersebut, penelitian ini menyintesis *Conceptual Framework for AI-Assisted Persona Creation* sebagai implikasi praktis. Kerangka ini menegaskan bahwa validitas artefak persona mutlak membutuhkan pengawasan *human-in-the-loop* dan penegakan prinsip keterlacakan wawasan (*insight traceability*) ke transkrip asli guna mencegah degradasi kemampuan analitis (*deskilling*) pada periset.

Kata kunci: Analisis tematik refleksif, kerangka konseptual, *large language model* riset pengalaman pengguna, *user persona*.

ABSTRACT

M RAIHAN ALGHANI LEKSONO. *Exploring Generative AI Utilization in User Persona Creation: A Qualitative Study on UX Practitioners*. Supervised by FIRMAN ARDIANSYAH and SHELVE NIDYA NEYMAN.

Traditional process of creating user personas is often time-consuming and resource-intensive. The emergence of generative AI (GenAI) offers high efficiency to accelerate this process, yet it harbors the risk of hallucinations and biases that can mislead design decisions. This qualitative study aims to investigate the role of GenAI, map the taxonomy of prompting strategies, analyze methodological risks, and formulate an operational governance framework. Data were collected through in-depth interviews with eight User Experience (UX) practitioners from various industrial sectors and analyzed using Reflexive Thematic Analysis. The findings reveal that GenAI fundamentally functions as a cognitive assistant to offload the processing of raw data. Although efficient, the machine's utilization is proven to trigger inherent flaws, such as hallucinations and demographic biases that homogenize user profiles. As a mitigation response, practitioners apply constraint-based prompting. Based on these findings, this study synthesized the Conceptual

Framework for AI-Assisted Persona Creation as a practical implication. This framework asserts that the validity of persona artifacts absolutely requires human-in-the-loop oversight and the enforcement of insight traceability to the original transcripts to prevent the analytical degradation (deskilling) of researchers.

Keywords: Conceptual framework, large language model, reflexive thematic analysis, user experience research, user persona.

@Hak cipta milik IPB University

IPB University



IPB University
— Bogor Indonesia —

Hak Cipta Dilindungi Undang-undang

1. Dilarang mengutip sebagian atau seluruh karya tulis ini tanpa mencantumkan dan menyebutkan sumber :
 - a. Pengutipan hanya untuk kepentingan pendidikan, penelitian, penulisan karya ilmiah, penyusunan laporan, penulisan kritik atau tinjauan suatu masalah
 - b. Pengutipan tidak merugikan kepentingan yang wajar IPB University.
2. Dilarang mengumumkan dan memperbanyak sebagian atau seluruh karya tulis ini dalam bentuk apapun tanpa izin IPB University.

Perpustakaan IPB University



@Hak cipta milik IPB University

IPB University



IPB University
— Bogor Indonesia —

Hak Cipta Dilindungi Undang-undang

1. Dilarang mengutip sebagian atau seluruh karya tulis ini tanpa mencantumkan dan menyebutkan sumber :
 - a. Pengutipan hanya untuk kepentingan pendidikan, penelitian, penulisan karya ilmiah, penyusunan laporan, penulisan kritik atau tinjauan suatu masalah
 - b. Pengutipan tidak merugikan kepentingan yang wajar IPB University.
2. Dilarang mengumumkan dan memperbanyak sebagian atau seluruh karya tulis ini dalam bentuk apapun tanpa izin IPB University.

Perpustakaan IPB University

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

M RAIHAN ALGHANI LEKSONO

Skripsi
sebagai salah satu syarat untuk memperoleh gelar
Sarjana pada
Program Studi Ilmu Komputer

**PROGRAM SARJANA ILMU KOMPUTER
SEKOLAH SAINS DATA, MATEMATIKA, DAN INFORMATIKA
INSTITUT PERTANIAN BOGOR
BOGOR
2026**



@Hak cipta milik IPB University

Hak Cipta Dilindungi Undang-undang

1. Dilarang mengutip sebagian atau seluruh karya tulis ini tanpa mencantumkan dan menyebutkan sumber :
 - a. Pengutipan hanya untuk kepentingan pendidikan, penelitian, penulisan karya ilmiah, penyusunan laporan, penulisan kritik atau tinjauan suatu masalah
 - b. Pengutipan tidak merugikan kepentingan yang wajar IPB University.
2. Dilarang mengumumkan dan memperbanyak sebagian atau seluruh karya tulis ini dalam bentuk apapun tanpa izin IPB University.

Penguji pada Ujian Skripsi:
Hafidlotul Fatimah Ahmad, S.Kom., M.Kom.



@Hak cipta milik IPB University

Hak Cipta Dilindungi Undang-undang

1. Dilarang mengutip sebagian atau seluruh karya tulis ini tanpa mencantumkan dan menyebutkan sumber :
 - a. Pengutipan hanya untuk kepentingan pendidikan, penelitian, penulisan karya ilmiah, penyusunan laporan, penulisan kritik atau tinjauan suatu masalah
 - b. Pengutipan tidak merugikan kepentingan yang wajar IPB University.
2. Dilarang mengumumkan dan memperbanyak sebagian atau seluruh karya tulis ini dalam bentuk apapun tanpa izin IPB University.



Judul Skripsi : Eksplorasi Pemanfaatan *Generative AI* dalam Pembuatan *User Persona*: Studi Kualitatif pada Praktisi UX

Nama : M Raihan Alghani Leksono
NIM : G6401221052

@Hak cipta milik IPB University

Disetujui oleh

Pembimbing 1:
Firman Ardiansyah, S.Kom., M.Si.

Pembimbing 2:
Dr. Shelve Nidya Neyman, S.Kom., M.Si.

Diketahui oleh

Ketua Program Sarjana Ilmu Komputer:
Dr. Sony Hartono Wijaya, S.Kom., M.Kom.
NIP 19810809 200812 1 002

Tanggal Ujian:
30 Juli 2026

Tanggal Lulus:



PRAKATA

Puji dan syukur penulis panjatkan kepada Allah *subhanaahu wa ta'ala* atas segala karunia-Nya sehingga karya ilmiah ini berhasil diselesaikan. Penelitian yang dilaksanakan sejak bulan Desember 2025 sampai bulan Juli 2026 ini berjudul “Eksplorasi Pemanfaatan *Generative AI* dalam Pembuatan *User Persona*: Studi Kualitatif pada Praktisi UX”. Penyelesaian karya ilmiah ini tentu tidak terlepas dari bantuan, bimbingan, serta doa dari berbagai pihak.

Rasa terima kasih dan penghargaan setinggi-tingginya penulis sampaikan kepada Bapak Firman Ardiansyah, S.Kom., M.Si. selaku dosen pembimbing I dan Ibu Dr. Shelve Nidya Neyman, S.Kom., M.Si. selaku dosen pembimbing II sekaligus dosen pembimbing akademik yang telah meluangkan waktu dan kesabarannya untuk memberikan bimbingan, arahan, serta masukan yang sangat berharga dan mendalam terkait topik penelitian ini. Ucapan terima kasih juga disampaikan kepada Ibu Hafidlotul Fatimah Ahmad, S.Kom., M.Kom. selaku dosen penguji yang telah memberikan saran, kritik membangun, dan masukan berharga sehingga karya ilmiah ini menjadi jauh lebih baik. Apresiasi turut penulis haturkan kepada para narasumber praktisi UX yang telah bersedia meluangkan waktunya untuk berdiskusi, membagikan wawasan, serta pengalamannya yang sangat berharga bagi kelengkapan data penelitian ini.

Penghargaan dan rasa cinta yang mendalam penulis persembahkan kepada Mama, Papa, Adik, dan seluruh keluarga yang senantiasa memberikan dukungan tiada henti, baik berupa doa, dukungan materi, maupun kekuatan mental kepada penulis selama masa perkuliahan hingga tahap akhir penyelesaian skripsi ini.

Dukungan moril selama proses perjuangan ini juga tidak terlepas dari kehadiran teman-teman terdekat. Penulis mengucapkan terima kasih kepada Ferdinand, Aszriel Teddy, Firoos Likan, Cindy, dan Shafaya sebagai teman seperjuangan yang selalu sedia mendengarkan keluh kesah serta saling memberikan semangat dan dukungan moril selama perkuliahan dan pengerjaan penelitian ini. Apresiasi juga disampaikan kepada teman-teman KKN Tempursari 2025 (Mentari) yang selalu saling mendoakan, menyemangati, dan menjadi *support system* yang baik, serta Althaf dan Faqih selaku adik-adik tingkat perkuliahan yang sangat membantu, menemani, dan senantiasa memberikan dukungan nyata kepada penulis dalam berbagai dinamika perkuliahan hingga di titik penyelesaian penelitian ini.

Rasa terima kasih juga penulis haturkan kepada Ihsan, Monica, Rizka, Rosyid, Indri, Ridho, Ican, Nanta, Brilliant, Gita, Nasywa, teman-teman Program Studi Ilmu Komputer angkatan 59, teman-teman PKU ST31, mahasiswa IPB angkatan 59, teman-teman SMA seperjuangan skripsi tahun ini, serta seluruh pihak yang tidak dapat penulis sebutkan satu per satu atas segala pertemanan, bantuan, doa, dan kenangan selama masa perkuliahan. Semoga karya ilmiah ini dapat memberikan manfaat bagi pihak-pihak yang membutuhkan dan bagi kemajuan ilmu pengetahuan ke depannya.

Bogor, Juli 2026

M Raihan Alghani Leksono



DAFTAR ISI

DAFTAR TABEL	xi
DAFTAR GAMBAR	xi
DAFTAR LAMPIRAN	xi
PENDAHULUAN	1
1.1 Latar Belakang	1
1.2 Rumusan Masalah	2
1.3 Tujuan	2
1.4 Manfaat	3
1.5 Ruang Lingkup	3
TINJAUAN PUSTAKA	5
2.1 <i>User Persona</i> dalam Riset UX	5
2.2 Iterasi Desain dan Artefak Riset UX	5
2.3 <i>Human-AI Collaboration</i> dan Dinamika Socioteknis	6
2.4 <i>Generative AI</i> dalam Desain dan Riset UX	7
2.5 Teknik <i>AI Prompting</i>	8
2.6 Validitas dan Reliabilitas Artefak GenAI	8
METODE	10
3.1 Waktu dan Lokasi Penelitian	10
3.2 Jenis dan Pendekatan Penelitian	10
3.3 Partisipasi Penelitian dan Posisionalitas Peneliti	11
3.4 Instrumen Penelitian dan Prosedur Pengumpulan Data	12
3.5 Teknik Analisis Data	13
3.6 Kredibilitas Data	15
HASIL DAN PEMBAHASAN	17
4.1 Gambaran Umum Pelaksanaan Penelitian	17
4.2 Penerapan <i>Reflexive Thematic Analysis</i> dan Pemetaan Tematis	18
4.3 Dimensi 1: Manfaat Operasional dan Interaksi Kognitif	21
4.4 Dimensi 2: Keterbatasan Struktural dan Risiko Socioteknis	26
4.5 Dimensi 3: Tata Kelola Manusia dan Batasan Operasional	32
4.6 <i>Conceptual Framework for AI-Assisted Persona Creation</i>	37
SIMPULAN DAN SARAN	40
5.1 Simpulan	40
5.2 Saran	40
DAFTAR PUSTAKA	42
LAMPIRAN	46
RIWAYAT HIDUP	61



DAFTAR TABEL

1	Profil delapan partisipan dalam penelitian	17
2	Cuplikan <i>audit trail</i> data transkrip menuju tema final	19
3	Pemetaan tema hasil analisis data kualitatif	20
4	Taksonomi strategi <i>prompting</i> praktisi UX	26
5	Klasifikasi anomali struktural pada artefak persona buatan AI	29
6	Klasifikasi risiko sosioteknis pada penggunaan AI dalam riset UX	31
7	Komparasi batasan fundamental AI dalam perancangan <i>user persona</i>	33
8	Taksonomi strategi mitigasi privasi data	34

DAFTAR GAMBAR

1	Alur pelaksanaan penelitian	10
2	Visualisasi alur kolaborasi <i>Hybrid Intelligence</i> antara praktisi UX dan <i>Generative AI</i>	22
3	Visualisasi <i>backward tracing</i> untuk menegaskan prinsip <i>insight traceability</i> pada luaran AI	24
4	Visualisasi mekanisme transformasi data pada tiga varian halusinasi algoritma GenAI	28
5	Visualisasi eskalasi risiko kognitif pada interaksi <i>human-AI</i>	31
6	Ilustrasi <i>backward tracing</i> dalam audit validitas artefak persona	36
7	Visualisasi SOP mitigasi risiko sosioteknis pada alur riset UX	36
8	<i>Conceptual Framework for AI-Assisted Persona Creation</i>	37

DAFTAR LAMPIRAN

1	Daftar 20 pertanyaan wawancara	47
2	Formulir persetujuan partisipan penelitian	49
3	Isian formulir persetujuan partisipasi penelitian	50
4	<i>Audit trail</i> proses <i>reflexive thematic analysis</i> fase <i>coding</i> menuju tema final	51

Hak Cipta Dilindungi Undang-undang

1. Dilarang mengutip sebagian atau seluruh karya tulis ini tanpa mencantumkan dan menyebutkan sumber :
a. Pengutipan hanya untuk kepentingan pendidikan, penelitian, penulisan karya ilmiah, penyusunan laporan, penulisan kritik atau tinjauan suatu masalah
b. Pengutipan tidak merugikan kepentingan yang wajar IPB University.
2. Dilarang mengumumkan dan memperbanyak sebagian atau seluruh karya tulis ini dalam bentuk apapun tanpa izin IPB University.



@Hak cipta milik IPB University

Hak Cipta Dilindungi Undang-undang

1. Dilarang mengutip sebagian atau seluruh karya tulis ini tanpa mencantumkan dan menyebutkan sumber :
 - a. Pengutipan hanya untuk kepentingan pendidikan, penelitian, penulisan karya ilmiah, penyusunan laporan, penulisan kritik atau tinjauan suatu masalah
 - b. Pengutipan tidak merugikan kepentingan yang wajar IPB University.
2. Dilarang mengumumkan dan memperbanyak sebagian atau seluruh karya tulis ini dalam bentuk apapun tanpa izin IPB University.

I PENDAHULUAN

1.1 Latar Belakang

Dalam praktik desain UX modern, proses desain sering menggunakan pendekatan *User-Centered Design* (UCD), yaitu metode iteratif yang menempatkan pengguna sebagai pusat dalam mendesain. Metode UCD umumnya melibatkan beberapa tahapan penting seperti *understand context of use*, *specify user requirements*, *design solutions*, dan *evaluate design* (Siregar *et al.* 2025). Dari kerangka tahapan UCD itulah konsep *user persona* muncul sebagai alat riset UX yang krusial. Pembuatan *user persona* adalah salah satu komponen penting dalam proses desain UX karena persona mewakili segmen pengguna secara fiktif namun berlandaskan pada data nyata, sehingga membantu tim desain untuk lebih berempati dan membuat keputusan berdasarkan kebutuhan pengguna (Lauer *et al.* 2025). Cooper *et al.* (2014) menjelaskan bahwa persona yang efektif harus dibangun berdasarkan data perilaku nyata yang diperoleh melalui riset kualitatif yang ketat untuk merepresentasikan kebutuhan pengguna secara akurat. Dalam metode tradisional, *user persona* biasanya dibuat melalui riset kualitatif seperti wawancara, observasi, dan studi etnografi yang prosesnya sangat memakan waktu dan sumber daya, terutama ketika tim desain ingin melakukan iterasi cepat.

Kemajuan teknologi digital yang sangat pesat telah mengenalkan manusia dengan kecerdasan buatan atau *artificial intelligence* (AI). Salah satu bentuk AI yang dekat dengan kehidupan sehari-hari dan banyak dikenal saat ini adalah *Generative AI* (GenAI). GenAI, khususnya *Large Language Model* (LLM) menawarkan solusi revolusioner dalam mempercepat proses pembuatan *user persona* karena kemampuannya menghasilkan teks yang kompleks berdasarkan hasil permintaan dari penggunaannya. LLM merupakan sebuah model pembelajaran mesin yang dilatih dengan jumlah data teks yang sangat besar untuk mempelajari, mengenali konteks dan maksudnya, lalu menarik data yang telah dipelajari sebelumnya yang kemudian menghasilkan jawaban dengan bahasa alami secara efektif (Limantara 2025). Kualitas luaran dari GenAI seperti LLM ini sangat bergantung pada kejelasan instruksi atau *prompt* yang diberikan. *Prompt* yang dirancang secara naratif dan alami dapat secara signifikan meningkatkan kemampuan model dalam menyelesaikan *task* dari *user* tanpa contoh secara eksplisit, hal ini membantah asumsi yang menyatakan bahwa *prompting* yang efektif hanya *prompting* yang menyertakan contoh secara eksplisit (Reynolds dan McDonell 2021).

Teknologi seperti LLM ini tentu berpotensi juga membantu berbagai tugas dalam tahapan iterasi desain. Namun, meskipun LLM memiliki kemampuan yang kuat dalam mendukung tugas-tugas generatif dan penalaran kompleks yang bermanfaat bagi proses kerja, penerapannya masih menghadapi tantangan signifikan terkait reliabilitas, halusinasi, dan bias data yang belum sepenuhnya teratasi (Yang *et al.* 2023). Studi terkini telah mulai mengeksplorasi potensi GenAI dalam ranah pembuatan *user persona*, tetapi dengan fokus yang berbeda dari praktik industri. Jiang *et al.* (2023) dalam penelitiannya berfokus pada pengujian kapabilitas teknis LLM untuk mengekspresikan sifat kepribadian (*personality traits*) tertentu secara konsisten dalam simulasi persona. Sementara itu, Strahler (2025) mengeksplorasi pemanfaatan GenAI dalam konteks pendidikan, yaitu

sebagai alat bantu bagi mahasiswa untuk menyimulasikan pembuatan profil pengguna dan meningkatkan kreativitas di kelas komunikasi. Meskipun kedua studi tersebut menunjukkan potensi teknologi ini, literatur yang ada masih terbatas pada ranah eksperimen teknis model atau simulasi pembelajaran.

Masih terdapat kekurangan literatur yang secara spesifik membedah bagaimana praktisi UX profesional menavigasi risiko metodologis seperti bias stereotip dan validitas data saat mengintegrasikan GenAI ke dalam alur kerja di industri yang menuntut akurasi tinggi. Jika persona yang dihasilkan mengandung bias atau halusinasi, penerapan teknologi ini dapat menjadi sia-sia karena pembuatan iterasi desain selanjutnya akan didasarkan pada asumsi yang salah.

Dalam realitas operasional industri saat ini, adopsi teknologi tersebut diwujudkan secara masif melalui pemanfaatan platform komersial berbasis *Large Language Models* (LLM) seperti ChatGPT, Claude, atau Gemini yang diandalkan oleh para praktisi UX sebagai instrumen pengolah teks. Dengan demikian, muncul kebutuhan untuk mengeksplorasi bagaimana platform-platform tersebut berperan konkret sebagai “asisten kognitif” dalam proses pembuatan *user persona*, serta memahami strategi apa yang diterapkan praktisi untuk mengatasi tantangan metodologis terkait reliabilitas, validitas, dan bias. Penelitian ini bertujuan mengidentifikasi praktik penggunaan serta menyintesis sebuah kerangka konseptual tata kelola (*governance framework*) yang mengintegrasikan kecerdasan buatan dengan pengawasan manusia (*human-in-the-loop*).

1.2 Rumusan Masalah

Rumusan masalah dalam penelitian ini mencakup:

- Bagaimana praktik dan strategi pemanfaatan GenAI oleh praktisi UX dalam proses pembuatan *user persona*, serta bagaimana bentuk taksonomi pola pemanfaatan tersebut, khususnya terkait teknik *prompting* dan evaluasi hasil AI?
- Apa saja dimensi kualitas artefak persona yang dipertimbangkan praktisi UX, serta bagaimana bentuk risiko-risiko seperti bias, halusinasi, atau generalisasi berlebihan yang muncul dan dikenali dalam proses tersebut?
- Apa saja kebutuhan, kondisi, dan prinsip dasar yang diperlukan untuk menyusun kerangka konseptual (*conceptual framework*) tata kelola risiko dalam penggunaan GenAI untuk pembuatan *user persona*?

1.3 Tujuan

Adapun tujuan dari penelitian ini adalah:

- Mengidentifikasi praktik dan variasi strategi *prompting* yang digunakan oleh praktisi UX serta mengklasifikasikannya ke dalam taksonomi pola pemanfaatan GenAI dalam pembuatan *user persona*.
- Menganalisis dimensi kualitas artefak persona serta memetakan bentuk-bentuk risiko, seperti bias, halusinasi, dan generalisasi berlebihan, yang muncul dalam proses pembuatan persona berbasis GenAI.
- Merumuskan kerangka konseptual (*conceptual framework*) yang mengintegrasikan penggunaan GenAI dengan mekanisme validasi berbasis manusia untuk meningkatkan reliabilitas artefak persona.

1.4 Manfaat

Penelitian ini diharapkan mampu memberikan kontribusi yang komprehensif, baik secara akademik dalam pengembangan disiplin ilmu maupun secara operasional dalam praktik profesional di industri. Secara lebih terinci, manfaat penelitian ini terbagi ke dalam dua aspek utama:

a Manfaat Teoritis

Secara keilmuan, hasil penelitian ini diharapkan dapat memperkaya literatur akademik pada bidang desain pengalaman pengguna (*User Experience*) dan *Human-Computer Interaction* (HCI). Kontribusi teoretis utamanya terletak pada pendalaman pemahaman mengenai dinamika sistem sosioteknis dalam kolaborasi antara manusia dan kecerdasan buatan (*human-AI collaboration*), khususnya dalam menelaah fenomena *cognitive offloading*, risiko ilusi kecerdasan (*illusion of competence*), serta batasan otonomi algoritma pada proses riset kualitatif. Temuan ini dapat dijadikan rujukan konseptual bagi studi lanjutan yang mengkaji tata kelola etika dan reliabilitas artefak berbasis AI.

b Manfaat Praktis

Secara penerapan, penelitian ini menawarkan rekomendasi tata kelola operasional dan kerangka kerja konseptual (*Conceptual Framework for AI-Assisted Persona Creation*) yang dapat diimplementasikan secara langsung oleh para *UX researcher*, *product designer*, maupun praktisi industri. Melalui panduan metodologis dan mekanisme *human-in-the-loop* yang dirumuskan, praktisi dapat meningkatkan efisiensi waktu sintesis data kualitatif sekaligus memitigasi risiko krusial seperti halusinasi data, bias demografi, serta ancaman degradasi kemampuan analitis (*deskilling*), sehingga validitas artefak riset UX yang dihasilkan tetap dapat dipertanggungjawabkan kebenarannya secara empiris.

1.5 Ruang Lingkup

Ruang lingkup dari penelitian ini di antaranya:

- Penelitian ini berfokus pada pemanfaatan *Large Language Models* (LLM) melalui interaksi berbasis teks dalam proses pembuatan *user persona*. Eksplorasi teknologi ini terbuka pada berbagai tools LLM komersial yang umum digunakan di industri, seperti ChatGPT, Claude, maupun Gemini, dengan penekanan pada praktik penggunaan, strategi *prompting*, serta evaluasi kualitas artefak yang dihasilkan. Analisis artefak dibatasi pada elemen-elemen inti *user persona*, yaitu *goals*, *pains*, *needs*, latar belakang pengguna, serta perilaku penggunaan (*behaviors*).
- Partisipan penelitian merupakan *UX researcher* atau praktisi UX dengan pengalaman kerja profesional di bidang riset atau perancangan UX minimal 1 tahun, serta sudah pernah menggunakan GenAI dalam proses pembuatan *user persona*, sehingga dapat memberikan wawasan berbasis pengalaman nyata.
- Penelitian tidak mencakup *high-fidelity prototyping*, *usability testing*, maupun implementasi sistem agar tetap fokus pada penggunaan GenAI untuk pembuatan *user persona*.
- Data dibatasi pada wawancara mendalam (*in-depth interview*) dengan praktisi UX dan analisis difokuskan pada pemetaan pola perilaku, strategi penggunaan

GenAI, identifikasi risiko, serta perancangan kerangka konseptual tata kelola berbasis temuan kualitatif.



@Hak cipta milik IPB University

IPB University



IPB University
— Bogor Indonesia —

Hak Cipta Dilindungi Undang-undang

1. Dilarang mengutip sebagian atau seluruh karya tulis ini tanpa mencantumkan dan menyebutkan sumber :
 - a. Pengutipan hanya untuk kepentingan pendidikan, penelitian, penulisan karya ilmiah, penyusunan laporan, penulisan kritik atau tinjauan suatu masalah
 - b. Pengutipan tidak merugikan kepentingan yang wajar IPB University.
2. Dilarang mengumumkan dan memperbanyak sebagian atau seluruh karya tulis ini dalam bentuk apapun tanpa izin IPB University.

Perpustakaan IPB University

II TINJAUAN PUSTAKA

2.1 *User Persona* dalam Riset UX

User persona merupakan representasi fiktif namun realistis dari beberapa *target user* yang berfungsi untuk memodelkan tujuan, motivasi, dan perilaku *user* secara spesifik. Cooper *et al.* (2014) menekankan bahwa persona yang efektif tidak boleh dibangun berdasarkan asumsi (*assumption-based*), melainkan harus disintesis secara ketat dari data perilaku nyata. Dalam praktiknya, persona berfungsi sebagai alat komunikasi strategis untuk membangun empati tim desain dan mencegah bias perancangan yang berorientasi pada selera desainer itu sendiri (*self-referential design*).

Dalam pendekatan tradisional, pembuatan persona yang memenuhi standar efektivitas tersebut memerlukan proses penelitian etnografi yang sangat ketat (*rigorous*) dan memakan waktu. Proses ini bermula dari pengumpulan data kualitatif berupa wawancara, analisis pola perilaku, hingga sintesis manual menjadi narasi karakter yang hidup. Kualitas persona ini sangat bergantung pada validitas data yang mendasarinya. Persona yang dibangun tanpa landasan data empiris sering kali dikritik sebagai produk yang “*fictional but not realistic*”, karena gagal merepresentasikan pengguna aktual sehingga lemah untuk menurunkan wawasan desain yang dapat diandalkan (Paoli 2023). Persona semacam ini pada akhirnya bermuara pada asumsi yang justru dapat menyesatkan keseluruhan proses desain (Cooper *et al.* 2014).

Untuk mengatasi keterbatasan metode tradisional yang rentan bias dan sering kali hanya bersandar pada sampel kualitatif yang kecil, praktik perumusan persona kini mengalami pergeseran masif menuju konstruksi berbasis *big data* dan AI generatif. Pengelompokan atribut pengguna tidak lagi murni mengandalkan sintesis manual, melainkan mulai memanfaatkan metode statistik pada set data berskala besar untuk mengubah data empiris menjadi profil yang jauh lebih objektif dan terukur (Molléri dan Marculescu 2025). Pendekatan *data-driven persona* ini secara langsung memperkuat prinsip dasar Cooper *et al.* (2014) dengan memastikan bahwa fondasi perancangan benar-benar berakar pada pola perilaku aktual dalam skala yang lebih komprehensif.

Evolusi *data-driven persona* ini pada hakikatnya merupakan manifestasi langsung dari visi *Human-Centered Artificial Intelligence* (HCAI). Mengacu pada paradigma *Human-Centered Artificial Intelligence* (HCAI) (Shneiderman 2022; Bevilacqua *et al.* 2025), penggunaan teknologi cerdas seperti algoritma dan LLM dalam perumusan persona tidak ditujukan untuk mensubstitusi empati atau peran riset peneliti UX, melainkan berfungsi untuk mengamplifikasi dan memberdayakan kapabilitas analitis manusia. Dengan demikian, teknologi hadir sebagai mitra strategis yang menjembatani data empiris skala besar dengan kedalaman pemahaman emosional, sehingga keputusan perancangan yang diambil tetap berada dalam ruang kendali dan intuisi manusia yang bermakna.

2.2 Iterasi Desain dan Artefak Riset UX

Dalam siklus pengembangan produk digital, kecepatan dan kemampuan untuk melakukan iterasi desain adalah kunci keberhasilan. Nielsen (1994) melalui konsep *usability engineering* menekankan pentingnya siklus desain iteratif ketika

evaluasi dilakukan seawal mungkin dalam proses desain. Pendekatan ini bertujuan mengidentifikasi dan memperbaiki masalah kegunaan sebelum biaya perbaikan menjadi terlalu tinggi pada tahap pengembangan lanjut.

Pada konsep iterasi ini, artefak riset UX seperti *user persona* berfungsi sebagai panduan statis yang harus terus diperbarui seiring dengan bertambahnya pemahaman tentang pengguna. Keterlambatan dalam penyusunan artefak ini sering kali menjadi hambatan dalam siklus iterasi yang cepat. Sebelum masuknya era kecerdasan buatan generatif, hambatan waktu dalam penyusunan artefak ini diatasi melalui pergeseran paradigma metodologis dan adopsi peranti lunak kolaboratif. Gothelf (2013) menginisiasi pendekatan *Lean UX* yang secara radikal memangkas ketergantungan pada dokumentasi statis yang tebal, mendorong tim untuk memfokuskan energi pada siklus eksperimentasi dan pembaruan artefak secara cepat. Senada dengan gagasan tersebut, integrasi *User-Centered Design* (UCD) ke dalam lingkungan Agile sangat bertumpu pada peranti kolaboratif untuk memfasilitasi pembaruan artefak (Salah *et al.* 2014), dengan pemanfaatan artefak terbagi (*shared artifacts*) tersebut terbukti secara empiris mampu memediasi penyelarasan pemahaman (*alignment work*) serta mereduksi hambatan komunikasi antara desainer dan pengembang (Brown *et al.* 2011).

Meskipun pergeseran metodologi dan adopsi peranti kolaboratif telah berhasil mengeliminasi hambatan komunikasi serta mempercepat distribusi artefak, tantangan mendasar pada tahapan kognitif tetap belum terselesaikan. Proses menyintesis data riset kualitatif yang mentah menjadi sebuah narasi *user persona* yang utuh masih murni mengandalkan proses analitis manual yang memakan banyak waktu. Celah efisiensi operasional inilah yang pada akhirnya menjadi katalis utama bagi industri untuk mulai mengeksplorasi potensi *Generative AI*.

2.3 Human-AI Collaboration dan Dinamika Sosioteknis

Integrasi *Generative AI* dalam alur kerja desain bukan sekadar penggantian alat, melainkan sebuah perubahan mendasar dalam sistem sosioteknis atau *socio-technical systems* (Mumford 2006). Dalam kerangka *human-AI collaboration*, praktisi UX pada dasarnya mendelegasikan tugas-tugas generatif kepada mesin untuk memperluas ruang eksplorasi ide (Dellermann *et al.* 2019). Dalam era kecerdasan buatan modern, LLM tidak lagi bertindak pasif, melainkan menjadi agen yang berpartisipasi aktif dalam pemecahan masalah yang lebih kompleks (Raisch dan Fomina 2024).

Delegasi tugas-tugas generatif yang dilakukan secara masif ini memicu fenomena *cognitive offloading*, yaitu pemindahan beban kerja mental internal kepada alat eksternal (Risko dan Gilbert 2016). Dalam konteks penyusunan artefak riset seperti *user persona*, fenomena ini berisiko tinggi memicu *automation bias*. Kondisi tersebut terjadi ketika praktisi mempercayai *output* dari mesin secara membabi buta dan mengabaikan informasi kontradiktif dari data lapangan (Parasuraman dan Manzey 2010).

Risiko tersebut semakin nyata karena *Generative AI* dirancang untuk memiliki pengetahuan instrumental yang fasih dalam merangkai bahasa dan menyimulasikan struktur tugas, meskipun sistem tersebut sebenarnya tidak memiliki model mental tentang realitas dunia atau *world models* (Yildirim dan Paul 2024). Kefasihan instrumental mesin ini sering kali menciptakan ilusi pemahaman

di mata peneliti manusia, sehingga manusia memangkas proses evaluasi kritisnya demi mengejar kecepatan operasional.

Keberhasilan konsep *human-AI collaboration* sangat bergantung pada tingkat kepercayaan yang terkalibrasi (Hoff dan Bashir 2015). Kalibrasi kepercayaan artinya praktisi secara sadar menerapkan mekanisme pengawasan kritis terhadap hasil keluaran mesin dan tidak memercayai GenAI secara mutlak. Tanpa kalibrasi ini, praktisi berisiko terjebak dalam fenomena otoritas algoritmik (*algorithmic authority*), ketika manusia cenderung memberikan legitimasi kebenaran kepada mesin dan menerima keluaran algoritma sebagai sebuah fakta tanpa peninjauan empiris (Dwivedi *et al.* 2023). Oleh karena itu, pemanfaatan *Generative AI* dalam riset UX memerlukan *human-in-the-loop* yang aktif untuk menjaga otonomi strategis manusia dalam pengambilan keputusan desain.

Hal ini sejalan dengan batasan fundamental kecerdasan buatan yang dikemukakan oleh Floridi (2023), bahwa model bahasa besar sama sekali tidak berpikir, bernalar, atau memiliki pemahaman sesungguhnya. Keberhasilan AI saat ini hanyalah bentuk pelepasan agensi dari kecerdasan (*liberated agency from intelligence*), yaitu ketika mesin memproses teks secara statistik berdasarkan struktur formalnya, bukan pada maknanya. Karena AI beroperasi murni tanpa adanya pemahaman, Floridi (2023) menegaskan bahwa manusialah yang harus bertindak sebagai “*shepherds of AI systems*” yang memegang kendali penuh atas agensi buatan tersebut untuk menghindari bias dan halusinasi.

2.4 *Generative AI* dalam Desain dan Riset UX

Perkembangan *Generative AI* (GenAI), khususnya yang berbasis *Large Language Models* (LLM), menawarkan potensi baru dalam mendukung aktivitas desain dan riset. Yang *et al.* (2023) menyajikan panduan komprehensif bagi praktisi yang bekerja dengan LLM dalam berbagai tugas aplikasi pemrosesan bahasa (*downstream NLP tasks*), yang mencakup pembuatan teks dan kemampuan penalaran. Kapabilitas ini memungkinkan GenAI berperan dalam menyintesis informasi, menghasilkan ide-ide awal, hingga membantu pembuatan draf artefak desain.

Namun, penerapan teknologi AI ini tentu tidak lepas dari risiko yang melekat pada arsitektur model komputasinya. Yang *et al.* (2023) menyoroti tantangan signifikan terkait reliabilitas model, kecenderungan untuk menghasilkan informasi fiktif (*hallucination*), serta adanya bias data yang diwarisi dari korpus pelatihan. Peringatan ini secara empiris diperkuat oleh kajian komprehensif dari Ji *et al.* (2023), yang menegaskan bahwa fenomena halusinasi pada *Natural Language Generation* (NLG) terjadi ketika model memproduksi teks yang terdengar fasih secara tata bahasa namun melenceng dari sumber atau konteks yang diberikan. Lebih jauh lagi, Bender *et al.* (2021) dalam studinya mengkritisi bahaya laten dari model bahasa berskala besar yang beroperasi semata-mata sebagai beo stokastik (*stochastic parrots*). Mereka memperingatkan bahwa LLM sekadar merangkai bentuk linguistik berdasarkan probabilitas statistik tanpa sedikit pun memiliki pemahaman semantik, sehingga sangat rentan memproduksi ulang bias stereotip yang tertanam di dalam data pelatihannya.

Mengacu pada konvergensi ketiga literatur tersebut, ketergantungan penuh pada GenAI tanpa adanya mekanisme verifikasi manusia di dalam riset UX menjadi praktik yang sangat berbahaya. Praktik tersebut berpotensi menghasilkan wawasan

yang tampak logis dan meyakinkan, padahal secara faktual cacat dan bias, yang pada akhirnya justru akan mendegradasi kualitas keputusan desain produk.

Selain isu validitas dan bias, pemanfaatan LLM juga membawa risiko kritical terhadap keamanan privasi. Literatur keamanan informasi membuktikan bahwa arsitektur LLM memiliki memori laten dengan kapasitas untuk menghafal (*memorize*) dan tanpa sengaja memproduksi ulang data sensitif atau teks rahasia yang pernah masuk ke dalam korpus pelatihannya (Carlini et al. 2021; Brown et al. 2022). Karakteristik komputasional ini memicu ancaman kebocoran data, sehingga menuntut adanya protokol pengecualian data yang ketat ketika model generatif digunakan dalam ranah riset industri yang melibatkan informasi rahasia.

2.5 Teknik AI Prompting

Kualitas *output* yang dihasilkan oleh model GenAI sangat bergantung pada formulasi instruksi atau *prompt* yang diberikan oleh pengguna. Reynolds dan McDonell (2021) mendefinisikan *prompt programming* sebagai teknik memprogram model bahasa melalui instruksi teks naratif untuk menyelesaikan tugas spesifik tanpa perlu mengubah parameter model itu sendiri.

Efektivitas *prompting* sering kali diasosiasikan dengan pemberian contoh (*few-shot*). Namun, ditemukan bahwa instruksi yang bersifat naratif dan alami juga dapat menghasilkan kinerja yang setara atau bahkan lebih baik dalam kondisi *zero-shot* (tanpa contoh), asalkan konteks dan batasan tugas didefinisikan dengan jelas. Dalam pembuatan persona, penerapan teknik *prompting* yang terstruktur seperti memberikan deskripsi peran dan batasan format yang spesifik, diduga dapat meningkatkan relevansi hasil dibandingkan dengan instruksi yang bersifat umum.

Sebagai ilustrasi praktis dalam ranah riset UX, instruksi yang terstruktur secara ideal meramu empat elemen taktis, yaitu penetapan peran (*role*), konteks operasional (*context*), tugas sintesis (*task*), dan batasan empiris (*constraint*). Sebagai contoh, seorang praktisi dapat merumuskan instruksi spesifik seperti: "Bertindaklah sebagai periset UX senior. Berdasarkan data transkrip wawancara terlampir, sintesis pola perilaku tersebut menjadi satu draf *user persona* fiktif yang mencakup rumusan *pain points*, *needs*, dan *goals*. Jangan menambahkan asumsi apa pun di luar teks terlampir, dan wajib sertakan kutipan langsung (*verbatim*) dari pengguna untuk memvalidasi setiap poin masalah." Dengan menggunakan kerangka instruksi yang berorientasi pada batasan empiris seperti ini, periset dapat meminimalisasi kebebasan algoritma dalam memanipulasi teks, sehingga draf karakter yang direproduksi tetap memiliki validitas yang dapat dipertanggungjawabkan.

2.6 Validitas dan Reliabilitas Artefak GenAI

Meskipun LLM mampu memproduksi teks yang fasih, validitas kontennya sebagai artefak riset masih menjadi perdebatan. Zamfirescu-Pereira *et al.* (2023) dalam studi eksperimentalnya menemukan bahwa interaksi dengan LLM sering kali menemukan kasus bahwa perubahan kecil pada *prompt* dapat menghasilkan variasi luaran yang drastis. Studi tersebut juga mengungkap fenomena *over-trust* di kalangan pengguna non-ahli, yang sering kali gagal mengidentifikasi kesalahan logika atau fakta karena terbuai oleh kefasihan bahasa yang dihasilkan AI.

Tantangan utama dalam penggunaan GenAI untuk pembuatan persona adalah risiko *overgeneralization* dan stereotip. Pengguna sering kali kesulitan

mengevaluasi kualitas luaran AI secara kritis, yang berpotensi meloloskan artefak desain yang bias atau tidak akurat ke dalam proses pengembangan produk (Zamfirescu-Pereira *et al.* 2023). Oleh karena itu, validasi manual dengan data lapangan tetap menjadi pertimbangan dalam pemanfaatan GenAI.

Sebagai respons terhadap kerentanan tersebut, literatur terkini telah merumuskan berbagai kerangka evaluasi kualitatif untuk meninjau kebenaran faktual luaran mesin. Truss (2026) memperkenalkan kerangka evaluasi berbasis keterlacakan wawasan (*insight traceability*), yang menetapkan bahwa setiap elemen kognitif pada artefak desain buatan AI wajib diuji silang menggunakan pemetaan langsung terhadap suara pengguna (*Voice of Customer*). Apabila sebuah klaim algoritma tidak memiliki jejak rujukan pada data empiris, klaim tersebut harus dieliminasi. Pendekatan ini diperkuat oleh parameter evaluasi yang diajukan oleh Salminen *et al.* (2024), yang mensyaratkan pengukuran tingkat representasi dan akurasi faktual luaran AI dengan membandingkannya secara langsung terhadap distribusi demografi dan perilaku pengguna asli di lapangan. Melalui penerapan metrik evaluasi yang mewajibkan pembuktian empiris ini, proses validasi manual akan bergerak menjadi sebuah protokol audit yang terstruktur guna memastikan tata kelola manusia tetap berdaulat atas keputusan desain (Liao dan Vaughan 2024).





III METODE

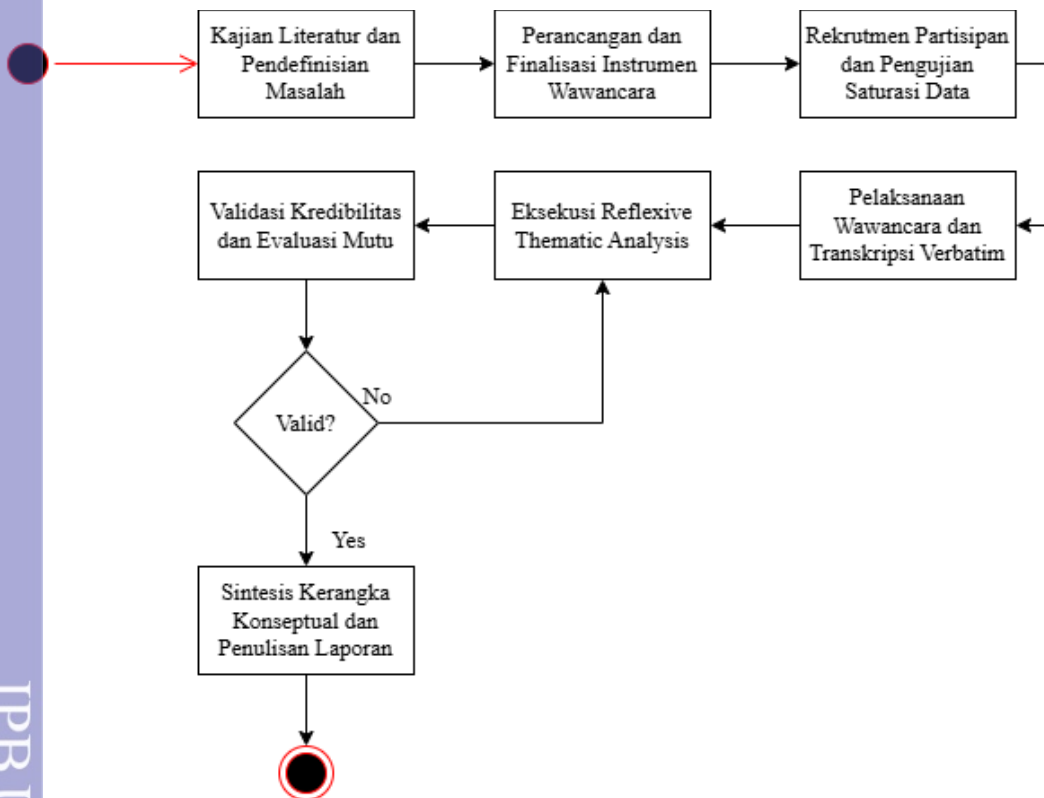
3.1 Waktu dan Lokasi Penelitian

Penelitian ini dilaksanakan selama lebih kurang delapan bulan, terhitung sejak bulan Desember 2025 hingga Juli 2026. Tahapan penelitian mencakup perancangan instrumen, pengumpulan data, hingga proses analisis data dan penulisan laporan akhir. Pengumpulan data kualitatif berupa wawancara mendalam dilakukan sepenuhnya secara daring melalui platform *video conference*.

3.2 Jenis dan Pendekatan Penelitian

Penelitian ini menggunakan pendekatan kualitatif eksploratif sebagai landasan utama. Pendekatan eksploratif dipilih karena fenomena penggunaan GenAI untuk menghasilkan artefak riset UX masih relatif baru dan belum memiliki kerangka teoritis yang mapan. Penelitian ini difokuskan untuk menggali dan memetakan pola perilaku, strategi *prompting*, persepsi kritis praktisi terhadap kualitas dan risiko artefak persona yang dihasilkan, serta tantangan yang dihadapi praktisi dalam mengintegrasikan AI ke dalam alur kerja desain.

Pelaksanaan penelitian ini secara filosofis berpijak pada paradigma konstruktivis dan interpretivis yang diselaraskan dengan metode *Reflexive Thematic Analysis* (Braun dan Clarke 2020). Guna mencapai tujuan tersebut, pelaksanaan penelitian dirancang ke dalam sebuah alur kerja sistematis yang terdiri atas tujuh tahapan operasional yang saling berkesinambungan. Visualisasi dari keseluruhan tahapan tersebut ditunjukkan pada Gambar 1.



Gambar 1 Alur pelaksanaan penelitian

- a Kajian Literatur dan Pendefinisian Masalah
Tahapan ini diawali dengan proses pencarian dan penyeleksian literatur melalui berbagai basis data akademik, seperti ACM Digital Library, ScienceDirect, SpringerLink, IEEE Xplore, serta Google Scholar. Pencarian literatur dieksekusi menggunakan kombinasi kata kunci spesifik, contohnya seperti "*Generative AI*", "*Large Language Models*", "*User Persona*", "*UX Research*", atau "*Human-AI Collaboration*". Hasil kurasi dari proses penelusuran ini kemudian digunakan untuk merumuskan kesenjangan riset (*research gap*) dan mendeduksi batasan masalah secara lebih presisi.
- b Perancangan dan Finalisasi Instrumen Wawancara
Instrumen berupa panduan wawancara semi-terstruktur disusun dan dievaluasi. Pertanyaan dikelompokkan ke dalam empat dimensi esensial, yaitu praktik penggunaan, strategi *prompting*, evaluasi risiko, dan mekanisme mitigasi, guna memastikan proses penggalian data di lapangan tetap terarah. Adapun jumlah pertanyaan wawancara sejumlah 20 butir terdapat pada Lampiran 1.
- c Rekrutmen Partisipan dan Pengujian Saturasi Data
Pencarian narasumber yang memenuhi kriteria dilakukan menggunakan kombinasi teknik *purposive sampling* dan *snowball sampling*. Perekrutan dan pengumpulan data dievaluasi secara konstan hingga mencapai titik saturasi data (*data saturation*) (Guest *et al.* 2006; Hennink dan Kaiser 2022).
- d Pelaksanaan Wawancara dan Transkripsi *Verbatim*
Pengumpulan data empiris dilaksanakan melalui wawancara daring secara langsung bersama para praktisi UX. Seluruh percakapan kemudian didokumentasikan ke dalam bentuk transkrip *verbatim* (kata per kata) dan disandingkan dengan catatan lapangan (*field notes*) untuk menjaga keutuhan konteks informasi.
- e Eksekusi *Reflexive Thematic Analysis*
Transkrip data yang telah terkumpul diolah menggunakan metode analisis tematik reflektif (Braun dan Clarke 2006, 2020). Teks dibedah secara induktif melalui tahapan familiarisasi, pemberian kode awal (*initial coding*), pencarian tema, peninjauan ulang, hingga pendefinisian tema-tema final yang merepresentasikan realitas pengalaman partisipan.
- f Validasi Kredibilitas dan Evaluasi Mutu
Upaya penjaminan keabsahan data dilakukan secara paralel melalui penyusunan *audit trail* (Nowell *et al.* 2017), penyusunan *analytic memos*, serta evaluasi struktur tema bersama dosen pembimbing (*peer debriefing*). Jika masih terdapat ketidaksesuaian pada tahap ini, tahap Eksekusi *Reflexive Thematic Analysis* akan dijalankan Kembali.
- g Sintesis Kerangka Konseptual dan Penulisan Laporan
Tahap terakhir berupa sintesis seluruh temuan dan tema final menjadi sebuah model tata kelola risiko yang utuh, yaitu *Conceptual Framework for AI-Assisted Persona Creation*. Model ini kemudian dijabarkan secara komprehensif ke dalam laporan akhir penelitian.

3.3 Partisipasi Penelitian dan Posisionalitas Peneliti

Subjek dalam penelitian ini adalah praktisi di bidang *User Experience* (UX) yang memiliki pengalaman langsung menggunakan *Generative AI* pada tahapan riset desain. Perekrutan partisipan dilakukan menggunakan teknik *purposive*

sampling pada fase awal, yang kemudian diperluas melalui *snowball sampling*. Teknik tersebut dipilih untuk menjangkau narasumber spesifik yang dapat memberikan wawasan mendalam serta memanfaatkan rekomendasi jaringan praktisi dengan pengalaman serupa di industri. Kriteria inklusi yang ditetapkan bagi partisipan meliputi:

- a. Berprofesi aktif sebagai *UX Researcher*, *Product Designer*, atau peran sejenis minimal 1 tahun dalam kegiatan riset atau desain UX.
- b. Memiliki pengalaman menggunakan *Generative AI* (seperti ChatGPT, Claude, atau Gemini) secara spesifik untuk menghasilkan, mengedit, atau mengevaluasi elemen *user persona*.
- c. Bersedia mengikuti sesi wawancara mendalam dan menyetujui perekaman data untuk keperluan analisis akademik.

Dalam proses pengumpulan data tersebut, posisionalitas (*reflexivity statement*) sebagai instrumen utama (*human instrument*) perlu dideklarasikan guna menjaga transparansi dan kredibilitas penelitian. Peneliti berstatus sebagai mahasiswa sarjana Ilmu Komputer di Sekolah Sains Data, Matematika, dan Informatika (SSMI) Institut Pertanian Bogor (IPB). Latar belakang akademis ini memberikan pemahaman teknis fundamental mengenai kecerdasan buatan, batasan model LLM, serta pengetahuan teoretis dalam perancangan UX, yang memberikan keuntungan analitis dalam mencerna istilah teknis dan logika *prompting* dari para partisipan.

Meskipun demikian, latar belakang tersebut berpotensi memunculkan kecenderungan pro-teknologi yang dapat memengaruhi proses interpretasi data. Sejalan dengan paradigma konstruktivis dan metode *Reflexive Thematic Analysis* yang menuntut pelibatan subjektivitas peneliti secara sadar, pengelolaan terhadap potensi bias ini dilakukan melalui praktik reflektivitas aktif. Peneliti menggunakan catatan reflektif (*reflexive memoing*) secara kontinu di sepanjang proses pengodean hingga penyusunan tema untuk menyadari, mengelola, dan mengendalikan asumsi pribadi. Tidak terdapat hubungan hierarkis atau konflik kepentingan profesional antara peneliti dan partisipan, sehingga seluruh transkrip dianalisis secara kritis agar narasi yang dikonstruksi benar-benar berakar pada realitas empiris di lapangan.

3.4 Instrumen Penelitian dan Prosedur Pengumpulan Data

Instrumen utama dalam penelitian ini yaitu peneliti yang bertindak sebagai *human instrument* (instrumen manusia). Peneliti secara spesifik berperan untuk memandu arah wawancara, menyesuaikan alur diskusi secara adaptif, serta melakukan *probing* (penggalian data) guna mengeksplorasi wawasan mendalam terkait pengalaman praktis partisipan dalam menggunakan *Generative AI*. Sebagai instrumen pendukung, panduan wawancara semi-terstruktur digunakan dan disusun berdasarkan rumusan masalah serta tujuan penelitian. Panduan ini memberi kerangka yang konsisten bagi seluruh sesi wawancara, sekaligus memberikan ruang fleksibilitas bagi peneliti untuk mengeksplorasi isu teknis yang muncul secara spontan. Struktur pertanyaan pada panduan dikelompokkan ke dalam empat dimensi esensial, yaitu:

- a. Praktik dan motivasi penggunaan GenAI: mengeksplorasi kapan, mengapa, dan *tools* apa yang digunakan partisipan dalam pembuatan persona.

- b Strategi *prompting*: membedah teknik spesifik interaksi yang diterapkan (seperti *role-playing*, pemberian konteks, atau penyertaan data mentah) serta iterasi instruksi yang dilakukan.
- c Evaluasi kualitas dan risiko artefak persona: mengukur bagaimana partisipan menilai akurasi hasil AI, mendeteksi halusinasi, dan melakukan validasi dengan data riil.
- d Mekanisme mitigasi risiko: menggali pandangan dan protokol partisipan untuk mengatasi risiko bias algoritma dan stereotip pada persona buatan AI. Keseluruhan butir pertanyaan dalam panduan wawancara tersebut dilampirkan pada Lampiran 1.

Prosedur pengumpulan data di lapangan dieksekusi melalui tiga tahapan utama. Tahap pertama adalah pra-wawancara. Pada tahap ini, calon partisipan yang memenuhi kriteria dihubungi dan diberikan penjelasan mengenai tujuan penelitian. Protokol etika diterapkan melalui pemberian formulir persetujuan (*informed consent*). Formulir ini berfungsi untuk memastikan bahwa partisipan menyetujui perekaman sesi wawancara dan penggunaan data untuk keperluan analisis akademik.

Tahap kedua berfokus pada pelaksanaan wawancara mendalam yang diselenggarakan sepenuhnya secara daring melalui platform *Zoom Meeting*. Setiap sesi wawancara ini memakan durasi berkisar antara 40 hingga 60 menit per responden, dengan alur penggalan data yang dipandu secara dinamis menggunakan instrumen semi-terstruktur. Selama proses interaksi tersebut berlangsung, seluruh percakapan didokumentasikan ke dalam format rekaman audio dan video dengan berpijak pada persetujuan awal partisipan. Bersamaan dengan itu, penyusunan catatan lapangan (*field notes*) turut dilaksanakan sebagai langkah metodologis pendukung untuk menangkap nuansa kontekstual, perubahan ekspresi, serta berbagai dinamika informasi non-verbal yang berpotensi luput apabila hanya mengandalkan rekaman bawaan dari aplikasi *Zoom Meeting*.

Tahap ketiga adalah transkripsi. Seluruh rekaman wawancara ditranskripsikan secara *verbatim* (kata per kata) untuk menjamin kelengkapan dan keaslian data. Guna mengakselerasi proses dokumentasi teks, peneliti memanfaatkan perangkat lunak *PDFSimpli AI Transcriber* sebagai instrumen transkripsi draf awal. Sebelum proses unggah file audio dilakukan ke platform eksternal tersebut, peneliti terlebih dahulu menerapkan protokol privasi ketat dengan menyamarkan atau menghapus seluruh informasi sensitif, nama perusahaan, serta atribut identitas partisipan guna memitigasi risiko kebocoran data. Meskipun demikian, untuk menjaga integritas data kualitatif secara mutlak, peneliti memegang kendali penuh melalui proses validasi manual dengan mendengarkan ulang seluruh rekaman dan memperbaiki anomali translasi mesin. Transkrip final ini kemudian disandingkan dengan catatan lapangan guna memelihara akurasi konteks non-verbal.

3.5 Teknik Analisis Data

Analisis data dalam penelitian ini dilaksanakan menggunakan metode *Reflexive Thematic Analysis (Reflexive TA)* yang berakar pada kerangka kerja Braun dan Clarke (2020). Pendekatan induktif (*bottom-up*) diterapkan secara ketat agar tema yang dihasilkan benar-benar merepresentasikan realitas empiris

partisipan mengenai strategi *prompting*, evaluasi kualitas artefak, dan mekanisme mitigasi risiko, tanpa memaksakan kerangka teori yang sudah ada sebelumnya.

Guna membedah transkrip data tersebut, proses analisis dieksekusi melalui enam tahapan sistematis sebagai berikut:

a Familiarisasi Data

Tahapan ini diawali dengan pembacaan seluruh transkrip *verbatim* dan catatan lapangan secara berulang dan mendalam. Pada fase ini, pemahaman terhadap konteks utuh dari setiap jawaban partisipan dibangun, diiringi dengan pencatatan gagasan awal mengenai pola-pola pemanfaatan *Generative AI* yang berulang.

b Pembangunan Kode Awal

Teks yang telah dipahami, kemudian dibedah melalui proses *inductive open coding*. Proses pengodean dilakukan berbasis *meaning unit*, dengan setiap segmen percakapan yang mengandung wawasan spesifik diekstrak dan diberikan label (*coding*). Pengodean ini dibantu menggunakan perangkat lunak Google Spreadsheet sebagai instrumen manajemen data kualitatif. Pendekatan semantik digunakan untuk memetakan taktik operasional eksplisit (seperti strategi *prompting* dan prosedur anonimisasi data), sedangkan pendekatan laten digunakan untuk mengonstruksi fenomena psikologis yang tersirat (seperti ilusi kecerdasan atau degradasi kemampuan analitis/*deskilling*).

c Pencarian Tema

Puluhan kode awal yang telah terbentuk dan memiliki kedekatan makna konseptual kemudian dikelompokkan ke dalam beberapa tema potensial. Pemetaan ini mencakup klaster mengenai peran operasional AI, pola halusinasi artefak, hingga batasan fundamental sistem.

d Peninjauan Ulang Tema

Peninjauan ulang dilakukan secara ketat dengan mengembalikan setiap draf tema ke dalam konteks percakapan utuh pada saat wawancara. Data yang terindikasi berada di luar fokus riset pembuatan *user persona* dieliminasi. Segala bentuk konflik interpretasi atau ambiguitas diselesaikan melalui metode perbandingan konstan (*constant comparison method*) (Charmaz 2014), yakni dengan mengeksekusi uji silang secara langsung antara kode yang ambigu pada satu transkrip praktisi terhadap transkrip praktisi lainnya untuk memverifikasi konsistensi pola empiris mereka.

e Pendefinisian dan Penamaan Tema

Tema-tema potensial disintesis dan dirumuskan dari hasil peninjauan ulang tema secara konkret dan sesuai dengan konteks temuan penelitian. Setiap tema diformulasikan secara konseptual dan diberikan penamaan yang secara presisi mencerminkan realitas sosioteknis yang terjadi di lapangan.

f Penulisan Laporan

Tahap akhir dalam *Reflexive Thematic Analysis* yaitu penyusunan narasi analitik yang menceritakan keseluruhan temuan secara terstruktur. Tema-tema yang telah terbentuk tersebut kemudian diorganisasikan ke dalam dimensi-dimensi makro yang logis untuk menyusun narasi hasil penelitian. Laporan akhir ini akan diperkuat dengan penyertaan kutipan *verbatim* yang paling representatif dari proses pengumpulan data.

Sebagai tahapan sintesis akhir, tema-tema yang telah divalidasi tersebut dipetakan kembali untuk merumuskan kerangka tata kelola operasional. Tema yang

mencerminkan taktik keberhasilan diterjemahkan menjadi *best practices*, sementara luaran yang merepresentasikan kegagalan algoritma dikonversi menjadi *common pitfalls*. Hasil sintesis dari analisis tematik ini membuahkan *Conceptual Framework for AI-Assisted Persona Creation* yang akan dibahas secara komprehensif pada bab hasil dan pembahasan.

3.6 Kredibilitas Data

Keabsahan data merupakan aspek fundamental dalam penelitian kualitatif guna memastikan bahwa temuan yang dihasilkan benar-benar merepresentasikan realitas empiris partisipan dan tidak hanya sekadar proyeksi bias peneliti. Dalam penelitian ini, serangkaian protokol evaluasi kritis diterapkan secara komprehensif di sepanjang fase pengumpulan hingga analisis data.

Langkah pertama adalah penerapan validasi silang. Validasi ini dieksekusi dengan cara menyilangkan narasi, pengalaman, dan perspektif antara kedelapan praktisi UX yang berasal dari latar belakang sektor industri yang berbeda. Melalui perbandingan ini, pencarian terhadap kasus negatif (*negative cases*) dilakukan secara aktif. Kasus negatif diidentifikasi ketika terdapat anomali perilaku atau pandangan yang bertentangan dengan mayoritas pola data, seperti perbedaan strategi *zero-shot prompting* di fase awal. Kehadiran kasus negatif ini tidak dieliminasi, melainkan dianalisis untuk memperkaya pemahaman bahwa interaksi dengan *Generative AI* memiliki dinamika operasional yang beragam.

Langkah kedua adalah dokumentasi kognitif dan prosedural melalui penyusunan *audit trail* serta *reflexive memoing* (Nowell *et al.* 2017). Keterlacakan data (*insight traceability*) dipastikan tersedia secara transparan melalui *audit trail*, yakni setiap transformasi dari kutipan *verbatim* di dalam transkrip menuju draf kode awal, hingga menjadi tema final, selalu didokumentasikan. Bersamaan dengan itu, catatan analitis (*reflexive memoing*) disusun sepanjang proses *coding* berlangsung. Memo ini berfungsi untuk melacak perkembangan interpretasi peneliti, mencatat konflik pemaknaan yang muncul, serta menstandarisasi istilah-istilah operasional teknis menuju abstraksi teoritis yang lebih mapan.

Langkah ketiga berfokus pada mitigasi bias interpretasi tunggal melalui *peer debriefing*. Mengingat analisis tematik reflektif sangat melibatkan subjektivitas peneliti, pengujian terhadap logika pengodean dan stabilitas struktur tema dilakukan bersama dosen pembimbing selaku *expert reviewer*. *Coding* beserta tema-tema yang sudah disintesis dalam fase ini ditinjau dan ditelaah oleh dosen pembimbing yang berfungsi sebagai audit internal untuk mengalibrasi arah analisis agar tetap berada pada batasan ruang lingkup riset *user persona*.

Langkah keempat adalah validasi makna dari sumber data utama. Mengingat tingginya keterbatasan waktu partisipan yang berstatus sebagai praktisi profesional aktif di industri, proses *member checking* formal pascaanalisis tidak memungkinkan untuk diselenggarakan. Sebagai antisipasi, teknik *real-time member checking* telah diterapkan saat sesi wawancara berlangsung melalui tindakan memparafrase dan mengonfirmasi ulang setiap wawasan krusial kepada partisipan. Secara metodologis disadari bahwa prosedur tersebut hanya berfungsi untuk memvalidasi akurasi pemahaman data mentah di awal wawancara dan tidak dapat mengonfirmasi ketepatan interpretasi tema final yang baru dikonstruksi secara induktif setelah keseluruhan proses analisis selesai. Oleh karena itu, sebagai mitigasi atas keterbatasan tersebut, validitas dan ketepatan interpretasi struktur tema final

dialihkan dan ditumpukan secara penuh pada ketajaman pengujian melalui mekanisme *peer debriefing* yang dilakukan secara intensif bersama dosen pembimbing selaku *expert reviewer*.

@Hak cipta milik IPB University

IPB University



Hak Cipta Dilindungi Undang-undang

1. Dilarang mengutip sebagian atau seluruh karya tulis ini tanpa mencantumkan dan menyebutkan sumber :
 - a. Pengutipan hanya untuk kepentingan pendidikan, penelitian, penulisan karya ilmiah, penyusunan laporan, penulisan kritik atau tinjauan suatu masalah
 - b. Pengutipan tidak merugikan kepentingan yang wajar IPB University.
2. Dilarang mengumumkan dan memperbanyak sebagian atau seluruh karya tulis ini dalam bentuk apapun tanpa izin IPB University.

IV HASIL DAN PEMBAHASAN

4.1 Gambaran Umum Pelaksanaan Penelitian

Secara keseluruhan, proses perekrutan melibatkan 13 praktisi UX yang dihubungi sebagai calon partisipan. Dari jumlah tersebut, delapan praktisi bersedia mengikuti wawancara dan menjadi partisipan penelitian, sedangkan lima lainnya tidak dapat berpartisipasi karena berbagai alasan. Proses pengumpulan data lapangan melalui wawancara mendalam berhasil dilaksanakan dan secara sadar dihentikan pada partisipan kedelapan (P8) setelah tercapainya titik saturasi data (*data saturation*). Berdasarkan parameter empiris yang dirumuskan oleh Hennink dan Kaiser (2022), saturasi teoretis terkonfirmasi ketika penambahan subjek baru tidak lagi menghasilkan kode, tema, atau kebaruan konseptual. Gejala saturasi ini mulai terdeteksi pada analisis partisipan keenam (P6), dan terwujud secara absolut ketika narasi mengenai prosedur anonimisasi data sensitif serta risiko distorsi artefak yang disampaikan oleh praktisi ketujuh (P7) menunjukkan pola yang sepenuhnya identik dengan temuan pada partisipan kedelapan (P8). Profil dan latar belakang kedelapan *User Experience Practitioner* yang berpartisipasi dalam penelitian ini disajikan pada Tabel 1.

Tabel 1 Profil delapan partisipan dalam penelitian

ID partisipan	Peran	Lama pengalaman di bidang UX (tahun)	Intensitas penggunaan GenAI	Spesifikasi tools GenAI	Skala perusahaan/sector industri
P1	<i>Product Designer</i>	±3	Kondisional (Saat fase <i>redesign</i>)	ChatGPT plus (<i>custom plugin</i> : UX Guru)	Skala menengah/Teknologi dan informasi
P2	<i>Lead Product Researcher</i>	±9	Rutin (Fase pra dan pascariset)	ChatGPT <i>enterprise</i> , Claude pro	Skala besar/Transportasi
P3	<i>UI/UX Designer</i>	±2	Kondisional (Instruksi manajemen)	ChatGPT go	Skala besar/ <i>Outsourcing</i>
P4	<i>Senior Product Designer</i>	±4	Rutin (Fase pencarian ide dan evaluasi)	ChatGPT plus (<i>custom plugin</i> : UX Guru)	Skala menengah/Teknologi dan informasi
P5	<i>Product Designer</i>	±4	Rutin (Sintesis data lapangan)	ChatGPT <i>free</i>	Skala menengah/Teknologi, informasi, dan komunikasi
P6	<i>Product Owner</i>	±6	Rutin (Saat klasterisasi <i>insight</i>)	ChatGPT <i>free</i>	Skala besar/Perbankan

Tabel 1 Profil delapan partisipan dalam penelitian (*lanjutan*)

ID partisipan	Peran	Lama pengalaman di bidang UX (tahun)	Intensitas penggunaan GenAI	Spesifikasi tools generative AI	Skala perusahaan/sektor industri
P7	UX Designer	±4	Mingguan (Tiap rilis fitur baru)	ChatGPT plus	Skala besar/ <i>Fintech</i>
P8	Product Designer	±4,5	Harian (Masuk alur kerja utama)	ChatGPT pro	Skala besar/ <i>Fintech</i>

Tabel 1 menyajikan variasi demografis yang komprehensif dari subjek penelitian, yang mencakup rentang pengalaman profesional di bidang *User Experience* (UX) antara dua hingga sembilan tahun, dengan distribusi peran mulai dari tingkat *junior*, *middle*, *senior*, hingga level *lead*. Selain itu, tabel ini juga menyoroti keragaman spesifikasi *tools Generative AI* yang digunakan oleh partisipan, yang terdistribusi mulai dari model versi gratis dengan batasan komputasi, hingga ekosistem berbayar dan penggunaan *custom plugin* khusus desain. Ragam latar belakang sektor industri tempat partisipan beroperasi, meliputi perusahaan teknologi menengah, korporasi transportasi skala besar, perbankan, hingga sektor *fintech*, memberikan kekayaan perspektif sosioteknis yang beragam.

Variasi tingkat keahlian, skala industri, serta perbedaan kapabilitas komputasional AI yang digunakan ini memperkuat validitas temuan melalui validasi silang. Praktisi senior dan *lead* (seperti P2 dan P6) memberikan wawasan strategis tingkat makro mengenai tata kelola risiko organisasi, ancaman otoritas algoritmik, dan implikasi bisnis dari penggunaan AI. Di sisi lain, praktisi di tingkat *junior* hingga *senior* (seperti P1, P3, P4, P5, P7, dan P8) memberikan pemahaman taktis tingkat mikro terkait operasional harian *prompt engineering*, variasi halusinasi teks, serta taktik pelepasan beban kognitif (*cognitive offloading*) dalam proses sintesis data. Meskipun terdapat perbedaan tingkat kematangan, konteks bisnis industri, maupun limitasi teknis dari masing-masing *tools*, analisis tematik menunjukkan adanya konvergensi pola yang konsisten, terutama dalam hal pemanfaatan GenAI sebagai asisten kognitif sekunder dan keharusan penerapan mekanisme *human-in-the-loop* guna menjaga validitas artefak riset UX.

Keseluruhan data transkrip yang terkumpul dari kedelapan partisipan tersebut kini tertata secara utuh dan siap untuk dieksplorasi serta dianalisis secara mendalam. Berdasarkan fondasi data empiris dan variasi wawasan lapangan tersebut, proses ekstraksi makna dilanjutkan untuk mengonstruksi temuan penelitian menuju tema final, sebagaimana diuraikan secara mendalam pada Subbab 4.2.

4.2 Penerapan *Reflexive Thematic Analysis* dan Pemetaan Tematis

Keseluruhan data mentah berupa transkrip *verbatim* dan catatan lapangan dari kedelapan partisipan telah dibedah dan dianalisis secara tematis menggunakan metode *Reflexive Thematic Analysis*. Proses ekstraksi makna ini dieksekusi secara sistematis untuk mengubah data kualitatif lapangan menjadi struktur wawasan yang rapi tanpa memaksakan kerangka teori eksternal.

Dalam praktiknya, potongan percakapan partisipan diekstrak dan diberikan label kode awal (*initial coding*) berdasarkan makna semantik yang eksplisit maupun makna laten yang tersirat secara psikologis. Guna mendemonstrasikan transparansi transformasi interpretasi dari kutipan mentah menjadi sebuah tema, cuplikan rekam jejak analitis (*audit trail*) disajikan pada Tabel 2. Adapun rincian matriks pengodean secara menyeluruh untuk semua partisipan direkam dan disajikan pada Lampiran 4.

Tabel 2 Cuplikan *audit trail* data transkrip menuju tema final

ID	Kutipan verbatim	Makna semantik/laten	Initial coding	Tema final
P1	"Timeline itu tuh selalu mepet... PM itu jarang mendefinisikan fase riset... Untuk menghemat waktu... ya kami perlu mempersingkat waktu."	Tekanan ritme kerja dari manajemen yang memaksa riset harus selesai sangat cepat.	<i>Time-constrained operational pressure</i>	Peran AI sebagai Asisten Kognitif dalam Alur Riset
P2	"Dia lumayan mengambil atau menggeneralisasikan konteks persona ke ekstrim kanan, yang kaya, kaya banget, yang Jaksel... nggak merepresentasikan segmentasi yang saya punya."	AI hanya mengambil data dominan masyarakat urban kelas atas yang cenderung elitis.	WEIRD ^a Bias	Keterbatasan Struktural pada Artefak Persona Berbasis AI
P8	"Apakah skill gue di area ini lama-lama jadi tumpul? karena tidak dilatih lagi... kemampuan gue buat sintesis manual sekarang mungkin tidak se-tajam dua tahun lalu."	Ketakutan akan hilangnya insting dan ketajaman riset akibat terlalu sering dibantu mesin.	<i>Deskilling risk</i>	Ilusi Kecerdasan dan Risiko <i>Over-Trust</i> terhadap AI

^aWEIRD: *Western, Educated, Industrialized, Rich, and Democratic*.

Puluhan kode awal yang telah terbentuk dan memiliki kedekatan makna konseptual tersebut kemudian disintesis menjadi enam tema final yang saling eksklusif. Guna membangun struktur laporan yang logis dan sejalan dengan kausalitas sosioteknis operasional di lapangan, keenam tema final ini diorganisasikan lebih lanjut ke dalam tiga dimensi makro, yakni: (1) Manfaat Operasional dan Interaksi Kognitif, (2) Keterbatasan Struktural dan Risiko Sosioteknis, serta (3) Tata Kelola Manusia dan Batasan Operasional. Rangkuman pemetaan hierarkis dari kode awal hingga menjadi tema final dan dimensi makro ditunjukkan pada Tabel 3.

Tabel 3 Pemetaan tema hasil analisis data kualitatif

Dimensi	Tema final (Fase 5)	Kode awal (Fase 2)
Dimensi 1: Manfaat Operasional dan Interaksi Kognitif	Peran AI sebagai Asisten Kognitif dalam Alur Riset	• <i>Time-constrained operational pressure</i> • <i>Workload scalability</i> • <i>Desk research acceleration</i> • <i>Data synthesis automation</i> • <i>Pre/post-research utilization</i> • <i>Cognitive offloading.</i>
	Strategi Prompting yang Kontekstual dan Adaptif	• <i>Iterative prompting</i> • <i>Role-playing directives</i> • <i>Front-loaded prompt</i> • <i>Reverse-prompting</i> • <i>Bias-awareness / Zero-shot prompt</i> • <i>Empirical grounding necessity</i> • <i>Context-aware prompting</i> • <i>Constraint-based prompting.</i>
Dimensi 2: Keterbatasan Struktural dan Risiko Sisioteknis	Keterbatasan Struktural pada Artefak Persona Berbasis AI	• WEIRD^a Bias • <i>Behavioral overgeneralization</i> • <i>Bias mayoritas / overshadowing data</i> • <i>Homogenisasi artefak persona</i> • <i>Expansive hallucination</i> • <i>Context blindness</i> • <i>Reductive hallucination</i> • <i>Plausible hallucination</i> • <i>Unrealistic solutions</i> • <i>Unrelevant constraint</i> • <i>Tone and wording inaccuracy</i>
	Ilusi Kecerdasan dan Risiko Over-Trust terhadap AI	• <i>Illusion of competence</i> • <i>Automation bias / over-trust</i> • <i>Rejection risk akibat over-reliance</i> • <i>Cognitive offloading</i> • <i>Deskilling risk</i> • <i>Replaceability anxiety</i> • <i>Over-reliance</i>
Dimensi 3: Tata Kelola Manusia dan Batasan Operasional	Batasan Peran AI dalam Aspek Empati dan Pengambilan Keputusan	• <i>Empathetic limitation</i> • <i>Human strategic autonomy</i> • <i>Contextual accountability</i> • <i>Decision-making limitation</i> • <i>Professional accountability</i> • <i>Anti zero-data generation</i> • <i>Authenticity of user evidence</i>
	Mekanisme Mitigasi Risiko dalam Penggunaan AI	• <i>Data anonymization</i> • <i>Selective data disclosure</i> • <i>Privacy over-trust / absence of protocol</i> • <i>Avoiding AI framing</i> • <i>Insight traceability</i> • <i>Human-in-the-loop intervention</i> • <i>Human-driven baseline</i>

^aWEIRD: Western, Educated, Industrialized, Rich, and Democratic.

Pengelompokan ketiga dimensi makro tersebut merepresentasikan lingkungan sosioteknis yang dinamis, ketika aspek psikologis dan operasional manusia (sosio) bersinggungan secara langsung dengan arsitektur dan keterbatasan algoritma (teknis). Adopsi *Generative AI* pada tahap awal sangat didorong oleh kebutuhan taktis untuk mengakselerasi efisiensi operasional dan melepaskan beban kognitif praktisi saat memproses data mentah (Dimensi 1). Namun, pada saat yang bersamaan, pemanfaatan tersebut secara inheren dibayangi oleh cacat struktural mesin seperti bias demografi dan halusinasi yang memicu kerentanan psikologis berupa ilusi kecerdasan pada diri praktisi (Dimensi 2). Menghadapi kompleksitas limitasi mesin tersebut, praktisi UX secara sadar dituntut untuk menegakkan batasan otoritas dan menerapkan mekanisme intervensi berlapis agar validitas keputusan desain tetap berada di tangan manusia (Dimensi 3). Keseluruhan tema dan dimensi yang beroperasi secara simultan ini kemudian disintesis menjadi luaran akhir penelitian yang utuh berupa *Conceptual Framework for AI-Assisted Persona Creation* pada akhir bagian pembahasan.

4.3 Dimensi 1: Manfaat Operasional dan Interaksi Kognitif

Dimensi ini mengeksplorasi motivasi fundamental dan taksonomi interaksi taktis praktisi UX dalam mengadopsi teknologi *Generative AI*. Analisis pada dimensi ini bukan hanya sekadar mendeskripsikan kegunaan teknologi, melainkan mendalami pemaknaan bahwa adopsi AI di industri UX bukan didorong oleh aspirasi otomatisasi penuh. Adopsi tersebut pada hakikatnya bertindak sebagai mekanisme pertahanan operasional bagi para praktisi dalam menavigasi keterbatasan sumber daya manusia dan waktu. Dimensi ini menaungi dua tema utama yang saling berkesinambungan, yaitu peran AI sebagai asisten kognitif dan penerapan strategi *prompting* yang adaptif.

4.3.1 Peran AI sebagai Asisten Kognitif dalam Alur Riset

Tema pertama ini mengungkap posisi fundamental *Generative AI* yang secara ketat difungsikan sebagai alat bantu kognitif, bukan sebagai “pengganti” peneliti. Konstruksi pemaknaan dari temuan ini menunjukkan bahwa pemicu utama penggunaan teknologi tersebut berakar pada kondisi tekanan struktural di dalam industri. Kondisi tersebut meliputi tingginya tuntutan ritme kerja harian serta ketidakseimbangan antara beban proyek dan ketersediaan jumlah anggota tim riset.

Praktisi dari berbagai sektor, mulai dari perusahaan teknologi, *fintech*, hingga konsultan, secara seragam mengindikasikan adanya defisit waktu riset yang ideal di lapangan. Sebagai contoh, narasi dari Partisipan 1 (P1) dan Partisipan 8 (P8) secara konsisten menyoroti fenomena struktural ketika manajer proyek kerap kali mengesampingkan alokasi waktu khusus untuk fase riset, sementara skala proyek terus membesar tanpa diiringi penambahan jumlah anggota tim. Kesenjangan antara tenggat waktu operasional yang sangat ketat dan kapasitas energi manusia ini memaksa periset untuk mencari jalan pintas guna mempersingkat waktu pemrosesan data. Kondisi beban kerja yang dirasa semakin berat tersebut direpresentasikan secara tajam melalui keluhan P8:

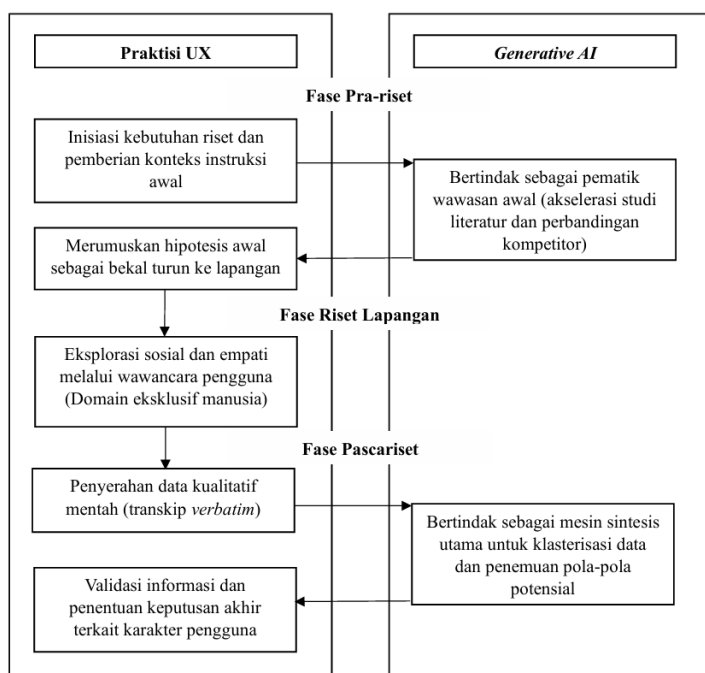
"skala proyeknya makin lama makin gede, dan tim risetnya nggak nambah... beban gue yang nanggung... ini nggak sustainable... gue mulai mikir

kalau gue terus-terusan pakai AI buat bantu sintesis... apakah skill gue di area ini lama-lama jadi tumpul?"

Tindakan praktisi yang mendelegasikan tugas berat seperti menyintesis puluhan halaman transkrip wawancara kepada mesin ini dilakukan secara strategis untuk menghemat energi mental agar dapat dialokasikan pada analisis tingkat lanjut yang menuntut kepekaan sosial, sebuah praktik operasional yang sejalan dengan fenomena *cognitive offloading* (Risiko dan Gilbert 2016).

Dalam alur kerja keseharian praktisi, delegasi kognitif ini terjadi secara terstruktur pada dua tahapan utama, yaitu fase sebelum riset lapangan dilakukan dan fase setelah pengumpulan data selesai. Pada tahap pra-riset, AI diposisikan sebagai pemantik wawasan awal untuk mengompensasi minimnya akses data praktisi. Partisipan 3 (P3) secara aktif menggunakan *generative AI* untuk mengakselerasi studi literatur dan perbandingan kompetitor guna mengatasi kendala birokrasi dan tenggat waktu sempit. Wawasan konseptual dan hipotesis awal yang dipantik oleh AI tersebut kemudian dibawa oleh praktisi sebagai bekal saat menggali informasi langsung dari pengguna di lapangan, sebelum akhirnya tumpukan data empiris yang diperoleh diserahkan kembali kepada AI pada fase pascariset.

Sementara itu, pada fase pascariset, *generative AI* bertindak sebagai mesin sintesis utama. Partisipan 2 (P2) dan Partisipan 5 (P5) mendelegasikan tugas klusterisasi data kualitatif mentah kepada AI untuk menemukan pola-pola permasalahan pengguna. Sinergi antara kecepatan komputasional mesin dalam memproses teks masif secara instan dan kebijaksanaan manusia dalam menilai validitas pola tersebut merupakan perwujudan nyata dari praktik *Hybrid Intelligence* di ruang lingkup desain (Dellermann *et al.* 2019). Visualisasi pembagian beban kognitif dan titik perpotongan tugas operasional antara praktisi dan mesin pada keseluruhan fase riset tersebut diilustrasikan secara lebih terperinci pada Gambar 2.



Gambar 2 Visualisasi alur kolaborasi *Hybrid Intelligence* antara praktisi UX dan *Generative AI*

Praktisi dalam hal ini memandang AI sebagai instrumen untuk mengolah kompleksitas data mentah menjadi sebuah informasi yang lebih terstruktur dan sederhana. Meskipun AI mampu merangkum data dengan cepat, keputusan akhir mengenai validitas informasi yang dihasilkan oleh AI tetap berada pada diri praktisi. P2 menyampaikan bahwa GenAI dapat berperan sebagai "rekan berpikir" yang membantu memperluas perspektif dan memperkaya opsi riset, namun tidak dapat dijadikan sebagai pengambil keputusan akhir dalam pendefinisian karakter pengguna.

4.3.2 Strategi *Prompting* yang Kontekstual dan Adaptif

Tema kedua berhasil diamati dari sisi interaksi teknis antara praktisi dan *Generative AI*. Temuan penelitian menunjukkan bahwa tidak ada satu standar baku dalam menyusun *prompt* atau instruksi pada AI. Praktisi memodifikasi instruksi mereka secara terus-menerus berdasarkan jenis data mentah yang dimiliki serta tingkat mitigasi risiko halusinasi yang ingin dihindari. Dari penjabaran pengalaman para partisipan, penelitian ini mengonstruksi strategi pemberian instruksi ke dalam beberapa paradigma operasional yang berbeda.

Pendekatan pertama berfokus pada penciptaan batasan kontekstual sejak detik pertama interaksi melalui penerapan *front-loaded prompt*. Partisipan 5 (P5) menerapkan strategi pemberian instruksi berskala masif di awal. P5 memaparkan taktiknya secara lugas: "*pertama itu biasanya aku ngasihnya jebret sekali banyak... dari info yang ada itu kita kasihin semua*". Tujuan dari strategi ini adalah untuk mengunci algoritma AI agar tidak melenceng dari topik, karena menurut pengalaman P5, "*kalau cuma dapet sepotong teksnya... kadang dia entah kemana-mana gitu jawabnya. Ada improvisasinya sendiri*".

Paradigma masif di awal ini bertolak belakang secara operasional, namun memiliki kesamaan tujuan analitis dengan strategi *iterative prompting* yang diterapkan oleh Partisipan 4 (P4). P4 lebih memilih memberikan instruksi-instruksi yang berbentuk pertanyaan-pertanyaan kecil secara bertahap. Taktik pancingan ini digunakan untuk memastikan bahwa mesin telah menyelaraskan pemahamannya terhadap konteks produk secara benar, sebelum beban tugas analitis yang sesungguhnya diberikan kepada sistem.

Temuan berikutnya dalam penelitian ini menunjukkan bahwa mayoritas praktisi sangat menentang pembuatan persona yang dibangun dari ruang hampa, yaitu meminta AI mendeskripsikan pengguna tanpa memberikan landasan data wawancara sama sekali. Penolakan ini merupakan wujud upaya sadar praktisi untuk mempertahankan validitas empiris dalam perancangan produk (Nielsen 1994). Partisipan 3 (P3) menegaskan bahwa instruksi yang dipandu oleh data mentah akan menghasilkan identifikasi kebutuhan pengguna dan titik permasalahan yang jauh lebih akurat, alih-alih persona klise yang tidak realistis (Paoli 2023). Praktisi menegaskan bahwa validitas artefak AI sangat bergantung pada kemampuan peneliti dalam menyuntikkan data kualitatif ke dalam sistem guna memastikan persona yang terbentuk merupakan representasi akurat dari perilaku pengguna asli.

Meskipun demikian, penelitian ini juga mengidentifikasi sebuah kasus anomali yang sangat kritis. Berkebalikan dengan partisipan lainnya, Partisipan 1 (P1) justru sengaja memulai interaksinya dengan instruksi kosong, tanpa memberikan konteks data di awal. Anomali ini dilakukan oleh P1 karena P1

memanfaatkan AI sebagai mekanisme psikologis untuk menguji bias konfirmasi yang mungkin menjebak pikiran peneliti itu sendiri. P1 menjabarkan alasannya secara transparan:

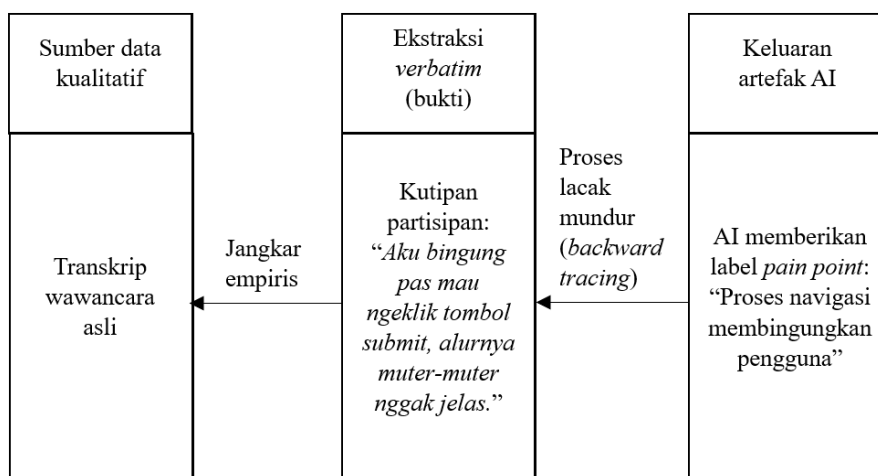
"takutnya pandanganku itu tuh ngebuat jawaban dia jadi biar seolah-olah sama kayak kesukaanku... padahal kan bisa jadi ChatGPT punya pandangan lain... aku butuh user persona versi dia dulu".

Perilaku anomali yang ditunjukkan oleh P1 ini memberikan interpretasi konseptual yang sangat esensial. Pendekatan seperti ini sejalan dengan gagasan bahwa paparan sudut pandang yang kontras atau *bias-awareness features* dapat memicu refleksi kritis (Zavolokina *et al.* 2025). Interaksi dengan *Generative AI* dalam ranah riset UX tidak melulu berpusat pada pemberian perintah kepada mesin untuk memvalidasi hipotesis manusia. Lebih jauh dari itu, praktisi juga memanfaatkan mesin sebagai sebuah cermin netral untuk mengkritisi asumsi subjektif mereka sendiri, sebelum akhirnya menyelaraskan temuan tersebut dengan data riil dari lapangan. Secara teoretis, taktik yang dieksekusi oleh P1 ini menandai sebuah pergeseran paradigma dalam kolaborasi *human-AI*, yakni dari sekadar pemanfaatan AI sebagai mesin penjawab otomatis, menjadi sebuah instrumen refleksi kognitif yang mematangkan ketajaman periset.

Sebagai penutup pada dimensi interaksi ini, strategi paling krusial untuk menahan laju halusinasi data adalah penerapan instruksi pembuktian secara eksplisit. Partisipan 6 (P6) selalu menyisipkan perintah wajib kepada sistem AI untuk melampirkan kutipan asli dari perkataan responden di samping setiap label titik masalah yang dihasilkan. P6 menjabarkan instruksi spesifik tersebut secara transparan:

"aku butuh grouping pain points, tapi bersamaan dengan itu... ada kebutuhan untuk verbatim atau kutipan dari jawaban responden yang menunjukkan pain point tersebut itu yang mana, terus dari responden yang mana".

Guna mengilustrasikan mekanisme pertahanan operasional ini secara konseptual, Gambar 3 mendemonstrasikan proses *backward tracing*. Dengan menggunakan teknik ini, setiap label atau titik masalah fiktif yang dihasilkan oleh mesin dapat dilacak mundur menuju jangkar empirisnya pada transkrip wawancara asli, sehingga validitas data tetap terjamin.



Gambar 3 Visualisasi *backward tracing* untuk menegaskan prinsip *insight traceability* pada luaran AI

Sejalan dengan hal itu, Partisipan 3 (P3) memaksa mesin untuk selalu menyertakan rujukan artikel yang valid setiap kali membuat sebuah kesimpulan wawasan tambahan. Hal ini dilakukan agar hasil sintesis tidak didasarkan pada asumsi algoritmik semata. P3 memberikan rasionalisasi operasional atas taktik pertahanan ini:

"tolong sertakan artikel atau jurnal atau apalah itu. Jadi, biar mereka itu nggak sembarang masukin saja, tapi beneran ada datanya".

Kendati demikian, dari sudut pandang arsitektur komputasi LLM, terdapat perbedaan tingkat keandalan yang sangat kontras antara taktik yang diterapkan oleh P6 dan P3. Taktik ekstraksi *verbatim* yang digunakan P6 merupakan mekanisme pertahanan yang sangat baik secara empiris, karena ia bekerja pada basis data internal yang secara faktual diberikan langsung oleh peneliti. Sebaliknya, instruksi P3 yang memaksa *generative AI* untuk mencari dan merujuk artikel atau jurnal eksternal tanpa menyuplainya dengan basis data khusus justru merupakan sebuah bentuk pemanfaatan GenAI yang sangat berbahaya. Secara arsitektural, tindakan ini tidak memperkuat kebenaran, melainkan memicu sistem untuk melakukan fabrikasi data atau *plausible hallucination*, yakni memproduksi judul literatur, penulis, dan tahun fiktif yang tampak sangat meyakinkan secara linguistik namun tidak pernah eksis di dunia nyata.

Oleh karena itu, penerapan prinsip *insight traceability* harus ditopang secara arsitektural oleh jangkar data internal yang terverifikasi. Praktik pemberian instruksi pembuktian yang berbasis pada transkrip nyata bertindak sebagai mekanisme pertahanan aktif untuk mencegah fenomena bias otomatisasi berlebihan, sekaligus menjaga kedaulatan penalaran ilmiah dari jeratan fabrikasi algoritmik.

Berdasarkan pemetaan tematis pada fase pengodean awal yang ada pada Tabel 3, teridentifikasi delapan variasi kode taktik operasional yang diterapkan oleh praktisi. Namun, guna menyusun sebuah taksonomi struktural yang komprehensif dan menghindari tumpang tindih konseptual, kedelapan taktik turunan tersebut disintesis dan diklasifikasikan ke dalam empat paradigma strategi utama berdasarkan kesamaan tujuan analitisnya. Pertama, taktik *context-aware prompting* dilebur ke dalam paradigma *front-loaded prompt*, karena keduanya berfokus pada penyediaan informasi latar belakang proyek secara masif di awal interaksi. Kedua, taktik *reverse-prompting* diklasifikasikan ke dalam strategi *iterative prompting*, mengingat taktik memancing pertanyaan balik dari AI tersebut merupakan bagian dari siklus penyelarasan pemahaman secara bertahap. Ketiga, taktik *bias-awareness* berafiliasi langsung dengan paradigma *zero-shot prompt* sebagai metode psikologis untuk menguji asumsi subjektif praktisi. Keempat, instruksi *role-playing directives* dan *empirical grounding necessity* dikategorikan ke dalam klasifikasi *constraint-based prompting*. Penggabungan ini didasarkan pada argumen bahwa pemberian peran pakar maupun kewajiban melampirkan kutipan *verbatim* pada hakikatnya merupakan bentuk penegakan batasan terhadap kebebasan algoritma mesin. Melalui proses sintesis tersebut, Tabel 4 menyajikan taksonomi final dari strategi *prompting* yang mencakup jenis strategi utama, tujuan analitis, serta ringkasan contoh penerapannya oleh para praktisi di lapangan.

Hak Cipta Dilindungi Undang-undang
1. Dilarang mengutip sebagian atau seluruh karya tulis ini tanpa mencantumkan dan menyebutkan sumber :
a. Pengutipan hanya untuk kepentingan pendidikan, penelitian, penulisan karya ilmiah, penyusunan laporan, penulisan kritik atau tinjauan suatu masalah
b. Pengutipan tidak merugikan kepentingan yang wajar IPB University.
2. Dilarang mengumunkan dan memperbanyak sebagian atau seluruh karya tulis ini dalam bentuk apapun tanpa izin IPB University.

Tabel 4 Taksonomi strategi *prompting* praktisi UX

Jenis strategi	Tujuan analitis	Contoh penerapan
<i>Front-loaded prompt</i>	Mengunci konteks sejak awal agar algoritma tidak melenceng atau melakukan improvisasi di luar ruang lingkup riset.	P5: Memberikan seluruh informasi dan konteks data secara masif pada instruksi pertama.
<i>Iterative prompting</i>	Menyelaraskan pemahaman mesin terhadap konteks produk secara bertahap sebelum memberikan beban tugas analitis berat.	P4: Memberikan instruksi berupa pertanyaan-pertanyaan kecil secara bertahap sebagai taktik pancingan.
<i>Zero-shot prompt</i>	Menjadikan mesin sebagai "cermin netral" untuk menguji bias konfirmasi subjektif yang mungkin dimiliki oleh periset.	P1: Memulai interaksi dengan instruksi kosong tanpa data awal untuk mengeksplorasi sudut pandang murni model.
<i>Constraint-based prompting</i>	Menahan laju fabrikasi data dan menegaskan prinsip <i>insight traceability</i> melalui jangkar pembuktian empiris.	P6: Memerintahkan AI secara wajib untuk melampirkan kutipan asli (<i>verbatim</i>) di samping setiap label masalah.

4.4 Dimensi 2: Keterbatasan Struktural dan Risiko Socioteknis

Dimensi kedua ini menginvestigasi sisi gelap dari interaksi antara praktisi UX dan teknologi *Generative AI*. Analisis pada dimensi ini tidak lagi berfokus pada kemudahan operasional dan efisiensi waktu, melainkan membongkar cacat struktural bawaan pada *Large Language Models* (LLM) serta dampak psikologis jangka panjang yang mengancam integritas profesional periset. Dimensi ini menaungi dua tema krusial, yaitu keterbatasan teknis artefak persona yang dihasilkan AI dan fenomena ilusi kecerdasan yang memicu pelepasan beban kognitif secara berlebihan.

4.4.1 Keterbatasan Struktural pada Artefak Persona Berbasis AI

Meskipun teknologi ini terbukti mampu mengakselerasi proses sintesis data, temuan empiris di lapangan menunjukkan bahwa artefak persona yang dihasilkan oleh AI mengandung kecacatan struktural bawaan. Keterbatasan ini dikategorikan ke dalam tiga bentuk anomali utama, yaitu bias demografi, generalisasi perilaku secara berlebihan, dan halusinasi data. Temuan ini secara langsung mengonfirmasi studi Yang *et al.* (2023) yang memperingatkan bahwa penerapan LLM dalam skenario dunia nyata masih dibayangi oleh masalah reliabilitas dan bias data turunan.

Anomali pertama adalah bias demografi dan geografis yang secara teoretis dipicu oleh ketimpangan distribusi data pelatihan global pada arsitektur LLM. Partisipan 2 (P2) yang merupakan periset senior di sektor transportasi berskala besar, secara gamblang mengidentifikasi fenomena *WEIRD bias* (*Western, Educated, Industrialized, Rich, and Democratic*) pada keluaran model AI. Menurut temuan P2, algoritma sering kali gagal memetakan realitas kelas pekerja atau masyarakat menengah ke bawah di Indonesia, dan justru

menciptakan persona yang bersifat sangat elitis. P2 membedah fenomena ini dengan bukti lapangan yang kuat:

"Dia lumayan mengambil atau menggeneralisasikan konteks persona ke ekstrim kanan, yang kaya, kaya banget, yang Jaksel... nggak merepresentasikan segmentasi yang saya punya."

Kondisi tersebut menyebabkan AI gagal menggambarkan realitas segmentasi *grassroots* secara akurat. Temuan ini didukung oleh Partisipan 4 (P4) yang menyatakan bahwa profil persona buatan AI sering kali terlalu meluas dan kurang merakyat, sehingga artefak tersebut menjadi kehilangan relevansinya apabila digunakan untuk merancang produk bagi demografi kelas sosial menengah ke bawah. Ketidakmampuan algoritma dalam menangkap realitas segmentasi lokal di Indonesia ini terkonfirmasi secara teoretis oleh Salminen *et al.* (2024), yang menemukan bahwa persona hasil luaran LLM sangat *US-centric*, yakni didominasi oleh atribut Amerika Serikat meskipun indikator negara tidak disebutkan di dalam instruksi awal.

Anomali kedua adalah kecenderungan mesin untuk melakukan generalisasi berlebihan terhadap perilaku pengguna dan homogenisasi artefak. Algoritma AI sering kali mereduksi keberagaman motivasi manusia ke dalam satu narasi tunggal yang klise. Partisipan 7 (P7), seorang *UX Designer* di perusahaan *fintech* menemukan bahwa AI memukul rata preferensi seluruh target penggunanya tanpa memedulikan konteks varian data. P7 menjabarkan fenomena reduksi tersebut secara mendetail:

"sebagian besar user yang aku interview itu emang punya concern soal bunga... nah pas aku masukin ke AI, dia langsung kayak nge-generate persona yang seolah-olah 'semua user sangat sensitif terhadap bunga'. Padahal kalau dilihat lebih dalam... ada juga user yang sebenarnya lebih peduli ke kecepatan cairnya dana. Jadi di situ AI kayak terlalu menyamaratakan."

Selain itu, Partisipan 8 (P8) juga menyoroti bahaya laten dari homogenisasi artefak ini secara makro. P8 memperingatkan apabila mayoritas praktisi UX di industri menggunakan model AI yang sama dengan logika instruksi yang serupa, persona yang dihasilkan berpotensi menjadi seragam secara algoritmik. P8 memberikan kritik tajam terhadap fenomena ini: *"ada kemungkinan persona yang dihasilin itu jadi seragam gitu. Kurang mencerminkan keunikan user di konteks spesifik masing-masing"*. Kekhawatiran empiris ini menunjukkan bahwa jika banyak praktisi memakai model dan logika *prompt* yang serupa, artefak yang dihasilkan berisiko terdorong ke pola representasi yang repetitif. Fenomena ini sejalan dengan temuan Truss (2026) mengenai *flatten identity groups* pada LLM, yaitu ketika algoritma cenderung meratakan keunikan konteks lokal demi menghasilkan keluaran yang persuasif berdasarkan bias data latih utamanya.

Anomali ketiga, yang juga merupakan permasalahan yang cukup berbahaya bagi validitas riset adalah fenomena halusinasi data. Berdasarkan temuan di lapangan, halusinasi ini tidak terjadi secara tunggal, melainkan terbagi ke dalam beberapa skenario operasional. Skenario pertama adalah halusinasi ekspansif, yaitu ketika mesin memfabrikasi asumsi fiktif yang tidak pernah ada di dalam transkrip wawancara. Partisipan 6 (P6) menemukan kecenderungan ini saat menginstruksikan AI untuk membuat *need statement*. P6 mengeluhkan: *"dia*

malah menambahkan asumsi-asumsi sendiri hasilnya ya, outputnya waktu itu" meskipun data yang diberikan sangat terbatas.

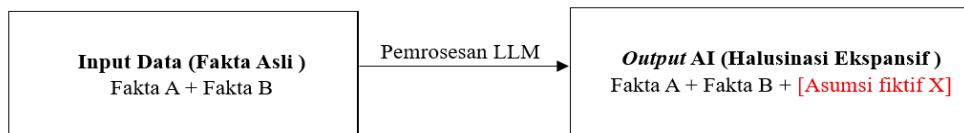
Skenario kedua yang ditemukan merupakan kebalikan dari sifat ekspansif, yakni halusinasi reduktif. Pada anomali ini, alih-alih menambahkan data fiktif, mesin justru menghilangkan atau membuang kriteria krusial dari pengguna saat proses peringkasan. Partisipan 5 (P5) mengeluhkan kecenderungan algoritma ini: *"kadang dia menghilangkan beberapa unsur penting dari infonya gitu... mungkin karena dia ngeringkasin, kadang ada yang hilang"*. Reduksi informasi oleh algoritma ini berbahaya karena berpotensi melenyapkan *nuance* atau konteks spesifik dari masalah pengguna yang sesungguhnya vital untuk perancangan fitur.

Skenario ketiga, yang sekaligus merupakan varian yang paling berbahaya adalah halusinasi halus (*plausible hallucination*). Pada anomali ini, AI menyajikan data fiktif yang dibungkus dengan struktur bahasa yang sangat rapi dan terdengar sangat logis di telinga periset. Partisipan 7 (P7) memberikan contoh telak ketika AI secara mandiri merangkai titik permasalahan (*pain points*) fiktif pada desain aplikasi *fintech*:

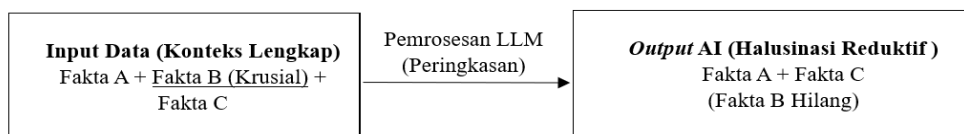
"dia nambahin detail yang kelihatannya masuk akal, tapi sebenarnya nggak ada di data. Misalnya dia nulis 'user merasa tidak percaya dengan institusi keuangan formal'. Padahal dari interview, nggak ada yang ngomong ke arah situ secara eksplisit gitu. Jadi itu kayak asumsi yang dibangun sendiri sama AI."

Guna mempertegas perbedaan karakteristik operasional dari ketiga kegagalan algoritma tersebut, Gambar 4 mengilustrasikan mekanisme transformasi input data menjadi output artefak pada masing-masing varian halusinasi.

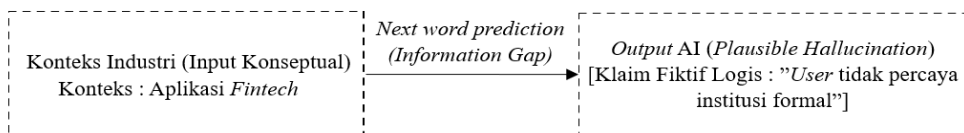
HALUSINASI EKSPANSIF (Injeksi Data Fiktif)



HALUSINASI REDUKTIF (Reduktif Konteks)



HALUSINASI HALUS/PLAUSIBLE (Fabrikasi Klaim Logis)



Gambar 4 Visualisasi mekanisme transformasi data pada tiga varian halusinasi algoritma GenAI

Ketiga bentuk halusinasi tersebut dapat dijelaskan secara konseptual melalui keterbatasan arsitektur LLM itu sendiri. Model kecerdasan buatan beroperasi berdasarkan perhitungan probabilitas statistik untuk memprediksi kemunculan kata berikutnya (*next-word prediction*), bukan melalui pemahaman semantik layaknya akal manusia (Yang *et al.* 2023). Fenomena ini membuktikan

bahwa ketika mesin dihadapkan pada *information gap*, algoritma memiliki tendensi laten untuk mengisi kekosongan tersebut menggunakan pola kalimat dari data pelatihannya yang terlihat koheren secara sintaksis, meskipun sebenarnya cacat secara empiris.

Guna menyintesis berbagai anomali struktural yang telah dijabarkan, pemetaan klasifikasi kegagalan algoritma tersebut dirangkum pada Tabel 5. Tabel ini membedah secara spesifik korelasi antara jenis anomali, mekanisme kegagalan secara arsitektural, dan manifestasinya di lapangan.

Tabel 5 Klasifikasi anomali struktural pada artefak persona buatan AI

Jenis anomali	Mekanisme kegagalan AI	Contoh kasus partisipan
Bias Demografi dan Geografis (<i>WEIRD Bias</i>)	Ketimpangan distribusi data pelatihan global yang gagal memetakan realitas kelas pekerja (<i>grassroots</i>).	P2: AI menghasilkan persona elitis yang mengabaikan segmentasi masyarakat menengah ke bawah.
Generalisasi Perilaku Secara Berlebihan dan Homogenisasi	Reduksi keberagaman motivasi pengguna ke dalam satu narasi tunggal yang klise dan seragam secara algoritmik.	P7: AI memukul rata bahwa seluruh pengguna aplikasi <i>fintech</i> sensitif terhadap suku bunga.
Halusinasi Data	Prediksi probabilitas statistik (<i>next-word prediction</i>) yang mengisi celah informasi (<i>information gap</i>) tanpa pemahaman semantik nyata.	• Halusinasi ekspansif: Penambahan asumsi fiktif pada <i>need statement</i> dari P6. • Halusinasi reduktif: Lenyapnya unsur penting pengguna saat AI melakukan peringkasan dari P5. • Halusinasi halus (<i>plausible</i>): Perangkaian <i>pain points</i> fiktif yang terdengar logis terkait institusi keuangan dari P7.

4.4.2 Ilusi Kecerdasan dan Risiko *Over-Trust* terhadap AI

Tema keempat menyingkap dampak sosioteknis yang lebih dalam yang melampaui sekadar kesalahan teks pada artefak persona. Penting untuk dideklarasikan bahwa dari kedelapan partisipan dalam penelitian ini, tidak ditemukan satupun kasus ekstrem terkait kepercayaan absolut praktisi terhadap AI, pandangan bahwa halusinasi data bukan masalah, ataupun anggapan bahwa AI mampu melakukan *persona generation* secara mandiri tanpa intervensi manusia. Namun demikian, analisis pada tema ini menunjukkan bahwa tingginya kefasihan linguistik yang dihasilkan oleh AI sering kali menciptakan sebuah ilusi kecerdasan (*illusion of competence*). Kefasihan sintaksis ini mampu mengecoh praktisi maupun pemangku kepentingan untuk mempercayai kualitas substansi keluaran AI secara absolut, tanpa menyadari bahwa wawasan tersebut dangkal atau bahkan cacat secara empiris.

Risiko utama yang dapat terjadi dari ilusi ini adalah ancaman ketergantungan yang berlebihan (*over-reliance*) terhadap teknologi AI. Partisipan 7 (P7) dan Partisipan 8 (P8) mengamati bahwa luaran AI selalu terlihat sangat rapi, terstruktur, dan meyakinkan, sehingga tampak seperti hasil pemikiran analitis yang matang. Kefasihan semacam ini dinilai berbahaya karena dapat mengecoh pihak manajemen yang tidak memiliki latar belakang riset untuk langsung menelan hasil tersebut sebagai kebenaran faktual. Partisipan 2 (P2) menceritakan pengalaman empirisnya ketika ia hampir terkecoh oleh laporan strategi produk buatan staf magangnya yang disusun sepenuhnya menggunakan AI. Laporan tersebut terlihat sangat komprehensif dan meyakinkan. Namun, saat ditelaah lebih dalam oleh periset senior, strategi tersebut terbukti hanyalah retorika dangkal tanpa pemahaman konteks bisnis yang nyata. P2 memberikan peringatan keras mengenai bahaya ilusi kefasihan ini:

"Kita bisa dimabukkan dengan segala ke-fancy-an AI, padahal konteksnya itu ngasal aja si AI. Nggak tahu dia ngambilnya dari mana, tapi dia bisa meyakinkan kita."

Meskipun pengalaman tersebut konteksnya terkait laporan strategi bisnis, bahaya ilusi kefasihan ini terbukti sama destruktifnya jika dibiarkan menyusup ke dalam proses penyusunan karakter dan motivasi pada *user persona*. Fenomena pengecohkan oleh kefasihan sistem ini merupakan bentuk nyata dari jebakan *algorithmic authority*. Dalam kondisi ini, manusia secara tidak sadar memberikan legitimasi kebenaran mutlak kepada keluaran algoritma semata-mata karena struktur bahasanya yang rapi, sehingga menyingkirkan proses peninjauan empiris yang seharusnya menjadi fondasi riset.

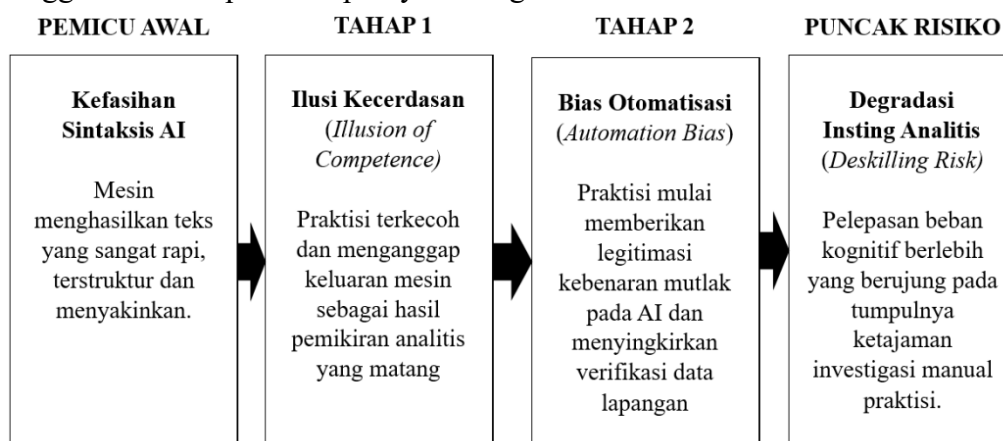
Dampak lebih lanjut dari ilusi kecerdasan ini adalah munculnya *automation bias*, yaitu suatu kondisi ketika praktisi secara perlahan menggeser otoritas kebenaran risetnya kepada mesin dan mengabaikan insting analisis manusianya (Parasuraman dan Manzey 2010). Partisipan 5 (P5) mengingatkan bahwa bias otomatisasi ini dapat berujung pada risiko penolakan produk secara masif di pasar (*rejection risk*). Apabila seorang periset mengandalkan solusi *persona* buatan AI tanpa melakukan verifikasi ke lapangan, produk akhir yang dirancang dipastikan akan meleset jauh dari kebutuhan riil pengguna, sehingga berujung pada kerugian bisnis perusahaan.

Lebih jauh lagi, ketergantungan seperti ini memicu ancaman psikologis laten berupa degradasi keterampilan analitis (*deskilling risk*) pada diri praktisi itu sendiri. Pelepasan beban kognitif (*cognitive offloading*) yang berlebihan ke dalam sistem eksternal lambat laun menyebabkan insting investigasi periset menjadi tumpul (Risiko dan Gilbert 2016). Kekhawatiran profesional ini diungkapkan secara mendalam dan sangat reflektif oleh Partisipan 8 (P8), yang mulai merasakan dampak psikologis akibat terlalu sering mendelegasikan tugas sintesisnya kepada AI:

"Gue mulai mikir, kalau gue terus-terusan pakai AI buat bantu sintesis... apakah skill gue di area ini lama-lama jadi tumpul? Karena otot itu nggak dilatih lagi seperti dulu. Dan ini gue rasain sendiri... kemampuan gue buat sintesis manual sekarang mungkin nggak se-tajam dua tahun lalu."

Ketakutan empiris yang disuarakan oleh P8 tersebut bukanlah sebuah kekhawatiran tunggal oleh dirinya sendiri, melainkan sebuah fenomena

struktural yang terkonfirmasi oleh literatur mutakhir pada disiplin ilmu interaksi *human-AI*. Pernyataan reflektif P8 ini mengonstruksi makna filosofis yang sangat tajam bagi diskursus keilmuan perancangan UX. Kehadiran *Generative AI* ternyata tidak hanya mengubah metode teknis produksi *user persona*, tetapi secara fundamental berpotensi mendegradasi arsitektur kognitif generasi periset di masa depan. Gambar 5 mendemonstrasikan eskalasi efek domino dari risiko kognitif ini secara visual, dimulai dari kefasihan sintaksis mesin yang menipu, hingga bermuara pada tumpulnya insting analitis manusia.



Gambar 5 Visualisasi eskalasi risiko kognitif pada interaksi *human-AI*

Ancaman struktural yang divisualisasikan pada alur eskalasi tersebut secara empiris divalidasi oleh studi literatur terdahulu. Li *et al.* (2024) dalam studinya menemukan bahwa *UX designer* berpengalaman melihat GenAI berguna untuk tugas repetitif, tetapi juga khawatir pada *skill degradation*, terutama pada desainer junior yang bisa kehilangan kesempatan mempelajari proses desain secara sistematis dan malah “terlatih” sebagai *prompter*.

Guna merangkum dan memetakan eskalasi konsekuensi dari interaksi sosioteknis tersebut secara terstruktur, Tabel 6 mengklasifikasikan ragam ancaman ini ke dalam dua level dampak yang berbeda. Pemetaan ini memisahkan antara dampak eksternal yang mengancam integritas produk dan bisnis, serta dampak internal yang mengikis kompetensi fundamental praktisi.

Tabel 6 Klasifikasi risiko sosioteknis pada penggunaan AI dalam riset UX

Kategori dampak	Manifestasi risiko sosioteknis
Dampak Eksternal (Produk/Bisnis)	Risiko Penolakan Pasar (<i>Rejection Risk</i>), yaitu kerugian bisnis yang terjadi akibat produk dirancang berdasarkan artefak persona yang meleset dari kebutuhan riil pengguna karena ketiadaan verifikasi lapangan.
Dampak Internal (Kompetensi Praktisi)	- Bias Otomatisasi (<i>Automation Bias</i>), yaitu hilangnya daya kritis ketika periset menyerahkan otoritas kebenaran riset dan melegitimasi keluaran algoritma secara mutlak. - Degradasi Keterampilan (<i>Deskilling Risk</i>), yaitu menumpulnya insting investigasi dan menurunnya ketajaman sintesis manual akibat pelepasan beban kognitif (<i>cognitive offloading</i>) yang berlebihan ke sistem kecerdasan buatan.

4.5 Dimensi 3: Tata Kelola Manusia dan Batasan Operasional

Sebagai respons mutlak untuk memitigasi eskalasi risiko sosioteknis yang telah dijabarkan pada dimensi sebelumnya, dimensi ketiga ini menganalisis batasan fundamental yang membatasi pendelegasian tugas kepada *Generative AI* serta mekanisme pertahanan yang diterapkan oleh praktisi guna menjaga integritas riset. Analisis dalam dimensi ini menegaskan alasan fundamental mengapa aspek kemanusiaan dalam riset UX tidak dapat diotomasi. Dimensi ini menaungi dua tema utama, yaitu batasan peran AI dalam aspek empati dan pengambilan keputusan, serta mekanisme mitigasi risiko yang diterapkan oleh praktisi.

4.5.1 Batasan Peran AI dalam Aspek Empati dan Pengambilan Keputusan

Tema kelima membahas mengenai batasan-batasan operasional yang tidak boleh dilampaui oleh *Generative AI* dalam proses perancangan *user persona*. Meskipun kita ketahui bahwa AI memiliki kemampuan analitis dan sintesis data yang sangat cepat, para partisipan secara konsisten menekankan bahwa terdapat area fundamental yang harus tetap dikendalikan sepenuhnya oleh akal dan pikiran manusia. Batasan ini terutama berkaitan dengan aspek empati, interpretasi emosi pengguna, dan pengambilan keputusan kritis yang membutuhkan tanggung jawab moral.

Aspek pertama yang menjadi batasan mutlak adalah empati dan interaksi antarmanusia. Partisipan 6 (P6) menekankan bahwa AI tidak memiliki kapasitas untuk merasakan emosi atau memahami kedalaman perasaan pengguna saat sesi wawancara berlangsung. Kesimpulan mengenai apa yang benar-benar dirasakan oleh pengguna hanya dapat ditarik oleh peneliti manusia yang berinteraksi langsung. Pandangan ini sejalan dengan Partisipan 4 (P4) yang menegaskan bahwa elemen yang berhubungan langsung dengan interaksi manusia tidak dapat digantikan oleh mesin. Batasan ini sangat logis mengingat model bahasa besar hanya beroperasi pada pemrosesan sintaksis linguistik tanpa memiliki pengalaman interaksi sosial di dunia nyata (Yildirim dan Paul 2024). Oleh karena itu, fungsionalitas AI murni hanya sebatas sebagai instrumen pemroses data teks, bukan untuk mengambil alih kepekaan sosial seorang periset.

Batasan kedua terdapat pada tahapan pengambilan keputusan kritis dan perumusan masalah. Partisipan 4 (P4) berprinsip untuk tidak mengandalkan AI dalam hal yang berkaitan dengan regulasi perusahaan maupun pendefinisian problem utama tanpa konteks yang jelas. Pendefinisian masalah harus diambil alih sepenuhnya oleh peneliti manusia. P4 mengingatkan bahwa amanah utama seorang praktisi UX adalah merancang produk yang tepat sasaran untuk menyelesaikan masalah manusia, bukan sekadar memindahkan tanggung jawab kepada algoritma:

"harus tau dulu batas yang mana yang bisa di handle AI... kita diamanahin untuk bikin produk yang bener, bukan diamanahin untuk nge-prompting AI".

Pergeseran peran praktisi menjadi sekadar "operator prompt" sesungguhnya merupakan gejala awal dari hilangnya daya kritis manusia akibat penyerahan keputusan pada otoritas mesin. Pernyataan P4 ini menjadi bentuk kesadaran praktisi untuk memitigasi bahaya otoritas algoritmik (*algorithmic authority*), guna menjaga otonomi strategis manusia dalam pendefinisian masalah (Dwivedi *et al.* 2023).

Aspek terakhir yang menjadi batas fundamental adalah orisinalitas dan validitas empiris dalam penyajian bukti lapangan. Partisipan 2 (P2) secara proaktif menerapkan prinsip *anti zero-data generation*, yaitu penolakan keras untuk meminta AI menghasilkan persona dari ruang hampa tanpa adanya landasan data empiris yang disuplai oleh periset. Sejalan dengan hal tersebut, Partisipan 8 (P8) menggarisbawahi bahwa setiap kutipan yang dicantumkan dalam artefak persona wajib merupakan pernyataan *verbatim* langsung dari pengguna, bukan hasil fabrikasi dari bahasa mesin. Tindakan P2 dan P8 ini merupakan manifestasi langsung dari upaya untuk menghindarkan persona dari asumsi fiktif yang menyesatkan, sejalan dengan prinsip *User-Centered Design* (Cooper *et al.* 2014). Apabila periset membiarkan AI mengarang *pain points* atau kutipan pengguna, artefak tersebut dipastikan jatuh menjadi *self-referential design* yang sama sekali tidak berpijak pada masalah pengguna secara riil. Kegagalan mempertahankan validitas empiris ini pada akhirnya akan menciptakan efek domino yang bermuara pada risiko penolakan produk (*rejection risk*) secara masif di pasar, sebagaimana yang telah diidentifikasi pada dimensi risiko sebelumnya.

Guna menyintesis dan mempertegas batasan operasional masing-masing antara praktisi dan AI, pemetaan batasan wilayah kerja antara keduanya dirangkum menjadi tiga dimensi kognitif. Tabel 7 mengilustrasikan secara komparatif batasan fundamental AI yang dihadapkan dengan otoritas eksklusif yang wajib dipertahankan oleh praktisi manusia.

Tabel 7 Komparasi batasan fundamental AI dalam perancangan *user persona*

Dimensi kognitif	Kapasitas komputasional AI	Otoritas eksklusif praktisi (manusia)
Empati dan Interaksi Antarmanusia	Terbatas pada pemrosesan sintaksis linguistik teks tanpa memiliki pengalaman interaksi sosial di dunia nyata.	Merasakan kedalaman emosi pengguna dan menarik kesimpulan makna dari interaksi sosial di dunia nyata.
Pengambilan Keputusan Kritis	Tidak memiliki kapasitas untuk memahami batasan regulasi perusahaan dan konteks bisnis riil yang kompleks.	Memegang otonomi strategis untuk mendefinisikan problem utama secara tepat sasaran (menghindari <i>algorithmic authority</i>).
Orisinalitas dan Validitas Empiris	Berisiko melakukan fabrikasi bahasa atau mengarang <i>pain points</i> dari ruang hampa jika dibiarkan tanpa landasan.	Menegakkan prinsip <i>anti zero-data generation</i> dengan memberikan transkrip <i>verbatim</i> murni dari lapangan sebagai konteks kepada mesin.

4.5.2 Mekanisme Mitigasi Risiko dalam Penggunaan AI

Tema terakhir ini menjelaskan berbagai strategi pertahanan dan prosedur yang diterapkan oleh praktisi untuk memitigasi risiko keamanan data serta kecacatan luaran AI. Praktisi secara proaktif melakukan kontrol ketat sejak tahap persiapan data hingga validasi akhir sebelum hasil riset disajikan dalam bentuk artefak persona.

Langkah pertama yang menjadi prioritas utama bagi mayoritas praktisi adalah perlindungan privasi dan kerahasiaan data perusahaan maupun pengguna.

Partisipan 1 (P1), Partisipan 7 (P7), dan Partisipan 8 (P8) secara konsisten menerapkan anonimisasi data sensitif sebelum mengunggahnya ke platform AI. Partisipan 6 (P6) bahkan menerapkan strategi pemilihan data yang selektif (*selective data disclosure*). Ia hanya menyertakan catatan hasil riset lapangan dan secara tegas mengeksklusi dokumen penting perusahaan seperti *business requirements* dari pengunggahan ke platform AI. P6 menegaskan batasan keamanan data ini secara lugas:

"All file business requirement tentu nggak dikasih, tapi lebih ke high level aja sih... Misalnya, research ini dilakukan untuk membuat aplikasi banking yang baru, tapi tidak spesifik nama aplikasinya atau rencana bisnisnya itu tidak dikasih tahu... Beneran hasil note taking-nya aja sesuai pertanyaan yang udah kita rancang."

Selain itu, P1 memanfaatkan penggunaan fitur percakapan sementara (*temporary chat*) guna mencegah algoritma AI menyimpan informasi internal perusahaan ke dalam memori pelatihannya. Tindakan preventif P1 ini merupakan bentuk nyata dari mekanisme pengecualian data untuk memitigasi risiko laten kebocoran data (*data leakage*) pada pemanfaatan *Large Language Models* (Yao et al. 2024).

Penerapan mekanisme mitigasi privasi ini terbukti sangat krusial, mengingat penelitian ini juga menemukan adanya anomali tingkat kesadaran (*awareness*) pada sebagian praktisi. Sebagai contoh *negative case*, Partisipan 3 (P3) justru menunjukkan tingkat kepercayaan (*over-trust*) yang sangat tinggi terhadap platform AI dengan menolak melakukan penyaringan data. P3 secara eksplisit menyatakan, "nggak sih, gue serahin semuanya... pokoknya ini data gue loh, gue masukin semuanya biar hasilnya lebih maksimal aja sih mikirnya." Kontras perilaku antara P3 dengan mayoritas partisipan lainnya membuktikan bahwa standar tata kelola privasi data dalam penggunaan *Generative AI* di ranah UX masih sangat bervariasi dan belum memiliki protokol industri yang seragam. Kondisi empiris ini secara nyata menegaskan urgensi perancangan *Standard Operating Procedure* (SOP) perlindungan privasi yang terstandardisasi secara ketat pada level korporat. Keamanan data riset dan informasi sensitif pengguna tidak boleh dibiarkan bergantung semata-mata pada tingkat literasi atau intuisi keamanan masing-masing individu periset. Tanpa adanya kebijakan tata kelola industri yang terpusat dan tegas, keluguan operasional semacam ini berpotensi membuka celah fatal bagi pelanggaran etika dan kebocoran data di masa depan.

Untuk memperjelas variasi tingkat literasi keamanan dan pendekatan operasional dari masing-masing praktisi tersebut, Tabel 8 merangkum klasifikasi strategi mitigasi privasi data yang ditemukan di lapangan. Matriks ini mendemonstrasikan spektrum kewaspadaan periset, mulai dari taktik pertahanan berlapis hingga anomali kepercayaan yang berlebihan.

Tabel 8 Taksonomi strategi mitigasi privasi data

Partisipan	Jenis strategi pertahanan	Tindakan operasional
P1, P7, P8	Anonimisasi Data Sensitif	Menyensor atau menyamarkan informasi rahasia dan identitas sebelum data mentah diunggah ke platform AI.

Tabel 8 Taksonomi strategi mitigasi privasi data (*lanjutan*)

Partisipan	Jenis strategi pertahanan	Tindakan operasional
P6	<i>Selective Data Disclosure</i>	Hanya mengunggah catatan hasil wawancara dan menahan dokumen rahasia perusahaan (<i>business requirements</i>).
P1	Penggunaan <i>Temporary Chat</i>	Mencegah algoritma menggunakan percakapan riset sebagai korpus pelatihan memori laten mesin (mitigasi <i>data leakage</i>).
P3 (Kasus Anomali)	<i>Zero Filtering (Over-trust)</i>	Menyerahkan seluruh data mentah secara utuh tanpa penyaringan sama sekali karena tingginya kepercayaan terhadap mesin.

Langkah kedua berfokus pada sinkronisasi pengetahuan dan pencegahan bias kognitif. Partisipan 5 (P5) menekankan pentingnya penyelarasan pengetahuan (*knowledge synchronization*) antara peneliti dan AI untuk meminimalkan halusinasi yang berujung pada iterasi *prompting* yang tidak berkesudahan. Sementara itu, Partisipan 8 (P8) menerapkan prosedur pencegahan *framing bias* dengan cara merumuskan hipotesis atau pandangan pribadi terlebih dahulu sebelum berinteraksi dengan AI. P8 menjabarkan prinsip pertahanan kognitif ini dengan sangat jelas:

"jangan tanya AI dulu sebelum gue punya hipotesis sendiri. Jadi sebelum gue prompt, gue paksa diri gue sendiri buat mikir dulu... Tujuannya supaya gue nggak langsung terbawa sama framing AI-nya."

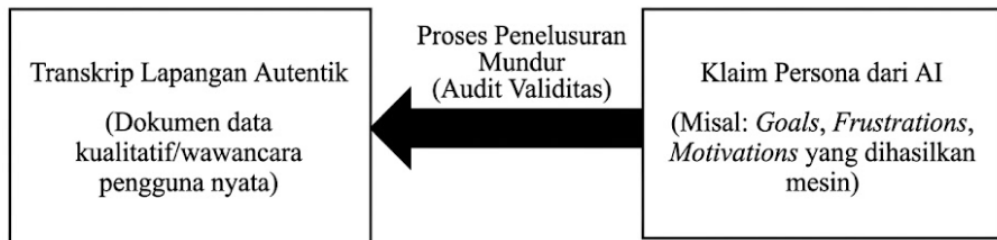
Pendekatan ini sangat penting agar peneliti tetap memiliki independensi analitis dan tidak mudah disetir oleh alur pemikiran awal yang didapatkan dari kecerdasan mesin.

Langkah ketiga sekaligus yang paling krusial adalah penerapan mekanisme *human-in-the-loop* sebagai basis penentuan standar kualitas. Partisipan 1 (P1), Partisipan 3 (P3), dan Partisipan 5 (P5) menegaskan bahwa setiap hasil dari AI tidak boleh langsung digunakan tanpa melalui proses peninjauan ulang yang komprehensif, mulai dari perbaikan tata bahasa hingga penghapusan poin-poin yang tidak relevan. Partisipan 2 (P2) menambahkan bahwa pemahaman teoretis yang mendalam terhadap fundamental riset merupakan syarat mutlak bagi periset sebelum mereka "diizinkan" menggunakan AI. Praktik pengawasan berlapis ini merepresentasikan tanggung jawab mutlak praktisi sebagai *"shepherds of AI systems"* (Floridi 2023), ketika periset memegang kendali penuh untuk meninjau ulang dan memvalidasi setiap luaran kecerdasan buatan guna menekan risiko bias dan halusinasi.

Terakhir, untuk menjamin akurasi informasi yang disajikan, Partisipan 7 (P7) dan Partisipan 8 (P8) menerapkan prinsip keterlacakan wawasan (*insight traceability*). Konsep ini mewajibkan agar setiap klaim sintetik yang masuk ke dalam artefak persona final harus dapat ditelusuri kembali validitasnya ke sumber data asli atau transkrip wawancara. Jika tidak ada bukti empiris yang mendukung, praktisi akan membuang informasi tersebut. P8 memberikan batasan yang sangat tegas terkait mekanisme pelacakan ini:

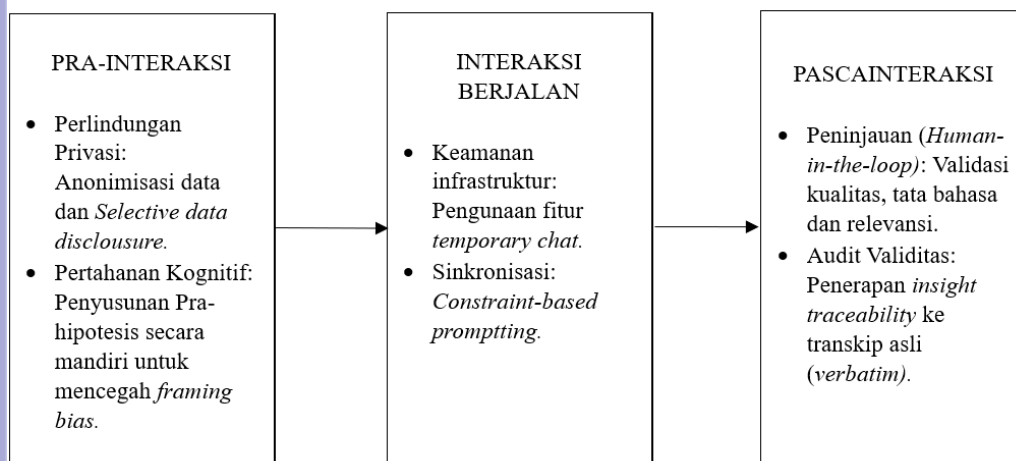
"setiap klaim yang masuk ke persona final harus bisa gue trace balik ke minimal satu transkrip... kalau nggak bisa... mending buang".

Gambar 6 mengilustrasikan skema *backward tracing* yang diterapkan oleh praktisi. Proses audit ini memastikan bahwa setiap elemen fiktif pada artefak AI ditelusuri mundur secara ketat hingga menemukan jangkar pembuktiannya pada dokumen transkrip yang autentik.



Gambar 6 Ilustrasi *backward tracing* dalam audit validitas artefak persona

Penerapan pelacakan data ini merepresentasikan arsitektur transparansi AI di level operasional. Liao dan Vaughan (2024) menyatakan bahwa pelacakan asal-usul data dan transparansi logika luaran model sangat penting agar manusia dapat melakukan audit secara efektif. Keputusan praktisi untuk membuang luaran AI yang tidak memiliki rekam jejak empiris ini membuktikan terjadinya proses kalibrasi kepercayaan (*trust calibration*) yang sangat ideal. Praktisi dalam hal ini secara proaktif menahan laju kepercayaan mereka agar tidak menjadi absolut, memastikan bahwa validitas artefak persona tetap berada di bawah kendali tata kelola manusia dan pengawasan kritis di setiap alur kerja riset. Keseluruhan langkah mitigasi operasional yang diterapkan oleh praktisi tersebut pada dasarnya membentuk sebuah *Standard Operating Procedure* (SOP) yang sekuensial dan berkesinambungan seperti digambarkan oleh Gambar 7.

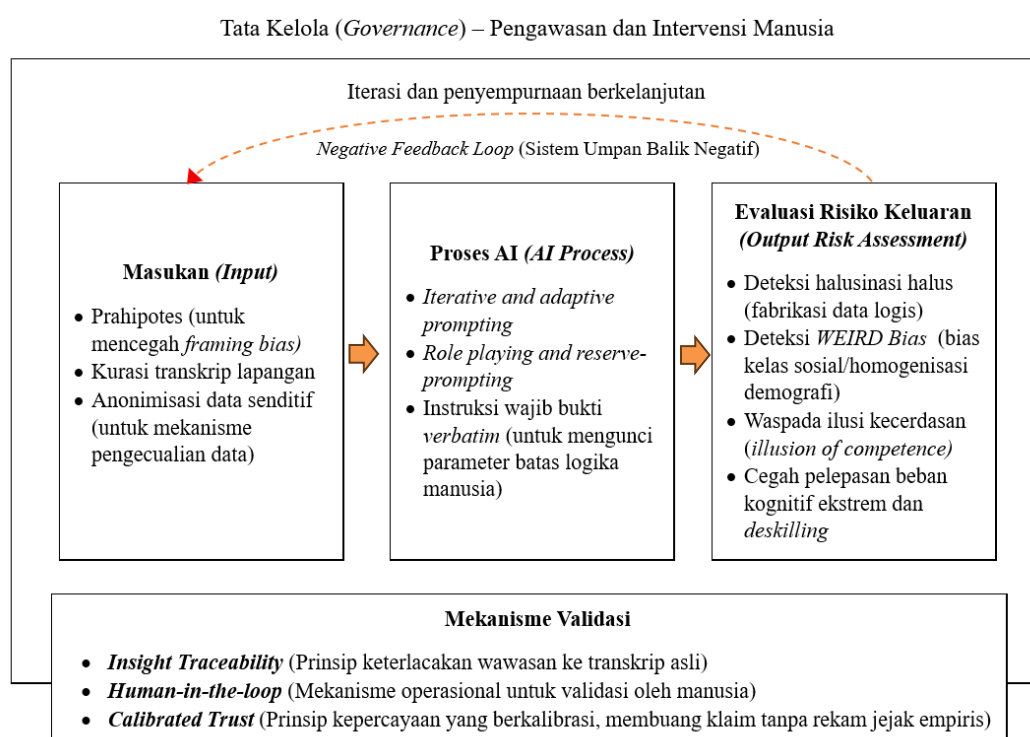


Gambar 7 Visualisasi SOP mitigasi risiko sosioteknis pada alur riset UX

Gambar 7 merangkum dan memvisualisasikan alur mitigasi tersebut ke dalam tiga fase kronologis: pra-interaksi, interaksi berjalan, dan pascainteraksi. Melalui visualisasi SOP ini, tergambar dengan jelas bahwa tata kelola manusia tidak hanya terjadi di satu titik evaluasi, melainkan membentang secara proaktif sebagai pagar pengaman yang mengapit seluruh siklus pendelegasian tugas kepada kecerdasan buatan.

4.6 Conceptual Framework for AI-Assisted Persona Creation

Sintesis menyeluruh dari ketiga dimensi tematik yang telah dibahas sebelumnya dikonstruksi menjadi sebuah kerangka kerja teoretis yang dinamakan *Conceptual Framework for AI-Assisted Persona Creation* pada Gambar 8. Kerangka ini disusun membentuk model tata kelola sosioteknis yang memetakan hubungan kausalitas antara pemrosesan algoritma dan tanggung jawab mutlak peneliti. Pemetaan ini bertujuan mengoptimalkan efisiensi asisten kognitif kecerdasan buatan sekaligus memberikan landasan teoretis untuk memitigasi bahaya otoritas algoritmik. Kerangka konseptual ini mentransformasikan temuan lapangan menjadi empat komponen utama yang saling berkesinambungan serta diikat oleh sebuah siklus iterasi dan penyempurnaan berkelanjutan seperti yang ditunjukkan pada Gambar 8.



Gambar 8 *Conceptual Framework for AI-Assisted Persona Creation*

a Masukan (Input)

Komponen pertama merepresentasikan gerbang awal yang menempatkan otoritas dan kendali penuh pada diri praktisi. Sebelum mendelegasikan tugas sintesis kepada sistem kecerdasan buatan, praktisi wajib menetapkan jangkar empiris melalui tiga langkah fundamental. Langkah pertama adalah pembentukan prahipotesis secara manual guna mencegah *framing bias*. Hal ini memastikan peneliti memiliki independensi analitis dan tidak terpengaruh oleh alur pemikiran awal yang disodorkan oleh mesin. Langkah kedua adalah kurasi transkrip lapangan yang riil sebagai bentuk penolakan mutlak terhadap praktik pembuatan data dari ruang hampa. AI tidak diizinkan untuk mengonstruksi realitas pengguna tanpa suplai data lapangan yang valid guna menghindari perancangan persona yang hanya berbasis pada asumsi. Langkah ketiga adalah

penerapan prosedur anonimisasi data sensitif secara ketat. Tindakan ini merupakan mekanisme pengecualian data yang krusial untuk mencegah risiko kebocoran informasi internal perusahaan ke dalam korpus pelatihan model LLM.

b Proses AI (*AI Process*)

Tahap kedua merupakan fase ketika model kecerdasan buatan beroperasi murni sebagai mesin pengolah statistik linguistik untuk merestrukturisasi data. Secara arsitektural, seluruh kumpulan transkrip yang telah lolos tahap anonimisasi di fase *Input* tersebut disuntikkan secara utuh sebagai *knowledge base* yang mengikat algoritma saat instruksi selanjutnya diberikan. Interaksi antara peneliti dan sistem tidak berjalan satu arah, melainkan dilakukan melalui taktik *prompting* yang iteratif dan adaptif. Peneliti menerapkan teknik *role-playing* untuk mengarahkan sudut pandang mesin, serta *reverse-prompting*, yakni sebuah taktik yang melibatkan instruksi untuk memaksa AI mengajukan pertanyaan balik kepada periset guna memvalidasi kelengkapan konteks sebelum AI mulai melakukan ekstraksi wawasan. Karena sistem kecerdasan buatan beroperasi sebagai entitas yang memiliki agensi tanpa pemahaman nyata, peneliti harus mengunci parameter pemrosesan mesin tersebut dengan batasan logika manusia. Penguncian ini dilakukan melalui instruksi yang mewajibkan penyertaan bukti *verbatim* pada setiap luaran yang dihasilkan. Kewajiban ini membatasi kebebasan algoritma dalam memanipulasi teks dan memaksa mesin untuk hanya mengekstrak wawasan yang benar-benar terdapat di dalam transkrip.

c Evaluasi Risiko Keluaran (*Output Risk Assessment*)

Komponen ketiga merepresentasikan tahapan analitis bagi praktisi untuk secara aktif mendeteksi dan mengidentifikasi potensi kegagalan dari pemrosesan kecerdasan buatan. Evaluasi ini memetakan kerentanan metodologis secara struktural maupun psikologis. Secara struktural, praktisi harus mengevaluasi keberadaan halusinasi halus (*plausible hallucination*), yaitu mendeteksi fabrikasi data yang terdengar sangat logis tetapi tidak memiliki akar empiris, serta mengidentifikasi ancaman bias kelas sosial (*WEIRD bias*) yang menyebabkan homogenisasi perilaku pengguna ke dalam stereotipe demografi yang tidak akurat. Secara psikologis, tahap ini menuntut periset untuk mewaspadaai kefasihan tata bahasa mesin yang berpotensi memicu fenomena ilusi kecerdasan (*illusion of competence*). Deteksi mandiri ini sangat penting agar praktisi tidak terdorong untuk menerima keluaran algoritma sebagai kebenaran objektif, yang dapat menjebak mereka ke dalam otoritas algoritmik, pelepasan beban kognitif yang ekstrem, hingga degradasi ketajaman analitis (*deskilling*) dalam jangka panjang.

d Tata Kelola (*Governance*)

Komponen keempat merupakan lapisan pertahanan yang mengembalikan kendali mutlak kepada manusia untuk menjustifikasi kebenaran artefak. Meskipun pada fase akhir komponen ini bertindak sebagai gerbang validasi melalui mekanisme *insight traceability*, kedudukan konseptual tata kelola manusia sejatinya bukanlah tahapan akhir yang terpisah. Secara konseptual pada ranah sistem sosioteknis, tata kelola merupakan sebuah instrumen pengawasan yang membungkus keseluruhan proses perancangan. Manifestasi kontrol berkesinambungan ini rohnya sudah ditanamkan sejak tahap Masukan

(*Input*) melalui kurasi data, dan terus mengawal setiap tahapan interaksi AI. Model LLM tidak memiliki pengalaman interaksi sosial di dunia nyata, sehingga mesin tidak akan pernah memiliki kapasitas untuk berempati atau mengambil keputusan yang menuntut pertanggungjawaban moral. Mitigasi secara operasional dilakukan melalui penerapan mekanisme *human-in-the-loop* sebagai basis penentuan standar kualitas akhir. Evaluasi mutu dilakukan dengan menerapkan validasi mutlak melalui prinsip keterlacakan wawasan (*insight traceability*). Setiap elemen persona yang dihasilkan oleh algoritma wajib ditelusuri kembali secara presisi ke dalam transkrip wawancara aslinya. Peneliti menjalankan prinsip kepercayaan yang terkalibrasi (*calibrated trust*) dengan cara membuang klaim yang diberikan oleh AI yang tidak memiliki rekam jejak empiris guna menjaga integritas desain produk.

e Iterasi dan Penyempurnaan Berkelanjutan

Seluruh komponen dalam kerangka kerja ini tidak bergerak dalam satu garis lurus yang statis, melainkan diikat oleh siklus sirkular berupa iterasi dan penyempurnaan berkelanjutan. Secara arsitektural komputasi, penolakan hasil artefak pada tahap Tata Kelola bertindak sebagai *negative feedback loop* (sistem umpan balik negatif) yang sangat krusial. Apabila keluaran algoritma tidak memenuhi syarat keterlacakan empiris atau terindikasi memuat bias, siklus umpan balik ini memaksa periset untuk melakukan recalibrasi instruksi dan penyelarasan konteks kembali pada tahap Masukan (*Input*). Proses recalibrasi sistem ini akan memperbaiki rumusan *prompt*, mempertajam konteks data, dan mengevaluasi kembali parameter batas yang sebelumnya ditetapkan. Siklus iterasi ini menegaskan bahwa validitas artefak persona dicapai melalui serangkaian negosiasi kognitif dan kalibrasi sistem yang ketat, bukan sekadar penerimaan pasif terhadap komputasi mesin.





V SIMPULAN DAN SARAN

5.1 Simpulan

Hasil penelitian ini menyimpulkan bahwa *Generative AI* difungsikan secara fundamental sebagai asisten kognitif dan bukan untuk menggantikan praktisi UX. Kehadiran teknologi ini memfasilitasi pelepasan beban kognitif (*cognitive offloading*) untuk mengakselerasi sintesis data mentah di tengah tekanan operasional dan keterbatasan waktu di industri. Dalam menavigasi kolaborasi tersebut, interaksi antara praktisi dan sistem kecerdasan buatan dieksekusi melalui paradigma instruksi (*prompting*) yang adaptif dan berbasis batasan. Taksonomi strategi operasional ini bertindak sebagai mekanisme pertahanan metodologis untuk mengunci kebebasan improvisasi algoritma, guna memastikan bahwa keluaran model tetap berjangkar secara mutlak pada bukti empiris dari lapangan.

Meskipun menawarkan akselerasi, pemanfaatan teknologi ini memicu berbagai risiko yang saling bereskalasi secara struktural dan psikologis. Secara struktural, artefak persona yang dihasilkan mengandung cacat bawaan berupa anomali halusinasi data dan bias kelas sosial (*WEIRD bias*) yang mereduksi keberagaman pengguna menjadi representasi artefak yang homogen. Secara psikologis, kefasihan tata bahasa AI sering kali memicu fenomena ilusi kecerdasan (*illusion of competence*). Kondisi ini mengancam ketahanan kognitif praktisi karena berisiko memunculkan bias otomatisasi (*automation bias*), menumbuhkan kepercayaan berlebihan (*over-trust*), dan pada puncaknya bermuara pada degradasi ketajaman insting analitis (*deskilling*) secara jangka panjang.

Sebagai kontribusi utama, penelitian ini menyintesis keseluruhan temuan tersebut ke dalam *Conceptual Framework for AI-Assisted Persona Creation*. Kerangka konseptual ini menegaskan bahwa elemen empati, interpretasi emosi pengguna, dan pengambilan keputusan strategis merupakan domain eksklusif manusia yang tidak dapat diotomatisasi. Validitas artefak UX di era *generative AI* hanya dapat dipertanggungjawabkan apabila ekosistem perancangan secara ketat mengintegrasikan intervensi manusia (*human-in-the-loop*) dan penegakan prinsip keterlacakan wawasan (*insight traceability*), guna memastikan otoritas keputusan desain tidak diambil alih oleh algoritma mesin.

5.2 Saran

Berdasarkan temuan dan keterbatasan dalam penelitian ini, terdapat beberapa saran strategis yang dirumuskan untuk agenda penelitian selanjutnya. Evaluasi terhadap implementasi *Generative AI* dalam proses perancangan persona ini membuka ruang yang luas bagi pengembangan studi di masa mendatang, sehingga beberapa rekomendasi yang diajukan mencakup:

1. Penelitian selanjutnya dapat menggunakan pendekatan *mixed methods* atau metode kuantitatif untuk membandingkan secara langsung tingkat akurasi dan kualitas *user persona* yang disintesis secara manual dengan yang dibantu oleh *Generative AI*. Pengukuran kualitas ini secara operasional dapat dieksekusi melalui metode pengujian *A/B testing* yang divalidasi melalui evaluasi pakar (*expert review*) untuk mengukur tingkat presisi wawasan mesin dibandingkan dengan ketajaman analitis manusia.

- 2 Melibatkan sampel partisipan yang lebih masif dari berbagai sektor industri dan tingkat pengalaman manajerial yang lebih bervariasi, guna memperluas cakupan dan keberagaman hasil penelitian.
- 3 Ruang lingkup pemanfaatan *Generative AI* dapat dieksplorasi lebih jauh signifikansinya pada tahapan proses kerja UX lainnya di luar pembuatan persona, seperti *usability testing*, *customer journey mapping*, atau penyusunan skenario penggunaan (*use cases*).
- 4 Penelitian berikutnya direkomendasikan untuk merumuskan kerangka standardisasi protokol kepatuhan privasi data dan etika AI yang seragam di tingkat industri dalam ruang lingkup riset pengguna, guna memastikan integritas praktik perancangan produk dan batasan tanggung jawab profesional tetap terjaga secara institusional.

@Hak cipta milik IPB University

IPB University



Hak Cipta Dilindungi Undang-undang
1. Dilarang mengutip sebagian atau seluruh karya tulis ini tanpa mencantumkan dan menyebutkan sumber :

- a. Pengutipan hanya untuk kepentingan pendidikan, penelitian, penulisan karya ilmiah, penyusunan laporan, penulisan kritik atau tinjauan suatu masalah
- b. Pengutipan tidak merugikan kepentingan yang wajar IPB University.

2. Dilarang mengumumkan dan memperbanyak sebagian atau seluruh karya tulis ini dalam bentuk apapun tanpa izin IPB University.



DAFTAR PUSTAKA

- Bender EM, Gebru T, McMillan-Major A, Shmitchell S. 2021. On the dangers of stochastic parrots: can language models be too big? Di dalam: Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FAccT '21); 2021 Mar 3–10; Virtual Event, Kanada. New York (NY): ACM. hlm 610–623. doi:10.1145/3442188.3445922.
- Bevilacqua R, Bailoni T, Maranesi E, Amabili G, Barbarossa F, Ponzano M, Virgolesi M, Rea T, Illario M, Piras EM, *et al.* 2025. Framing the human-centered artificial intelligence concepts and methods: scoping review. *JMIR Hum Factors*. 12:e67350. doi:10.2196/67350.
- Braun V, Clarke V. 2006. Using thematic analysis in psychology. *Qual Res Psychol*. 3(2):77–101. doi:10.1191/1478088706qp063oa.
- Braun V, Clarke V. 2020. One size fits all? What counts as quality practice in (reflexive) thematic analysis? *Qual Res Psychol*. 18(3):328–352. doi:10.1080/14780887.2020.1769238.
- Brown H, Lee K, Mireshghallah F, Shokri R, Tramèr F. 2022. What does it mean for a language model to preserve privacy?. Di dalam: Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency (FAccT '22); 2022 Jun 21–24; Seoul, Republik Korea. New York (NY): ACM. hlm 1–13. doi:10.1145/3531146.3534642.
- Brown JM, Lindgaard G, Biddle R. 2011. Collaborative events and shared artefacts: agile interaction designers and developers working toward common aims. Di dalam: Proceedings of the 2011 Agile Conference (AGILE '11); 2011 Agu 7–13; Salt Lake City, AS. Washington (DC): IEEE Computer Society. hlm 87–96. doi:10.1109/AGILE.2011.45.
- Carlini N, Tramèr F, Wallace E, Jagielski M, Herbert-Voss A, Lee K, Roberts A, Brown T, Song D, Erlingsson Ú, *et al.* 2021. Extracting training data from large language models. Di dalam: Proceedings of the 30th USENIX Security Symposium; 2021 Agu 11–13; Virtual Conference. Berkeley (CA): USENIX Association. hlm 2633–2650; [diakses 2026 Jun 29]. <https://www.usenix.org/conference/usenixsecurity21/presentation/carlini-extracting>.
- Charmaz K. 2014. *Constructing Grounded Theory*. Ed ke-2. London: 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. *Bus Inf Syst Eng*. [diakses 2026 Mei 12]; 61(5):637–643. doi:10.1007/s12599-019-00595-2.
- Dwivedi YK, Kshetri N, Hughes L, Slade EL, Jeyaraj A, Kar AK, Baabdullah AM, Koohang A, Raghavan V, Ahuja M, *et al.* 2023. “So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *Int J Inf Manage*. 71:102642. doi:10.1016/j.ijinfomgt.2023.102642.

@Hak cipta milik IPB University

Hak Cipta Dilindungi Undang-undang
1. Dilarang mengutip sebagian atau seluruh karya tulis ini tanpa mencantumkan dan menyebutkan sumber :
a. Pengutipan hanya untuk kepentingan pendidikan, penelitian, penulisan karya ilmiah, penyusunan laporan, penulisan kritik atau tinjauan suatu masalah
b. Pengutipan tidak merugikan kepentingan yang wajar IPB University.
2. Dilarang mengumumkan dan memperbanyak sebagian atau seluruh karya tulis ini dalam bentuk apapun tanpa izin IPB University.

- Floridi L. 2023. AI as agency without intelligence: on ChatGPT, large language models, and other generative models [editorial]. *Philos Technol.* 36:15. doi:10.1007/s13347-023-00621-y.
- Gothelf J, Seiden J. 2013. *Lean UX: Applying Lean Principles to Improve User Experience*. Sebastopol (CA): O'Reilly Media.
- Guest G, Bunce A, Johnson L. 2006. How many interviews are enough? An experiment with data saturation and variability. *Field Methods.* 18(1):59–82. doi: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. doi:10.1177/0018720814547570.
- Ji Z, Lee N, Frieske R, Yu T, Su D, Xu Y, Ishii E, Bang YJ, Madotto A, Fung P. 2023. Survey of hallucination in natural language generation. *ACM Comput. Surv.* 55(12):Article 248. doi:10.1145/3571730.
- 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. *J User Exp.* 20(4):149–165. <https://uxpajournal.org/leveraging-ai-toward-the-development-of-vector-personas-for-ux-research/>.
- Li J, Cao H, Lin L, Hou Y, Zhu R, Ali AE. 2024. User experience design professionals' perceptions of generative artificial intelligence. Di dalam: Proceedings of the CHI Conference on Human Factors in Computing Systems; 2024 Mei 11–16; Honolulu, Amerika Serikat. New York (NY): ACM. hlm 1–18. doi:10.1145/3613904.3642114.
- Liao QV, Vaughan JW. 2024. AI transparency in the age of LLMs: a human-centered research roadmap. *Harv Data Sci Rev.* 5(SP5). doi:10.1162/99608f92.8036d03b.
- 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: LPPM UK Petra. hlm 53–67.
- Molléri JS, Marculescu B. 2025. Data-driven personas for software engineering research. Di dalam: Mannion M, Mannisto T, Maciaszek L, editor. 20th International Conference on Evaluation of Novel Approaches to Software Engineering (ENASE 2025); 2025 Apr 4–6; Porto, Portugal. Setúbal: SCITEPRESS. hlm 798–805. [diakses 2026 Jul 22]. doi:10.5220/0013475100003928.
- Mumford E. 2006. The story of socio-technical design: reflections on its successes, failures and potential. *Inf Syst J.* 16(4):317–342. doi: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.
- Paoli S. 2023. Improved prompting and process for writing user personas with LLMs, using qualitative interviews: capturing behaviour and personality traits of users. *arXiv*. [diakses 2026 Jun 22]; doi:10.48550/arxiv.2310.06391.
- Parasuraman R, Manzey DH. 2010. Complacency and bias in human use of automation: an attentional integration. *Human Factors*. 52(3):381–410. doi:10.1177/0018720810376055.
- Raisch S, Fomina K. 2024. Combining human and artificial intelligence: hybrid problem-solving in organizations. *Acad of Manage Rev*. doi:10.5465/amr.2021.0421.
- 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. doi:10.1145/3411763.3451760.
- Risko EF, Gilbert SJ. 2016. Cognitive offloading. *Trends Cogn Sci*. 20(9):676–688. doi:10.1016/j.tics.2016.07.002.
- Salah D, Paige RF, Cairns P. 2014. A systematic literature review for agile development processes and user centred design integration. Di dalam: Proceedings of the 18th International Conference on Evaluation and Assessment in Software Engineering (EASE '14); 2014 Mei 13–14; London, Inggris. New York (NY): ACM. Article 5. doi:10.1145/2601248.2601276.
- Salminen J, Liu C, Pian W, Chi J, Häyhänen E, Jansen BJ. 2024. Deus ex machina and personas from large language models: investigating the composition of AI-generated persona descriptions. Di dalam: Mueller FF, Kyburz P, Williamson JR, Sas C, Wilson ML, Touns Dugas P, Shklovski I, editor. Proceedings of the CHI Conference on Human Factors in Computing Systems; 2024 Mei 11–16; Honolulu, Amerika Serikat. New York (NY): ACM. hlm 1–20. doi:10.1145/3613904.3642036.
- Shneiderman B. 2022. *Human-Centered AI*. Oxford: Oxford Univ Pr.
- 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. doi:10.33795/jip.v11i4.7625.
- Strahler D. 2025. Generative AI (GenAI) for co-creating user personas. Di dalam: *Proceedings of the New York State Communication Association*. 2024(1):3. [diakses 2026 Mei 25]. <https://docs.rwu.edu/nyscaproceedings/vol2024/iss1/3/>.
- Truss M. 2026. PersonaCite: VoC-grounded interviewable agentic synthetic AI personas for verifiable user and design research. Di dalam: Proceedings of the Extended Abstracts of the 2026 CHI Conference on Human Factors in Computing Systems; 2026 Apr 13–17; Barcelona, Spanyol. New York (NY): ACM. hlm 1–7. doi:10.1145/3772363.3798543.
- 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>.

Hak Cipta Dilindungi Undang-undang

1. Dilarang mengutip sebagian atau seluruh karya tulis ini tanpa mencantumkan dan menyebutkan sumber :
a. Pengutipan hanya untuk kepentingan pendidikan, penelitian, penulisan karya ilmiah, penyusunan laporan, penulisan kritik atau tinjauan suatu masalah
b. Pengutipan tidak merugikan kepentingan yang wajar IPB University.
2. Dilarang mengumunkan dan memperbanyak sebagian atau seluruh karya tulis ini dalam bentuk apapun tanpa izin IPB University.

- Yao Y, Duan J, Xu K, Cai Y, Sun Z, Zhang Y. 2024. A survey on large language model (LLM) security and privacy: the good, the bad, and the ugly. *High-Confid Comput.* 4:100211. doi:10.1016/j.hcc.2024.100211.
- Yildirim I, Paul LA. 2024. From task structures to world models: what do LLMs know? *Trends Cogn Sci.* [diakses 2026 Mei 12]; 28(5):404–415. doi:10.1016/j.tics.2024.02.008.
- Zamfirescu-Pereira JD, Wong RY, Hartmann B, Yang Q. 2023. Why Johnny can't prompt: how non-AI experts try (and fail) to design LLM prompts. Di dalam: Schmidt A, editor. *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*; 2023 Apr 23–28; Hamburg, Jerman. New York (NY): ACM. hlm 1–21.
- Zavolokina L, Sprenkamp K, Katashinskaya Z, Jones DG. 2025. Biased by design: leveraging inherent AI biases to enhance critical thinking of news readers. Di dalam: Krasnova H, Kranz J, Lindgren R, editor. *Proceedings of the Thirty-Third European Conference on Information Systems (ECIS 2025)*; 2025 Jun 18; Amman, Yordania. Atlanta: Association for Information Systems. hlm 8; [diakses 2026 Jun 4]. <https://aisel.aisnet.org/ecis2025/hci/hci/8>.

