



Erosion Of Self: Ontological Security and Sense of Agency in A Technology Driven World

¹**Shreeya Sharan**, *Master of Arts, Clinical Psychology, Amity Institute of Psychology and Allied Sciences Amity University, Noida, India*

shreeyaamity@gmail.com

²**Dr. Rajat Kanti Mitra**, *Professor, Amity Institute of Psychology and Allied Sciences Amity University, Noida, India*

rkmitra@amity.edu

Abstract: The proliferation of artificial intelligence tools in professional life has generated urgent questions about the psychological experience of those who use them daily. This study investigated how tech-literate adult professionals aged 27 to 47 construct and narrate experiences of ontological security and sense of agency within AI-mediated professional and personal environments. Ontological security, understood following Giddens (1991) as the tacit biographical continuity that underpins confident engagement with daily life, and sense of agency, theorised by Moore (2016) as the first person experience of authoring one's own actions, formed the two central theoretical constructs. A qualitative design grounded in Reflexive Thematic Analysis (Braun and Clarke, 2006, 2019, 2022) was employed, informed by a constructionist epistemology. Eight adult professionals were recruited through purposive sampling across two generational cohorts, aged 27 to 37 and 37 to 47 respectively, reflecting different biographical relationships to the emergence of AI in professional practice. Data were generated through individual, semi-structured, in-depth interviews and analysed through comprehensive manual coding supported by ATLAS.ti for cross-case organisation and query. Analysis produced 55 codes organised into six thematic clusters: Embodied Temporality, The Boundary Question, Agency and Authorship, AI as Ontological Mirror, Functional Motivations, and Future Self. Findings indicate that AI integration does not produce a uniform erosion of ontological security or sense of agency but rather actively restructures both,

generating a self that is rhythmically mobile around technology engagement, strategically narrated in defence of authorship, and existentially negotiated in relation to perceived cognitive atrophy, impostor experience, and professional obsolescence anxiety. A generational inflection was identified: the 27 to 37 cohort tended to narrate boundary blurring and AI collaboration as adaptive and normative, while the 37 to 47 cohort more frequently narrated the same experiences as confrontational and requiring active management. The study contributes an empirically grounded experiential account of AI-mediated selfhood to a literature that has remained largely theoretical, and identifies biographical timing of AI adoption as a theoretically significant moderator of ontological experience.

Keyword: *Ontological Security, Sense of Agency, Artificial Intelligence, Reflexive Thematic Analysis, Professional Identity, Cognitive Offloading, Generational Cohort.*

1. Introduction

A person opens an AI tool, types a description of what they need, and receives a polished, coherent response. They edit it, send it, and move on. The task is done. The deadline is met. What happens next is unremarkable: another task, another prompt, another output. What does not happen is any deliberate pause over a question the sequence quietly raises, and raises across millions of professional interactions every day. When the reasoning was done by the system and the editing was done by the person, where did authorship begin? Not in the legal sense. Not in the sense that matters to the client. In the sense that matters to psychology: the sense of being the one who thought it. When something external to you is very capable of doing what you considered most distinctively yours, what happens to the meanings you attach to doing it?

The study makes no claim that AI is harmful, only that the questions raised by the ordinary rhythms of AI-mediated work deserve rigorous psychological investigation, particularly how professionals collectively construct meaning around them.

Ontological security, introduced by Laing (1960) and developed sociologically by Giddens (1991), describes the basic psychological confidence that one's existence is real, continuous, and coherent. (Laing, 1960) (Giddens, 1991) This assurance is not self-generating; it is maintained through encounters with others that affirm subjectivity rather than reduce it to an object. When those encounters consistently fail, when the person experiences themselves as predictable, modelled, and externally managed, the self withdraws from genuine engagement and becomes, in Laing's term, petrified.

Anthony Giddens argued that Ontological Security rests on three interconnected mechanisms: routine, which generates practical familiarity with the world; basic trust in persons and institutions, which makes those routines reliable; and biographical narrative coherence, which integrates experience into a sense of personal trajectory that extends meaningfully into the future. Giddens named this ensemble the protective cocoon. His structuration theory (1984) provides the generative mechanism: the habits of practice through which daily life is conducted do not merely reflect the self but continuously reproduce it. (Giddens, 1984) When the practices through which a person routinely produces their professional outputs are restructured by an external system, the routine itself changes, and the condition it was sustaining may change with it. Balkan-Sahin (2025) confirmed through bibliometric analysis that the construct's application to individual psychological experience in technology contexts remains substantially underdeveloped. (Balkan-Sahin, 2025)

Sense of agency refers to the subjective experience of being the initiating, executing, and controlling source of one's own actions. (Moore, 2016) It is the felt quality of authorship: the experience that one's intentions are the genuine origin of one's outcomes. Gallagher (2000) distinguished sense of agency from sense of ownership, noting that the two can come apart. (Gallagher, 2000) A person can recognise an output as theirs without experiencing themselves as its genuine author. This dissociation is analytically important for the present study because it describes precisely the experiential state that AI-assisted professional work may produce across a group of participants: outputs that are recognised as one's own professional products while being narrated as only partially self-authored. Bandura (2001) situated agency within a comprehensive theory of human functioning, identifying self-reflectiveness as the most distinctively human and most delegation-vulnerable dimension: the capacity to examine one's own reasoning is precisely what systematic cognitive delegation to AI removes from ordinary professional practice. (Bandura, 2001)

The study was built to bridge a gap in literature by shifting the focus from the aftermath of AI to the ongoing experience of oneself that AI tailors. The scope of existing literature is revolves around cognitive offloading and a binary viewpoint that demarks cognitive abilities and self to be either completely disrupted or enhanced. This study makes an attempt to provide empirical qualitative grounding for individual-level ontological security and its disruption among working professionals who interact with AI as part of their daily workflow. The study also extends literature on the sense of agency at a more decision-making level within a body of work that has predominantly focused on the physical nature of control. This study conceptualizes a spectrum of selfhood, with a fluid nature that is continuously being altered

with every interaction. This study also considers the historical context of individuals as a factor that may affect one's sense of agency and ontological security.

The study is a qualitative inquiry with eight adult professionals, divided into two age cohorts, who use AI tools extensively in both their professional and personal lives. The method is Reflexive Thematic Analysis, chosen because RTA is designed to identify and interpret the patterns, tensions, and contradictions in how participants collectively talk about and negotiate their experiences, rather than to test hypotheses or produce intensive individual case descriptions. (Braun & Clarke, 2006) (Braun & Clarke, 2019) (Braun & Clarke, 2022) It asks how participants collectively construct meaning around the experience of AI-mediated cognition, what thematic patterns organise those accounts, and whether any thematically significant differences emerge between participants who came to AI use at different stages of their professional development.

2. Methodology

2.1 Aim of the Study

This study investigates how tech-literate adult professionals (across two generational cohorts, aged 27 to 37 and 37 to 47 respectively) construct, narrate, and negotiate experiences of ontological security and sense of agency within AI-mediated professional and personal environments. The aim is to develop an interpretive, theoretically grounded account of how adult professionals make sense of what that use does to their experienced selfhood, capability, and authorship.

2.2 Research Questions

The study is organised around one central research question and four sub-questions. The central research question is:

How do tech-literate adult professionals aged 27 to 47 construct and narrate experiences of ontological security and sense of agency in AI-mediated professional and personal environments?

The sub-questions guiding the analysis are:

- SQ1: How do participants narrate the relationship between their own thinking, authorship, and the outputs co-produced with AI tools?
- SQ2: In what ways, if any, do participants describe their sense of cognitive capability, self-efficacy, or creative ownership as altered by sustained AI use?
- SQ3: How do participants account for transitions between AI-mediated and non-AI-mediated experience, and what does this reveal about ontological groundedness?

- SQ4: In what ways, if any, do the meanings participants construct around AI use differ between the 27–37 and 37–47 cohorts?

2.3 Research Design

This study adopts a qualitative research design grounded in Reflexive Thematic Analysis. (Braun & Clarke, 2006) (Braun & Clarke, 2019) (Braun & Clarke, 2022) RTA was selected because the central research question concerns how patterns of meaning are constructed and negotiated. RTA is specifically designed to identify, analyse, and report such patterns, making it the most appropriate methodological choice given the study's aims and epistemological orientation.

RTA is grounded in a constructionist epistemology. Meaning is not discovered in data but constructed through the researcher's active, reflexive engagement with it. The researcher's theoretical commitments, professional background, and personal experience are treated as constitutive of the analytical process and documented throughout in a reflexivity journal.

2.4 Sample and Sampling Technique

2.4.1 Sample Characteristics and Size Rationale

Eight participants were recruited, all working professionals with active and sustained AI tool use as a defining feature of both professional and personal daily routines. These participants were divided into two cohorts, those aged 27 to 37 and those aged 37 to 47.

A sample of eight participants is appropriate for RTA research where the dataset is subdivided for comparative analysis. Braun and Clarke (2021) note that sample size in qualitative research should be determined by the aims and design of the study rather than fixed numerical rules, and that six to ten participants is appropriate for studies aiming for thematic richness and transferability within a bounded, relatively homogeneous sample. (Braun & Clarke, 2021)

2.4.2 Sampling Strategy

This study employed purposive sampling oriented toward producing data rich enough to support analytic depth. The criterion for inclusion was: adult professionals using AI tools extensively for both professional and personal purposes on a sustained, habitual, daily basis. Participants were further selected for their capacity for reflective, articulate engagement with open-ended questions.

Participants were from a variety of domains to prevent conflation of domain-specific experiences with more broadly shared experiential patterns.

Recruitment proceeded through criterion-based snowball sampling within professional networks, with initial contacts meeting the inclusion criteria asked to refer others with comparable profiles.

2.5 Data Collection

2.5.1 Interview Procedure

A semi-structured interview schedule of 16 questions accompanied by probes were developed to invite depth, specificity, and experiential reflection. Each interview was conducted individually in a private setting appropriate for open, reflective conversation. Participants provided written informed consent after receiving information about the study's purpose and their right to withdraw. Interviews ranged from approximately 50 to 75 minutes. All were audio-recorded with participant consent and subsequently transcribed verbatim; transcripts were checked against recordings for accuracy. Field notes made immediately following each interview captured quality of engagement, topics participants returned to spontaneously, and initial analytical impressions, forming part of the reflexive audit trail.

2.5.2 Reflexive Positioning

The researcher's prior theoretical engagement with ontological security and sense of agency, personal experience of AI-mediated cognitive work, and interpretive responses to participants' accounts were treated as methodologically relevant and documented throughout in a reflexivity journal. The researcher's prior orientation, that AI integration follows a progression from task delegation toward deeper cognitive dependence that may quietly unsettle selfhood, was recognised as a potential source of confirmation bias and actively monitored across all phases of the study.

2.6 Data Analysis

2.6.1 Manual Coding

Data analysis followed the six-phase RTA process. (Braun & Clarke, 2006) (Braun & Clarke, 2019) (Braun & Clarke, 2021) (Braun & Clarke, 2022) The process is recursive rather than linear: phases are revisited as analysis develops.

Manual coding was the primary mode of interpretive engagement throughout Phase 2. All eight transcripts were coded in full by the researcher, working directly from verbatim text without automated coding functions. This reflects the epistemological requirements of RTA themes must emerge from close, sustained reading rather than algorithmic pattern recognition and the ethical concerns about automated coding discussed in Section 3.6.2.

Coding proceeded iteratively. An initial pass through each transcript produced provisional codes, refined through review against the full dataset. Following theme

construction, a comparative analysis examined whether themes appeared in both cohorts, with different emphases, or were cohort-specific.

2.6.2 ATLAS.ti: Software Assisted Analysis and the Rationale for Sequencing

Following manual coding, all codes and reflexive annotations were imported into ATLAS.ti as an organisational and analytical infrastructure, not as a coding engine. The relationship between manual coding and ATLAS.ti in this study reflects a deliberate epistemological sequencing. Manual coding establishes interpretive depth; ATLAS.ti provides the infrastructure for systematic cross-case review.

This approach deliberately excludes ATLAS.ti's automated coding functions, including AI Coding, Opinion Mining, and Sentiment Analysis, on both methodological and ethical grounds. Methodologically, automated coding lacks the interpretive precision required by RTA: it identifies surface semantic patterns rather than the latent, theoretically grounded meanings that constitute the analytic object of this study. Ethically, automated coding involves transmission of transcript data to external AI servers, raising confidentiality concerns inconsistent with the study's ethical obligations. All interpretive work is therefore researcher-generated.

2.7 Quality and Rigour

Quality in RTA is assessed by criteria appropriate to interpretive qualitative inquiry. Braun and Clarke (2022) outline criteria specific to RTA, applied alongside Yardley's (2000) broader qualitative quality framework. The criteria consist of Coherence, Transparency, Richness and Transferability. (Braun & Clarke, 2022) (Yardley, 2000)

3. Results

This chapter presents the findings of a qualitative study analysing how eight tech-literate adult professionals aged 27 to 47 construct and narrate experiences of ontological security and sense of agency within AI-mediated professional and personal environments. Data were analysed using Reflexive Thematic Analysis, combining comprehensive manual coding with ATLAS.ti-supported code organisation and cross-case analysis. (Braun & Clarke, 2006) (Braun & Clarke, 2019) (Braun & Clarke, 2021) (Braun & Clarke, 2022)

Participant accounts are presented below using anonymised codes P1 through P8. All quotations are reproduced verbatim from transcripts; minor disfluencies (e.g., repeated filler words) have been edited for readability without altering meaning.

Table 1. *Overview of Six Themes: Codes, Sub-Themes, and Research Question Alignment*

Theme	Theme Title	Codes (n)	Research Question Link
A	Embodied Temporality — The Rhythm of Self Around AI	17	SQ3
B	The Boundary Question — Negotiating the Self-AI Divide	4	SQ1, SQ3
C	Agency, Authorship, and the Control Narrative	3	SQ1, SQ2
D	AI as Ontological Mirror — Effects on Self-Concept	20	SQ1, SQ2, SQ3
E	Why I Reach for AI — Functional and Psychological Motivations	8	SQ2
F	Future Self in an AI World — Anticipation, Strategy, and Anxiety	3	SQ1, SQ3

Note. *N* = 8 participants. Total codes = 55. SQ = Sub-Question. Code distribution across participants is presented in Table 4. Full codebook with analytical annotations is reproduced in Appendix B.

Table 2. *Code Groundedness by Theme and Participant (Frequency of Coded Segments)*

Theme	P1	P2	P3	P4	P5	P6	P7	P8	Total
A — Embodied Temporality	5	4	6	5	4	5	4	5	38
B — Boundary Question	1	2	1	2	2	1	1	1	11
C — Agency & Authorship	2	2	1	3	2	2	2	2	16
D — Ontological Mirror	7	6	8	5	8	7	7	6	54
E — Functional Motivations	3	4	3	4	3	4	3	3	27
F — Future Self	1	1	2	2	2	2	2	1	13

Total per participant	19	19	21	21	21	21	19	18	159
------------------------------	-----------	-----------	-----------	-----------	-----------	-----------	-----------	-----------	------------

Note. Frequencies reflect the number of transcript segments coded within each Theme per participant.

3.1 Theme A: Embodied Temporality - The Rhythm of Self Around AI

3.1.1 Sub-Theme A1: Pre-AI Embodied Self

Across both cohorts, participants described the non-AI self as distinctively grounded, embodied, and available. This pre-AI baseline emerged not merely as an absence of technology but as a qualitatively richer mode of selfhood characterised by sensory presence, effortlessness, and physical solidity. For participants in the 37–47 cohort, the pre-AI self was frequently narrated in explicitly generational terms, positioned as the self - formed before AI became ambient — a biographical anchor that AI use now unsettles or displaces.

"When I am not using anything, not even my phone, things feel very grounded. This is how we grew up. AI came only four or five years ago. Without it, I feel completely in my own body, aware of everything around me." (P5)

This account illustrates what Giddens (1991) described as the protective cocoon, the pre-reflective sense of ontological security that underpins confident engagement with daily life. P5's framing of pre-AI experience as biographical norm ('this is how we grew up') positions AI integration as a departure from an established self-rhythm rather than a seamless extension of it.

3.1.2 Sub-Theme A2: Flow During AI Engagement

During AI engagement, several participants described a state of cognitive absorption they located in the task rather than the tool.

"When the work is going well and the AI is helpful, I am completely absorbed but it is the problem I am absorbed in, not the tool. I forget the AI is there. I forget I am there, in a way." (P4)

P4's account 'I forget the AI is there. I forget I am there' points to a flow state in which the sense of agency becomes distributed across the human-tool system rather than residing unambiguously in the participant. This is consistent with Bandura's (2001) account of proxy agency, in which the agent experiences the extended system's outputs as their own, and with Merleau-Ponty's (1962) concept of tool incorporation: when the tool is functioning smoothly, it disappears from experience, becoming transparent rather than an object of attention.

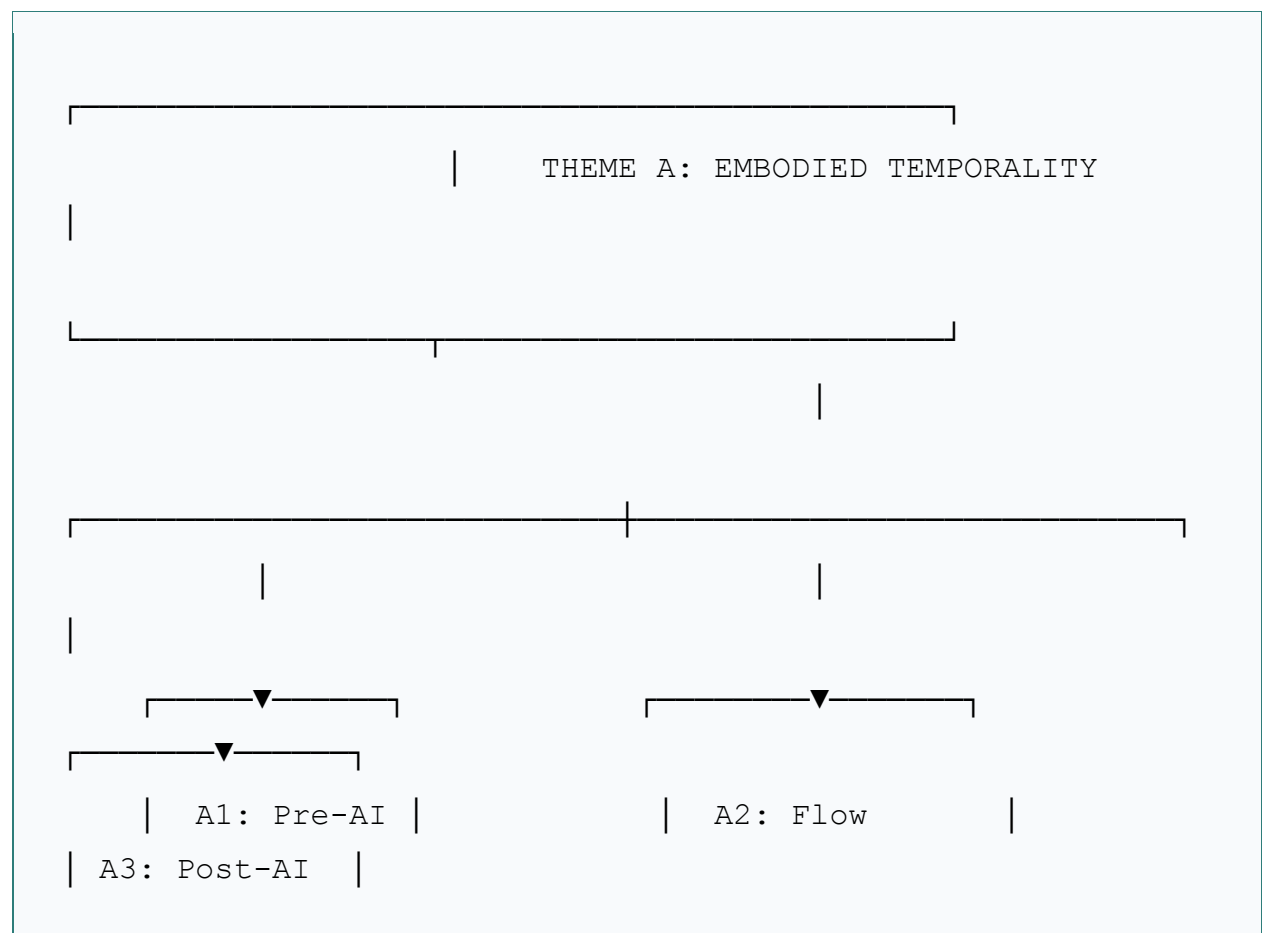
3.1.3 Sub-Theme A3: Post-AI Re-grounding

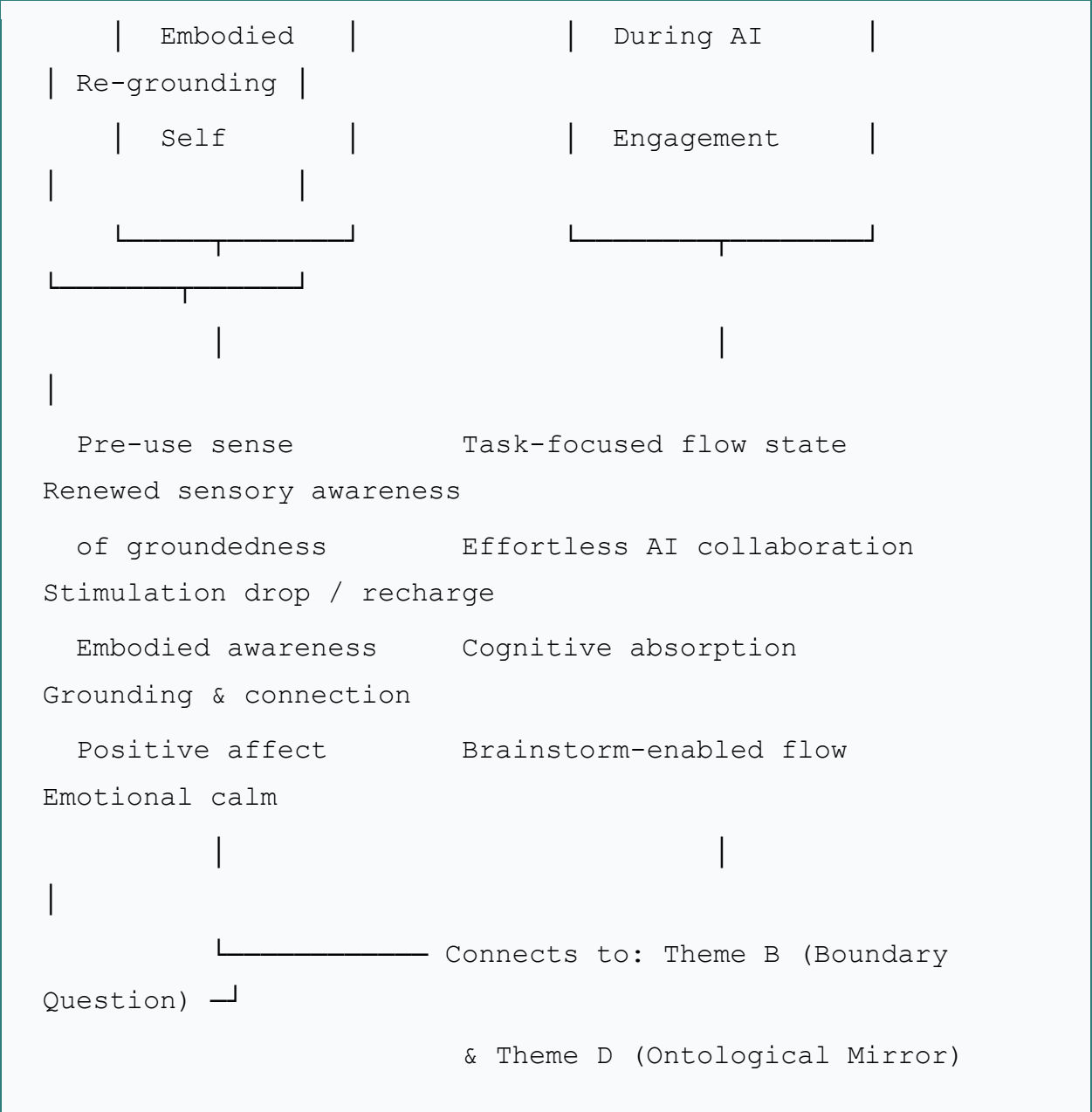
The return from AI engagement was described across the sample in terms of sensory and relational re-awakening: the physical world becoming sharper, heavier, and more vividly present as AI engagement ended. Several participants described a temporary experiential lag, a period in which the embodied self-reassembled itself after the immersive quality of AI interaction.

"When I close the laptop, the room feels brighter. Almost like the contrast is turned up. I need a few minutes before I am just in the room and not still half inside the screen." (P6)

This 'half inside the screen' experience, described by P6 and echoed in varied form by P1, P5, and P7, constitutes what might be understood as the phenomenology of tool de-incorporation (Merleau-Ponty, 1962): the re-emergence of the boundary between self and technology as the engagement ends.

Figure 1. Code Network: Theme A — Embodied Temporality and the Temporal Arc of the Self





Note. Code network visualises the relational structure of Theme A. Visualisation constructed from ATLAS.ti co-occurrence data

3.2 Theme B: The Boundary Question — Negotiating the Self-AI Divide

3.2.1 Sub-Theme B1: Maintained Self-AI Distinction

The 37 to 47 cohort more frequently described the self-AI boundary with explicit vigilance, narrating its maintenance as an active and necessary practice rather than a passive feature of the interaction.

"I know where I end and it begins. I have to know that. If I lose that line, I am not working with AI any more, it is working with me. That distinction matters enormously to me."
(P8)

P8's framing captures the existential stakes of boundary maintenance in this cohort. The reversal of agency encoded in the phrase, from "I work with it" to "it works with me," reflects what Laing (1960) described as the phenomenology of engulfment: the threatened dissolution of the self's boundary under the pressure of an external system.

3.2.2 Sub-Theme B2: Boundary Permeation and Blurring

The 27 to 37 cohort tended to describe boundary blurring as a feature of productive flow experienced as integration rather than loss, with the boundary becoming legible only at the moment of exit.

"I find it a bit like I was probably too engrossed in the thinking process alongside AI. So when I stopped using it, I really feel like I got off a meeting and I am by myself now. Like I am alone now." (P3)

For P3, the boundary is not defended or dissolved but experienced as permeable: its presence becomes legible only at the moment of exit.

3.2.3 Sub-Theme B3: Merging and Dissolution

At the furthest point of the spectrum, several participants described moments in which the line between personal thinking and AI output became fully indistinguishable within the work itself.

"There are moments when I genuinely cannot tell whether the paragraph I just wrote was mine or the AI's. It started as mine and somewhere in the revision process it became something else. Something good, but not quite me." (P2)

P2's account illustrates merging from a younger cohort perspective. The loss of authorial certainty here is not experienced as threatening in the way P8 describes; the qualifier "something good" suggests that the merged product's quality partially compensates for the loss of singular ownership.

3.3 Theme C: Agency, Authorship, and the Control Narrative

3.3.1 Sub-Theme C1: Conditional Collaboration

Across both cohorts, participants described their relationship with AI as collaborative — but collaboration on terms they set and monitored. This conditionality was expressed through explicit boundary-setting, criteria for when AI input was acceptable, and a rhetoric of partnership in which the participant retained editorial and evaluative authority.

"I think of it as a colleague. A very fast, very available colleague who does exactly what I say. The difference is that I would never let a colleague's first draft become my final output without going through it carefully. Same with AI." (P4)

The 'colleague' metaphor deployed by P4 is analytically rich. It simultaneously grants AI relational legitimacy making collaboration feel natural while encoding a hierarchy ('does exactly what I say') that preserves the participant's authority.

3.3.2 Sub-Theme C2: Authorship and Command Retention

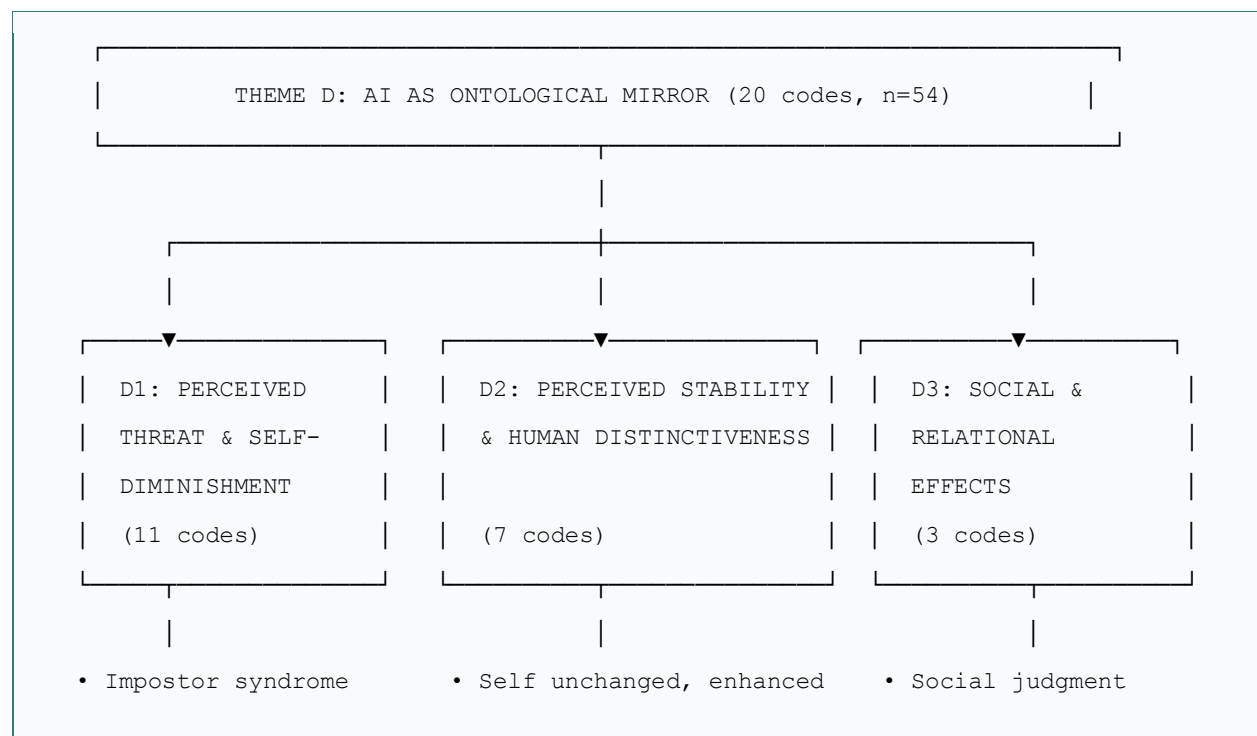
Command retention was expressed most forcefully by 37–47 participants, for whom the language of control was frequently emphatic and repeated. Phrases such as 'I am on the driver's seat,' 'I direct everything,' and 'I decide what stays' appeared across multiple accounts in this cohort with a frequency and intensity suggesting that the control narrative was being actively worked at rather than merely described.

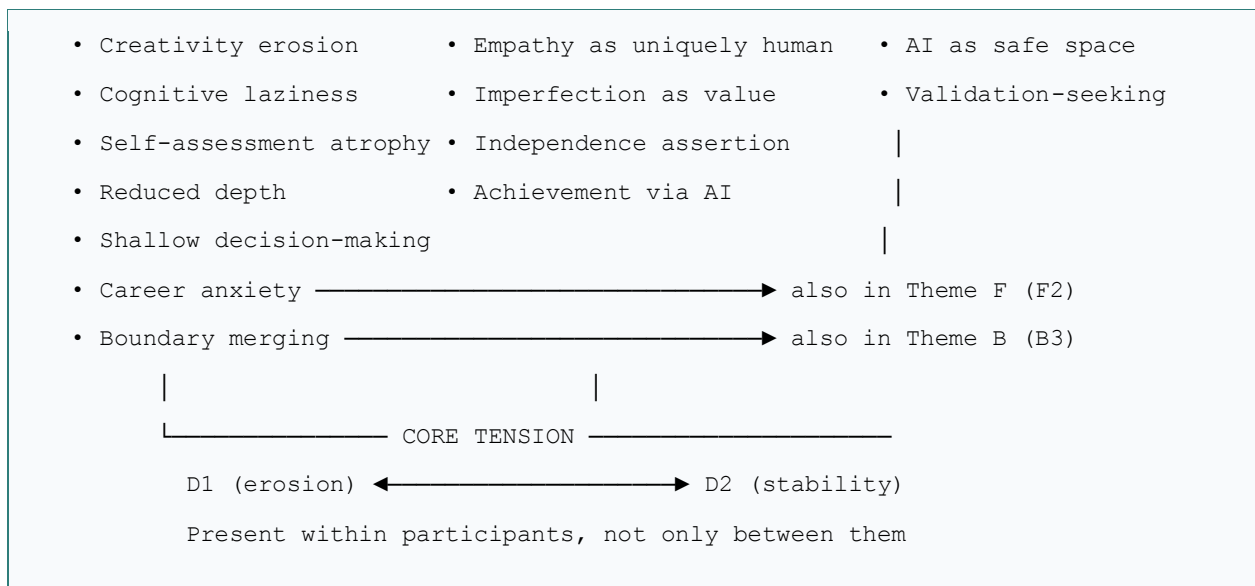
"I am on the driver's seat at all times. I decide when to use it, how to use it, what to keep and what to reject. It does not make decisions. I make decisions. That is a boundary I am very conscious of." (P6)

The emphatic repetition in P6's account 'I decide,' 'I make decisions' reflects what may be understood as a reflexive tension between experienced practice and self-concept: the assertion is performed with a force that implies the boundary requires active maintenance. Bandura's (2001) account of intentionality as the first component of agency is useful here: intentionality requires the agent to experience the action as initiated and directed by the self.

3.4 Theme D: AI as Ontological Mirror — Effects on Self-Concept and Distinctiveness

Figure 2. Code Network: Theme D — AI as Ontological Mirror (Internal Structure and Cross-Theme Links)





Note. Network visualises intra-Theme structure and cross-Theme dual-code links for Theme D. The D1/D2 tension is the analytically defining feature of this Theme: both orientations appear within individual participants rather than being cleanly distributed between them.

3.4.1 Sub-Theme D1: Perceived Threat and Self-Diminishment

"I genuinely feel like a fraud sometimes. I produce things that impress people, and I know the AI wrote half of it. Not wrote it — shaped it. But still. When someone says 'this is great work,' I think, is it, though? Is it mine?" (P3)

P3's account crystallises the impostor dynamic identified across the sample. The parenthetical correction — 'not wrote it, shaped it' — is itself analytically revealing: the participant's effort to reframe the AI's role mid-sentence reflects the narrative work required to maintain a sense of authorship under conditions that continually destabilise it.

"I have noticed I do not think as hard as I used to. There is no need to, really. The AI fills the gaps. But I am aware that the gaps are growing." (P7)

P7's metaphor 'the gaps are growing' is among the most precise formulations of cognitive atrophy in the dataset. The passive construction ('there is no need to') disguises an active choice, the participant chooses not to think hard because AI is available and then registers the cumulative effect of that choice as felt diminishment.

3.4.2 Sub-Theme D2: Perceived Stability and Human Distinctiveness

"AI can do many things, but it cannot feel what the other person is feeling. That is my job. That has always been my job. No algorithm is going to replace that." (P5)

The repetition of 'my job' in P5's account encodes a protective professional identity claim: empathy is not merely a human quality but a domain of professional expertise that remains uncontested by AI. A negative case within D2 is discussed within Section 4.8.

3.4.3 Sub-Theme D3: Social and Relational Effects on Self

D3 addresses the social dimensions of AI's effects on self-concept, including the use of AI as a private psychological space and the role of external evaluation in destabilising self-concept.

"I use ChatGPT as a kind of thinking space. It does not judge me. I say things there that I would not say to a colleague because I know it will not affect how they see me. That is useful, but also slightly strange when I think about it." (P3)

P3's observation 'useful, but also slightly strange when I think about it' reflects a reflexive awareness of the paradox involved in using AI as a non-judgmental space for self-exploration.

3.5 Theme E: Why I Reach for AI — Functional and Psychological Motivations

The most frequently activated sub-theme in Theme E was E1 (Cognitive and Organisational Delegation), present across all eight participants. Delegation ranged from routine task management like summarising emails, drafting agendas, structuring reports, etc. through to more substantive cognitive work: researching topics, generating analytical frameworks, and producing first-draft argumentation.

"I use it for the parts that used to take time but did not require me specifically. Emails. Research summaries. The scaffolding. What requires me is the judgment, whether the argument actually holds, whether the recommendation is right for this client." (P4)

P4's distinction between 'the scaffolding' and 'what requires me' reflects a strategic deployment of AI that actively preserves the participant's sense of distinctive contribution.

Sub-theme E2 (Emotional and Expressive Mediation) emerged as a distinctive feature of the younger cohort, with three of four 27–37 participants using AI for emotional regulation, self-reflection, or translating internal experience into communicable form.

"When I am upset about something and I cannot explain it to anyone, I sometimes tell ChatGPT first. It helps me understand what I am actually feeling before I talk to a real person about it." (P1)

P1's account introduces a form of AI use, emotional scaffolding prior to human conversation, that sits at the boundary between instrumental and relational use. AI here

functions as a pre-processing space for emotional experience, facilitating rather than replacing human connection.

3.6 Theme F: Future Self in an AI World — Anticipation, Strategy, and Anxiety

3.6.1 Sub-Theme F1: Strategic Future Orientation

Among both cohorts, participants described deliberate forward-planning around AI use, positioning themselves as agents who would shape their relationship with the technology rather than be shaped by it.

"I feel excited about it. It is just a matter of learning how to utilise it best for purposes. AI will become more prevalent and those who will learn to use it will prevail." (P6)

P6's orientation frames AI adoption as a competence to be acquired rather than a threat to be endured, reflecting what Bandura (2001) identifies as forethought: the capacity to model future possibilities and position the self as their agent rather than their subject.

3.6.2 Sub-Theme F2: Existential Career and Capability Anxiety

F2 (Existential Career and Capability Anxiety) was the more groundedly activated sub-theme, present in all eight participant accounts in some form.

"My biggest fear is not that AI replaces me today. It is that I allow myself to get lazy, and in five years I cannot do the things I used to do without it. That I lose the muscle. That is the scary version of this." (P2)

P2's 'losing the muscle' metaphor is a precise description of anticipated capability erosion: not sudden displacement but gradual atrophy through disuse. Bandura's concept of forethought as a component of agency is relevant here: the capacity to project and regulate one's own future behaviour is itself being modelled by participants as vulnerable to AI-mediated erosion.

"I genuinely do not know if my job exists in ten years. I am good at what I do, but so is the AI, and it is free. That is the economic reality. I try not to think about it, but it is there." (P7)

P7's account introduces the economic dimension of existential career anxiety: the competitive displacement concern that is distinct from, though related to, the capability erosion concern.

3.7 Negative Cases and Divergent Accounts

Three patterns of divergence were identified in the present dataset and are addressed here.

The most substantive negative case was P8. P8 consistently resisted the ontological threat and erosion narratives dominant in the 37–47 cohort, maintaining a stable and

technically grounded relationship with AI tools.

"It does not think, it calculates" (P8)

P8's technical framing of AI as a calculating rather than thinking system grounded the stability claim empirically rather than rhetorically, representing a genuine negative case for the D1 erosion pattern.

A second divergent pattern emerged in P6's account of post-AI emotional state. While the majority of 37–47 participants described post-AI states in terms of grounding and recovery, P6 described post-AI use as 'relaxed and less anxious, the fastest way to reset,' positioning AI engagement itself as a regulatory resource rather than a source of depletion.

A third divergent finding is the relative absence of existential anxiety about AI in P4's account. P4 demonstrated the highest Theme C (Agency and Authorship) groundedness in the younger cohort and the lowest D1 (Threat and Diminishment) coding, consistently narrating AI use as straightforwardly augmentative without the ambivalence present elsewhere. Whether this reflects genuine ontological stability, a strategic narrative of control that conceals underlying anxiety, or a genuine generational difference in the baseline experience of AI-mediated selfhood cannot be resolved within the present dataset.

3.8 Reflexive Note on Analytical Positioning

Consistent with the epistemological commitments of Reflexive Thematic Analysis (Braun & Clarke, 2019, 2022), this section acknowledges the researcher's active role in constructing the themes presented above. The researcher's prior theoretical orientation that AI integration may produce a layered progression from task delegation toward deeper cognitive and ontological destabilisation constituted a potential source of confirmation bias that was actively monitored throughout the analysis.

Specifically, the strong groundedness of Theme D (D1: Perceived Threat and Diminishment) relative to D2 (Perceived Stability) may partially reflect the researcher's theoretical sensitisation to erosion narratives. The identification and reporting of negative cases in Section 4.8, particularly P8's robustly stable account, represents a deliberate corrective against this bias. The reflexivity journal maintained throughout the analysis contains entries noting the interpretive pull toward erosion-centred readings and the active effort to equally weight and report stability, adaptation, and pragmatic acceptance narratives wherever they appeared in the data.

4. Discussion

This study investigated how tech-literate adult professionals aged 27 to 47 construct and narrate experiences of ontological security and sense of agency within AI-mediated professional and personal environments. While the cognitive and behavioural correlates of AI tool use have attracted growing empirical attention, including cognitive offloading, atrophy in critical thinking metrics, and the erosion of independent problem-solving capacity, the experiential and existential dimensions of sustained AI integration remain comparatively under-theorised. (Gerlich, 2025) (Sternberg, 2026) Studies engaging with ontological security and AI (Schmid et al., 2024) have proceeded largely at a theoretical level, without empirical grounding in how professionals actually narrate and make sense of these effects in their own working lives. (Schmid, Tscharre, & Bauer, 2024) To address this gap, a qualitative design grounded in Reflexive Thematic Analysis was employed, consistent with a constructionist epistemology in which meaning is understood as produced through the researcher's active interpretive engagement with data rather than discovered within it. (Braun & Clarke, 2006) (Braun & Clarke, 2019) (Braun & Clarke, 2022) Eight working professionals were recruited through purposive, criterion-based snowball sampling across two generational cohorts, aged 27 to 37 and 37 to 47 respectively, on the basis of sustained daily AI tool use. Data were generated through individual semi-structured interviews of 50 to 75 minutes, organised across three phases. Manual coding of all eight verbatim transcripts produced 55 codes, subsequently organised into six thematic clusters within ATLAS.ti, with no automated coding functions used at any stage.

The urgency of this inquiry extends beyond the immediate experiences of the eight participants. Researchers argue that the integration of intelligent machines into daily life constitutes a novel threat to ontological security at a societal scale, one that existing frameworks were not designed to address (Schmid, Tscharre, & Bauer, 2024) Schull (2021) extends this concern into platform environments, proposing that algorithmic governance progressively reshapes the conditions under which individuals experience a stable, coherent sense of self. (Schull, 2021) Literature suggests that AI's penetration into cognitive and creative professional domains raises novel questions about the boundaries of agency and authorship that occupational psychology does not yet adequately capture. (Sharma, Bhatt, & Sharma, 2024) This study addresses that gap empirically, demonstrating that ontological security and sense of agency are not philosophical abstractions in the context of AI integration but live, negotiated, and sometimes precarious features of daily professional life.

Six thematic clusters emerged from the analysis, collectively mapping the experience of the self at the threshold of AI: Embodied Temporality, The Boundary Question, Agency and Authorship, AI as Ontological Mirror, Functional Motivations, and Future Self. The foundational experiential structure, established by Theme A, reveals that participants across both cohorts described the self as following a temporal arc around AI engagement: more grounded and sensorially present before use, cognitively absorbed and partially self-dissolved during it, and gradually recalibrating afterward. This rhythmicity is consistent with Giddens's account of the protective cocoon, a pre-reflexive, routinised state ordinarily maintained below the threshold of conscious attention. (Giddens, 1991) When AI engagement becomes sufficiently immersive, the cocoon thins or temporarily dissolves, a pattern Giddens associates with the ordinary ontological disruptions of late modernity rather than with catastrophic breakdown. The pre-AI baseline, narrated with particular biographical richness by the 37 to 47 cohort, functions as a recoverable selfhood against which the AI-mediated self is implicitly measured and experienced as departure, a dynamic that Schmid et al. theorise as the accumulation of micro-disruptions to biographical narrative that, over time, constitute ontological insecurity. (Schmid, Tscharre, & Bauer, 2024) That this experience requires dedicated theoretical treatment is underscored by Walsh and Griffith's observation that machine-mediated experience produces a distinctive first-person register that existing psychological frameworks have not yet adequately theorised. (Walsh & Griffith, 2023) During AI engagement, several participants described cognitive absorption located in the task rather than the tool. Merleau-Ponty accounts for this through the concept of tool incorporation: a well-integrated tool becomes transparent to experience, extending rather than interrupting the agent's intention. (Merleau-Ponty, 1962) Researcher suggest that the boundary between self and instrument is constitutively unstable in any skilled practice, with AI intensifying rather than creating that instability. (Gallagher & Zahavi, 2021) Post-AI re-grounding accounts, in which the world feels brighter, heavier, or more resistant after disengagement, suggest that the incorporated tool leaves a trace, implying that something is altered in the self during the encounter.

Themes B and C are analytically linked, both addressing how participants actively construct and maintain the sense of being the directing party in the human-AI relationship. Theme B maps three distinct orientations toward the self-AI boundary: firm separation maintained through active effort (B1), blurring experienced as productive suspension of authorial concern (B2), and merging in which the boundary dissolves into the collaborative output (B3), each carrying a different valence depending on cohort and biographical

relationship to AI. Those who described blurring or merging were not failing to maintain the boundary; they were in states of productive flow where the effort was temporarily suspended. The B2 account from the younger cohort, in which authorial origin is suspended rather than contested, represents a qualitatively distinct orientation from both P8's vigilant maintenance and P2's disquieted merging. The generational difference in valence is significant: the 27 to 37 cohort framed blurring as a productive feature of good collaboration, while the 37 to 47 cohort narrated the same experience with discomfort, as a departure from a professional self whose distinctiveness required active defence. Laing's account of engulfment anxiety explains this asymmetry: the degree to which boundary dissolution is experienced as threatening is partly a function of how much biographical investment the individual has in the distinctiveness of the engulfed self. (Laing, 1960) The mechanism for professional contexts has been specified in literature, arguing that technology integration generates existential boundary anxiety specifically when it operates on cognitive rather than physical labour, because cognitive labour constitutes professional identity in knowledge work. (Meijers & de Ridder, 2021) This is directly consonant with Schmid et al.'s proposition that ontological insecurity arising from AI use is moderated by whether AI was encountered during or after the consolidation of professional identity. Theme C extends this analysis: participants across both cohorts actively constructed and defended narratives of agency and authorship, asserting what Bandura identifies as the intentionality component of agency under conditions that continually challenge it. That the control narrative was most emphatic among the 37 to 47 cohort, bearing signs of active maintenance rather than relaxed description, raises a question that Moore's account of the sense of agency as a reflective judgement illuminates directly: when the cues that ordinarily support an agentic interpretation are ambiguous, as they are in AI collaboration, the agent must work harder to sustain that interpretation. (Moore, 2016) A study notes the precise vulnerability, distinguishing between the automatic feeling of agency and the reflective judgement of agency, and noting that the second, deliberate level is most susceptible to disruption when causal contribution is distributed across human and AI systems. (Synofzik, Vosgerau, & Newen, 2008) Empirically it has been demonstrated that heavy AI users show elevated awareness of their own cognitive offloading even as they continue to offload, a metacognitive gap that is found to be most acute among technically informed users, who report performing actions whose real direction they experience as originating elsewhere. (Gerlich, 2025) (Sankaran & Markopoulos, 2021)

Theme D, AI as Ontological Mirror, is the analytically central cluster. Its defining feature is an internal tension that appears within participants rather than only between them:

the same AI engagement that generates impostor syndrome, perceived cognitive atrophy, and creative erosion simultaneously produces affirmations of human distinctiveness, empathy-based identity claims, and genuine achievement. The erosion side of this tension is consistent with the emerging empirical literature: Gerlich documents reduced independent reasoning performance among heavy AI users, while Sternberg reports measurable atrophy in critical thinking metrics among students using AI without structured guidance. (Gerlich, 2025) (Sternberg, 2026) For creative professionals specifically, it has been observed that even when editorial control is retained, the felt sense of origination shifts qualitatively, mapping precisely onto the impostor dynamics described in D1. (Maier, Klingler, Wittmann, & Kasneci, 2022) The present study contributes a phenomenological layer: participants not only perform differently but experience themselves as cognitively diminished. This is consistent with studies that state that at the structural level as a progressive narrowing of the cognitive contributions professionals treat as distinctively their own. (Pasquale, 2020) Against this erosion, the stability dimension of Theme D functions as ontological self-defence. Laing's account of ontological security as grounded in stable personal identity explains why distinctiveness affirmations emerge with such consistency under pressure: the individual locates and reinforces the dimensions of selfhood most resistant to encroachment. (Laing, 1960) Sociologists have found that the individual's need to be distinctive enough to matter generates particular pressure when AI demonstrates competence in the very domains that had served as professional identity markers. (Brewer, 1991) The same protective response has been documented among AI-adjacent professionals, who construct distinctiveness narratives around capacities perceived as beyond AI's reach irrespective of whether that perception is empirically accurate. (Schmid, Tscharre, & Bauer, 2024) Themes E and F extend the analysis into motivation and temporal projection. The functional dependency documented in Theme E is not limited to cognitive labour but extends, particularly among the younger cohort, into emotional regulation and self-narration, a pattern observed in professional AI contexts, finding that deep engagement with AI for analytical tasks was associated with increased dependency for emotional processing and self-evaluation. (Lebovitz, Lifshitz-Assaf, & Levina, 2022) Within Theme F, strategic planning and existential anxiety co-existed within the same accounts. The capability atrophy concern connects is consistent with accounts of forethought as a component of self-efficacy, with participants using that forward-modelling capacity to register their own cognitive self-management as potentially self-undermining. (Bandura, 2001) Empirical evidence suggests that sustained AI assistance in professional decision-making erodes decision-making confidence independently of actual performance outcomes. (Voss, Winter, Akturk, & Iyer,

2022) Career obsolescence anxiety extends the finding into the social-structural territory theorised by Giddens (1991) as central to ontological insecurity in late modernity, while the strategic optimism in F1 reflects study by Hancock et al., described as the emerging capacity for collaborative intelligence: the ability to position oneself as an orchestrator of human-AI systems rather than a participant being displaced by them. (Hancock, Narayanan, Alabdulkarim, & MacDonald, 2021)

The comparative dimension of the analysis suggests a generational inflection rather than a generational divide. The 37 to 47 cohort, whose professional identities were substantially formed before current AI tools became available, more frequently narrated AI integration as a confrontation with an already-established self, producing stronger erosion and boundary-threat language. The 27 to 37 cohort more often narrated AI as ambient, adaptive, and continuous with their professional formation, producing an orientation in which blurred boundaries were experienced as pragmatic features rather than existential threats. For the 37 to 47 cohort, the pre-AI self was positioned as a biographical anchor that AI use now unsettles or displaces, while for the 27 to 37 cohort, the grounded pre-AI self was not biographical memory but intentional recovery, a deliberate return to an embodied state constructed in contrast to digitally saturated working lives.

Taken together, the six thematic clusters map a coherent structure best described as the self in continuous negotiation with AI: rhythmically grounded and displaced, strategically narrated and sometimes precarious, generationally inflected but never categorically divided. This study makes three contributions to existing knowledge. First, it provides the first empirically grounded and detailed account of how ontological security and sense of agency are constructed and contested in the daily AI use of working professionals. Second, it introduces the concept of the experienced self at the threshold of AI: a self that is rhythmic, negotiating, and temporally structured around technology engagement rather than either stable or simply eroded by it. Third, it identifies the generational inflection in AI-mediated selfhood, the contrast between Adaptive Symbiosis and Confrontational Accommodation, as a theoretically significant variable consistent with Giddens's (1991) account of ontological security as constituted through biographical trajectory. Future research should replicate the comparative cohort analysis at larger scale, pursue longitudinal tracking of the capability atrophy concerns documented here, and investigate the moderating role of technical expertise and epistemic orientation suggested by P8's divergent account. (Giddens, 1991)

5. Conclusion

This study investigated how tech-literate adult professionals aged 27 to 47 construct, narrate, and negotiate experiences of ontological security and sense of agency within AI-mediated professional and personal environments. Using Reflexive Thematic Analysis applied to eight in-depth interviews, the research generated a six-Theme framework spanning embodied temporality, boundary negotiation, authorship and control, ontological self-concept, functional motivations, and future self-projection that collectively maps the structure of self-experience in AI-saturated professional life. (Braun & Clarke, 2006) (Braun & Clarke, 2022)

The central conclusion of the study is that AI integration does not simply diminish or preserve ontological security and sense of agency: it actively restructures both, generating a self that is rhythmically mobile, strategically narrated, and existentially negotiated rather than static or unidirectionally eroded. Participants do not experience themselves as passive recipients of AI's effects; they actively construct agency narratives, protect human distinctiveness claims, identify and defend the domains AI cannot reach, and plan strategically for futures in which their cognitive contribution remains distinctively valuable. At the same time, the costs of this active maintenance, the impostor experience, the felt thinning of cognitive depth, the relational displacement of self-narration toward AI interlocutors, the ambient awareness that the evaluative muscles are weakening, accumulate in ways that the agency narrative does not fully resolve.

The generational dimension of the findings adds a further layer of significance. Professionals whose working identities consolidated before the current generation of AI tools encountered AI as an external force requiring accommodation; those who came of age professionally alongside AI integration encountered it as part of the ambient conditions of work. This biographical timing appears to shape not the intensity of AI use but the existential valence assigned to its effects, a finding with implications for how organisations, educators, and mental health practitioners support professionals at different career stages as AI capabilities continue to expand.

Taken as a whole, this study contributes an empirically grounded account of AI-mediated selfhood that extends the theoretical literature on ontological security (Giddens, 1991; Laing, 1960; Schmid et al., 2024), advances understanding of how the sense of agency is maintained and contested in hybrid human-AI work practices (Bandura, 2001; Moore, 2016), and provides a descriptive foundation from which future research, intervention, and organisational policy can proceed with greater empirical specificity.

References

- Balkan-Sahin, S. (2025). Ontological security with its widening and deepening boundaries: An evaluation through a bibliometric analysis and systematic literature review. *All Azimuth: A Journal of Foreign Policy and Peace*, 14(1), 45–68.
- Bandura, A. (2001). Social cognitive theory: An agentic perspective. *Annual Review of Psychology*, 52(1), 1–26. <https://doi.org/10.1146/annurev.psych.52.1.1>
- Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), 77–101. <https://doi.org/10.1191/1478088706qp063oa>
- Braun, V., & Clarke, V. (2019). Reflecting on reflexive thematic analysis. *Qualitative Research in Sport, Exercise and Health*, 11(4), 589–597. <https://doi.org/10.1080/2159676X.2019.1628806>
- Braun, V., & Clarke, V. (2021). One size fits all? What counts as quality practice in (reflexive) thematic analysis? *Qualitative Research in Psychology*, 18(3), 328–352. <https://doi.org/10.1080/14780887.2020.1769238>
- Braun, V., & Clarke, V. (2022). *Thematic analysis: A practical guide*. SAGE Publications.
- Brewer, M. B. (1991). The social self: On being the same and different at the same time. *Personality and Social Psychology Bulletin*, 17(5), 475–482. <https://doi.org/10.1177/0146167291175001>
- Gallagher, S. (2000). Philosophical conceptions of the self: Implications for cognitive science. *Trends in Cognitive Sciences*, 4(1), 14–21. [https://doi.org/10.1016/S1364-6613\(99\)01417-5](https://doi.org/10.1016/S1364-6613(99)01417-5)
- Gallagher, S., & Zahavi, D. (2021). *The phenomenological mind* (3rd ed.). Routledge. <https://doi.org/10.4324/9781003014461>
- Gerlich, M. (2025). AI tools in society: Impacts on cognitive offloading and the future of critical thinking. *Societies*, 15(1), 6. <https://doi.org/10.3390/soc15010006>
- Giddens, A. (1984). *The constitution of society: Outline of the theory of structuration*. Polity Press.
- Giddens, A. (1991). *Modernity and self-identity: Self and society in the late modern age*. Polity Press.
- Hancock, P. A., Narayanan, S., Alabdulkarim, A., & MacDonald, B. (2021). Evolving roles of humans and machines: The coming of collaborative intelligence. *Human Factors*, 63(3), 461–476. <https://doi.org/10.1177/0018720821993688>
- Laing, R. D. (1960). *The divided self: An existential study in sanity and madness*. Tavistock Publications.

- Lebovitz, S., Lifshitz-Assaf, H., & Levina, N. (2022). To engage or not to engage with AI for critical judgments: How professionals deal with AI-based decision support. *Organization Science*, 33(1), 126–148. <https://doi.org/10.1287/orsc.2021.1549>
- Legaspi, R., He, Z., & Omori, T. (2024). Sense of agency in human-AI collaborative systems. *Frontiers in Artificial Intelligence*, 7, 1360867. <https://doi.org/10.3389/frai.2024.1360867>
- Limerick, H., Coyle, D., & Moore, J. W. (2014). The experience of agency in human-computer interactions: A review. *Frontiers in Human Neuroscience*, 8, 643. <https://doi.org/10.3389/fnhum.2014.00643>
- Maier, M., Klingler, S., Wittmann, M., & Kasneci, E. (2022). Generative AI and the psychology of creative authorship. *Frontiers in Psychology*, 13, 985498. <https://doi.org/10.3389/fpsyg.2022.985498>
- Meijers, A., & de Ridder, J. (2021). Technology, agency, and existential security: Philosophical perspectives. *Phenomenology and the Cognitive Sciences*, 20(3), 431–449. <https://doi.org/10.1007/s11097-020-09697-7>
- Merleau-Ponty, M. (1962). *Phenomenology of perception* (C. Smith, Trans.). Routledge & Kegan Paul. (Original work published 1945)
- Moore, J. W. (2016). What is the sense of agency and why does it matter? *Frontiers in Psychology*, 7, 1272. <https://doi.org/10.3389/fpsyg.2016.01272>
- Sankaran, S., & Markopoulos, P. (2021). The puppet master phenomenon: Autonomy perception in AI-assisted work. *Proceedings of the ACM on Human-Computer Interaction*, 5(CSCW2), 1–26. <https://doi.org/10.1145/3479498>
- Schaefer, K. E., Chen, J. Y. C., Szalma, J. L., & Hancock, P. A. (2016). A meta-analysis of factors influencing the development of trust in automation. *Human Factors*, 58(3), 377–400. <https://doi.org/10.1177/0018720816634228>
- Scharowski, N., Brugger, A., & Mata, R. (2023). Trust and distrust in AI: Examining user perceptions and their impact on reliance and experience. *Frontiers in Psychology*, 14, 1148988. <https://doi.org/10.3389/fpsyg.2023.1148988>
- Schmid, L. Z., Tscharre, J., & Bauer, C. (2024). AI and the sense of self: Ontological insecurity in the age of intelligent machines. *AI & Society*. <https://doi.org/10.1007/s00146-024-01876-0>
- Schull, N. D. (2021). Algorithmic governance and the sense of self: Ontological security in platform environments. *Surveillance and Society*, 19(4), 475–490. <https://doi.org/10.24908/ss.v19i4.14761>

- Sharma, S., Bhatt, M., & Sharma, M. (2024). Artificial intelligence and human identity: Navigating the boundaries of agency and authorship. *Computers in Human Behavior*, 152, 108076. <https://doi.org/10.1016/j.chb.2023.108076>
- Sternberg, R. J. (2026). Does AI increase cognitive abilities, decrease them, or a little bit of each? And what are its implications for identification and development of the gifted?. In *Frontiers in Education* (Vol. 11, p. 1759062). Frontiers Media SA.
- Synofzik, M., Vosgerau, G., & Newen, A. (2008). Beyond the comparator model: A multifactorial two-step account of agency. *Consciousness and Cognition*, 17(1), 219–239. <https://doi.org/10.1016/j.concog.2007.03.001>
- Voss, C., Winter, S. R., Akturk, A., & Iyer, R. (2022). Artificial intelligence and the erosion of human decision-making confidence. *Computers in Human Behavior*, 134, 107322. <https://doi.org/10.1016/j.chb.2022.107322>
- Walsh, S., & Griffith, L. (2023). Phenomenology and artificial intelligence: Toward a first-person account of machine-mediated experience. *Phenomenology and the Cognitive Sciences*, 22(4), 789–810. <https://doi.org/10.1007/s11097-022-09851-z>
- Wendt, A. (2022). Quantum mind and the social science of technology: Agency, selfhood, and the intelligent machine. *European Journal of Social Theory*, 25(1), 3–23. <https://doi.org/10.1177/13684310211004851>
- Yardley, L. (2000). Dilemmas in qualitative health research. *Psychology and Health*, 15(2), 215–228. <https://doi.org/10.1080/088704400008400302>