

The Agentic AI Adoption Paradox: Barriers to Enterprise Deployment in the Software Industry

Ricky Tanudjaja^{1*}, Sahat Hutajulu²¹² Bandung Institute of Technology, West Java, Indonesia

Abstract

Agentic AI adoption in Indonesia is stalling despite strong executive interest and substantial investment. No empirical study has investigated Agentic AI adoption within Indonesia's software development industry. Existing empirical research has primarily examined Generative AI, with limited attention to the distinct characteristics of autonomous Agentic AI systems. Drawing on UTAUT, this qualitative case study investigates this adoption paradox and extend adoption theory by repositioning ethical concern from a peripheral barrier to an endogenous structural factor. Sixteen semi-structured interviews and one focus group discussion were conducted with practitioners across executive, managerial, and technical tiers using purposive sampling to capture diverse perspectives. Data collection continued until thematic saturation was reached. The data were analysed thematically using the Gioia method, with the assistance of Ligre software to identify and organise emerging patterns and themes. The study identifies three key barriers to Agentic AI deployment in Indonesia. First, output unreliability and qualification limitation hinder its adoption in production environments. Second, Indonesia's relatively low labour costs reduce the economic incentive for investing in Agentic AI infrastructure. Third, while Agentic AI offers significant operational benefits, its capabilities also heighten ethical concerns. We suggest that effective governance and monitoring require Agentic AI to have its ethical guardrails embedded at design time rather than retrofitted post-deployment, have transparent and explainable audit logging, and have its autonomy be calibrated based on its reliability and task risk exposure.

KEYWORDS

agentic AI; technology adoption; change management; software development

Introduction

The field of Artificial Intelligence (AI) is evolving from Generative AI, which primarily produces outputs in response to human prompts or inputs, to Agentic AI, in which systems actively perceive, reason, and take autonomous actions to achieve goals and influence their environment (V et al., 2026). This transformation is significant for the software engineering industry, whose growth has often been constrained by shortages of skilled developers, limited coding capacity, and the complexity of managing large-scale projects (Charlotte Relyea et al., 2025). Agentic AI is expected to address these challenges through its ability to operate autonomously, thereby enhancing efficiency, productivity, and human-machine collaboration (Murugesan, 2025). However, despite its transformative potential, a substantial adoption gap remains. This study refers to this phenomenon as the Agentic AI Adoption Paradox: a condition in which strong executive support and significant investment in Agentic AI coexist with persistently low rates of successful production deployment and limited realisation of business value.

Emerging evidence suggests that this gap is not merely a technological issue but also reflects differences in contextual and organisational readiness. Yang et al. (2025) provide one of the first large-scale empirical studies of AI agent usage,

showing that adoption is concentrated among early adopters, professionals in countries with higher GDP per capita and educational attainment, and workers in digitally intensive and knowledge-based sectors such as technology, academia, finance, marketing, and entrepreneurship. These findings suggest that Agentic AI adoption depends on enabling conditions beyond technology availability, including environmental digital maturity and user capabilities.

This perspective is particularly relevant in Indonesia, where organisational ambition appears to outpace implementation capability. Although 61 percent of Indonesian companies are already deploying AI agents, only 27 percent report achieving the expected Return on Investment (ROI) (IBM, 2025). Similarly, while 92 percent of business leaders in Indonesia acknowledge that AI adoption is essential for maintaining competitiveness, nearly half (48 percent) report lacking a clear AI implementation strategy (Microsoft & LinkedIn, 2024). Together, these figures reveal a sector caught between aspiration and readiness.

These adoption difficulties are compounded by the inherent ethical risks that Agentic AI introduces. Agentic AI raises serious concerns around trust, privacy, security, bias, and workforce displacement (Murugesan, 2025). Gartner (2025) predicts that by the end of 2027, more than 40 percent of Agentic AI projects will be cancelled due to escalating costs, unclear business value, or inadequate risk controls. Without robust guardrails, organisations expose themselves to operational disruption, governance failures, and reputational damage (Aubret & Florentin, 2025).

This research explores the Agentic AI Adoption through qualitative case study research design with Software Engineering (SE) professionals across managerial tiers in Indonesia to investigate the following research questions: RQ1: How does performance expectancy of Agentic AI vary across different organisational tiers, and how does this variation contribute to adoption intention within the Software Development Life Cycle (SDLC)?

RQ2: What distinct cultural, structural, and technological barriers obstruct software organisations from successfully transitioning Agentic AI from pilot phases to enterprise-grade production deployment?

RQ3: What ethical risks are introduced by deploying Agentic AI, and what structured governance frameworks are required to mitigate these risks?

Core Constructs and Propositions

The study of technology adoption is among the most extensively researched domains in information systems and management science. One of the most widely applied frameworks is the Unified Theory of Acceptance and Use of Technology (UTAUT), developed by Venkatesh et al. (2003), which integrates the Theory of Reasoned Action (TRA) (Fishbein, 1976) and the Technology Acceptance Model (TAM) (Davis, 1989). In this thesis, the four core UTAUT constructs—Performance Expectancy (PE), Effort Expectancy (EE), Social Influence (SI), and Facilitating Conditions (FC)—are employed not as hypotheses in a predictive statistical model, but as interpretive analytical lenses for examining technology adoption behaviour.

This approach aligns closely with the theory generation approach by Russo (2024) in studying Generative AI adoption within software engineering. The multi-faceted nature of modern AI adoption demands a comprehensive capture of key influencing factors (Russo, 2024). While Russo achieves this multi-dimensional perspective by integrating the TAM, Diffusion of Innovations (DOI), and Social Cognitive Theory (SCT), this thesis achieves a

similarly holistic view by leveraging the integrated nature of UTAUT.

The novelty of this study is that we include Ethical Concerns (EC) a core factor that directly influences deployment decisions. Classical technology adoption frameworks treat EC as an external boundary condition. Murugesan (2025) identifies ethics as core structural factor in responsible adoption, given the technology's capacity to take consequential, irreversible actions without human verification. Therefore, this study incorporates Ethical Concern (EC) as an additional construct. EC captures developers' and managers' perceptions of accountability, transparency, fairness, and workforce impact associated with autonomous agentic systems.

Drawing on the foregoing theoretical synthesis and supported by empirical evidence from the existing studies, this thesis advances the following theoretically grounded propositions:

P1 – Practitioners who perceive Agentic AI as meaningfully enhancing job performance and development quality (Performance Expectancy) are expected to express stronger adoption intent.

P2 – Practitioners who perceive lower cognitive and technical effort in integrating Agentic AI into existing workflows (Effort Expectancy) are expected to express stronger adoption intent.

P3 – Peer, managerial, and industry endorsement of Agentic AI (Social Influence) is expected to positively shape adoption intent.

P4 – The availability of training, governance frameworks, and managerial support (Facilitating Conditions) is expected to predict whether adoption intent successfully converts to production deployment.

P5 – Ethical Concern, encompassing accountability, transparency, fairness, workforce displacement anxiety, and security vulnerability, function as a factor of adoption-to-deployment conversion.

Methods

Research Type

This study employs a qualitative case study research design, selected for its suitability in understanding complex social phenomena within real-life business environments and its ability to support theory development where existing theories remain inadequate (Iiris Aaltio & Pia Heilmann, 2010; Yin, 2003). This design is appropriate given the exploratory nature of the research problem, the emergent character of the Agentic AI adoption challenge, and the absence of prior empirical work in the Indonesian context. The aim of the study is to examine the Agentic AI adoption paradox at multiple levels of analysis corresponding to distinct organisational tiers (executive, managerial, and technical contributor) and four mutually reinforcing forces (technical, cultural, structural, and ethical).

Population and Informant

We began participant recruitment by using purposive sampling to capture diverse perspectives of Agentic AI adoption from their respective organisational roles (Patton, 2015). The initial selection criteria included software engineering professionals who possessed AI-related expertise, acquired through direct experience using AI in their organisational role. This ensured that participants could provide informed, experience-based insights into the organisational adoption of Agentic AI.

We conducted sixteen interviews between November 2025 and January 2026 whose experience in the software

industry varied from 2 to 37 years. The target sample of

sixteen individual interviews was determined in light of the

Table 1. Semi-Structured Interview Respondents

ID	Organisational Position	Years of Experience	Organisational Tier
R1	Co-founder & Director of Technology	5	Strategic
R2	Head of Software Engineer & Cloud Operations	15	
R3	Chief Technology Officer	27	
R4	CEO/Chief Executive Officer	37	
R5	Product Manager	6	Operational
R6	Data Analyst Risk Management	10	
R7	Post-Sales	10	
R8	Business Developer / Software Architect	22	
R9	Software Engineer	2	Practitioner
R10	Software Engineer	2	
R11	Software Engineer	3	
R12	Software Engineer	4	
R13	Software Engineer	11	
R14	Software Engineer	20	
R15	Software Quality Engineer	8	
R16	Digital Transformation Specialist (Automation)	3	

Table 2. Focus-Group Discussion Participants

ID	Organisational Position	Years of Experience
FG-1	Software Engineer	2
FG-2	Software Engineer	2
FG-3	Software Engineer	2
FG-4	Software Engineer	2
FG-5	Software Engineer	4
FG-6	Software Engineer	4
FG-7	Software Engineer	4
FG-8	Cloud and DevOps Engineer	4

principle of theoretical saturation — the point at which additional data no longer generate new conceptual categories or meaningfully extend existing ones (Glaser & Strauss, 2017). In practice, qualitative case studies of this scope typically achieve saturation within twelve interviews according to Guest et al., (2006).

The participants were drawn from three organisational tiers as described in Table 1 along with the participant's years of experience in SE and their job positions. The strategic tier included Chief Executive Officers, Chief Technology Officers, and equivalent senior decision-makers (n = 4). The operational tier included Project Managers, Product Managers, and equivalent middle-management roles (n = 4). The practitioner tier comprised Software Engineers, Quality Assurance Engineers, and related roles (n = 8).

One focus group discussion (FGD) as described in Table 2 was additionally conducted with eight participants — six software engineers, one DevOps engineer, and one project manager — following a live demonstration of Agentic AI capabilities. The FGD was designed to elicit the social construction of meaning in an interactive context and to capture the interpersonal adoption dynamics that individual interviews cannot access.

Data Collection Procedures

Interviews were conducted face-to-face or through virtual platforms, lasting approximately 20–60 minutes each. An interview protocol was developed from the theoretical framework constructs with open-ended questions exploring Agentic AI adoption intentions, deployment barriers, and ethical concerns. The FGD was conducted with similar question protocol but with emphasis on capturing shared meanings, emerging consensus, and productive disagreements. Both sessions were conducted primarily in a mix of Bahasa Indonesia and English,

reflecting participants' preferred language of professional expression. All interviews were audio-recorded with participant consent and transcribed verbatim in the original language of the interview prior to analysis. Initial coding (first-order codes) was conducted on the original-language transcripts to preserve semantic nuance and avoid premature loss of meaning through early translation (van Nes et al., 2010). Quotations selected for inclusion in Table 3 were subsequently translated into English by the researcher, who is bilingual in Bahasa Indonesia and English.

Data Analysis

Data were analysed using the Gioia Methodology (Gioia et al., 2013), through a transparent three-tiered coding structure. Analysis proceeded via three sequential coding stages: First-Order Concepts (open coding), in which raw data were segmented using in-vivo labels closely reflecting participants' own language; Second-Order Themes (axial coding), in which codes were grouped into categories and subcategories through constant comparative analysis; and Aggregate Dimensions (selective coding), in which overarching themes were abstracted across the complete dataset.

Data management was supported by Ligre, a web-based Computer-Assisted Qualitative Data Analysis Software (CAQDAS) platform. All interviews and discussion transcripts were imported, coded, and maintained within a structured codebook ensuring traceability of all interpretive decisions. The coded summaries were compared across cases, and the findings were repeatedly reviewed, grouped, and refined into broader categories until no new patterns emerged following the comparison approach of Glaser & Strauss (2017).

Trustworthiness

To establish rigor in accordance with qualitative case study conventions, this study addressed the four trustworthiness criteria (credibility, transferability, dependability, confirmability). Credibility was assured through data source triangulation across three organisational tiers (strategic, operational, practitioner) and across data collection methods (interviews and focus group discussion), enabling findings to be cross-checked for convergence and divergence. Transferability was supported through thick description of participant roles, years of experience, and verbatim quotations retained in [Table 3](#), together with a transparent account of the criteria and process underlying participant selection through purposive sampling. Dependability was addressed through maintenance of a detailed audit trail, supported by Ligre which preserved transcripts, codes, and codebook revisions within a traceable digital environment documenting the coding process from first-order codes through second-order themes to aggregate dimensions, consistent with the Gioia methodology. Confirmability was pursued through the retention of direct participant quotations in the [Table 3](#) to ground interpretive claims in raw data rather than researcher inference alone.

Ethical Approval

All participants provided informed consent before participating in the study.

Result and Discussion

This chapter presents the empirical findings and theoretical discussion regarding the adoption of Agentic AI within software development firms in Indonesia. The Gioia coding analysis in [Table 3](#) is built upon qualitative data extracted from 16 individual semi-structured interviews (R1–R16) and one Focus Group Discussion (FGD).

P1 – Job Performance and Software Quality Expectancy

Productivity and speed enhancement. The most consistently cited benefit was time compression. As one FGD participant (FG-6) explained, Agentic AI enables ‘basic (coding) functions to be done within a couple of minutes, where I think people usually take probably like at least like a day-ish,’ while also being able to ‘generate (the documents) within minutes.’ These observations are consistent with the reported 40% increase in the productivity among midlevel professionals from using ChatGPT (Noy & Zhang, 2023).

Elevated Code Quality. Several informants observed that agentic AI raises baseline quality of software outputs. R2 noted that agent assistance helped to reduce human error, while R14 highlighted its ability to analyse and learn from vast amounts of source code, allowing them to provide recommendations based on best practices. This suggests that Agentic AI contributes to code quality.

Cognitive Resource Reallocation. Repetitive coding, test case generation, documentation, and requirement verification were identified as automatable by AI Agent, freeing engineers to focus on higher-value work. R12 reported about an ‘increase in interesting, innovative jobs’ by having repetitive activities are automated by AI, while R9 also noted that Agentic AI adoption creates opportunities for employees to concentrate on more productive and strategically important activities.

P2 – Effort Expectancy: The First-Mile / Last-Mile Adoption Gap

Integration–Deployment Gap. R14 explained that ‘integrating (Agentic AI) is easy because we can just talk to the agents ... but making sure they don’t produce errors is a headache,’ showing that basic use is simple, but production use is much harder. So, while prompting and generating outputs is easy, ensuring accuracy, reliability, and safe use in real environments requires stronger testing, validation, and AI knowledge.

Hallucination and Output Unreliability.

Hallucination is one of the primary challenges with LLM-based agents for software development (Rasheed et al., 2026). Agentic AI systems may produce incorrect or misleading code in production systems, necessitating output verification through appropriate validation and testing practices. Participants consistently noted that AI-generated outputs cannot be fully trusted without oversight. R14 highlighted this issue directly, stating that ‘Agentic AI also has a lot of hallucinations,’ while R9 added that the overall ‘quality may not be good yet.’ Collectively, these findings suggest a practical upper bound on Agentic AI reliability in applied settings, reinforcing the continued need for human oversight in critical workflows.

Context Provision and Management.

Participants highlighted that effective use of Agentic AI depends heavily on how well users provide and manage context. FGD participants noted that reliable outputs require detailed and well-structured input, including relevant data and constraints. R10 further explained that the main challenge is not simply prompting for what to build, but communicating the underlying ‘why,’ including architectural reasoning, past decisions, and stakeholder constraints that are often not fully documented. This aligns with literature showing that industrial codebases and related documents can exceed the information capacity of current models, making context management a critical limitation in Agentic AI real-world deployment (Apostolou et al., 2026).

P3 – Social Influence: Leadership Directives and Testimony

Top-Down Leadership Mandate. Consistent with Asian high power-distance culture, Agentic AI adoption in Indonesia is primarily initiated through executive leadership. R2 explained that once senior management announced the ambition of becoming an AI-driven company, the directive was rapidly cascaded throughout the organisation. This suggests that top-down leadership support serves as a key catalyst for AI adoption, creating both organisational commitment and employee intent to engage with AI initiatives.

Social Proof and Peer Validation. Participants highlighted the importance of social proof and peer validation in encouraging Agentic AI adoption. R2 noted that observing colleagues achieve positive outcomes with AI often sparked interest among team members, explaining that ‘my experience in the team was probably looking at success stories first. “Oh, I can use AI like this.” And that makes people interested in using AI too.’ At the same time, R14 believed that social influence alone is a non-factor and that argued that ‘the internal drive plays a bigger role,’ highlighting the importance of individual motivation and self-efficacy in determining whether employees actively engage with Agentic AI. Together, these findings indicate that while a culture of knowledge sharing and mimetic pressure can stimulate interest in Agentic AI, sustained adoption also depends on employees’ confidence and willingness to explore the technology themselves.

Fear of Non-Adoption. On the other hand, a stronger mechanism of fear of non-adoption emerges, where R13 expressed this at the individual level: ‘if others are using AI

and you are not, you risk being left behind'. This suggests that adoption is influenced also by perceptions of competitive threat, extending traditional social influence mechanisms toward a fear of non-adoption driven by concerns over relative disadvantage.

P4 – Facilitating Conditions: Infrastructure Provision and Governance Gap

Training Deficit. A recurring barrier identified by participants concerns uneven levels of AI literacy and familiarity, which creates a bifurcation between individuals who actively experiment with AI tools and those who do not. These gaps manifest in practical usage challenges, particularly in prompting competence, context guidance, and the ability to critically evaluate outputs. R7 described this as stemming from a 'fear of the unknown' and a lack of knowledge, while R5 expressed concern about the difficulty of identifying bias in AI-generated outputs.

Cost and Infrastructure Constraints.

Participants highlighted cost as a key barrier to Agentic AI adoption. In the Indonesian context, relatively low labour costs can weaken the return on investment for AI tools (R1), while high or uncertain subscription expenses further deter uptake, particularly among SMEs. R9 noted that significant upfront and ongoing investment is still required, making adoption financially challenging. Additionally, when organisations do not subsidise tool access, usage declines as cost becomes a burden (Li et al., 2024). Overall, these findings suggest that financial constraints and pricing volatility remain important limitations to widespread adoption of Agentic AI.

Regulatory vacuum and governance lag.

The most striking macro-level finding was the near-complete absence of AI-specific legal or regulatory frameworks in Indonesia. R2 and R14 noted that there are little to no regulations even for software in Indonesia, let alone for AI. The FGD produced a succinct diagnosis: there is a significant lag between the legal aspects of AI and the actual technology, with the latter racing ahead at lightning speed.

P5 – Ethical Concern as Factors of Adoption-to-Deployment Conversion

Accountability Ambiguity and Distributed Liability. When autonomous agents cause production failures, liability becomes structurally obscured across the principal-agent hierarchy. The consensus, articulated by multiple respondents, is that to let the end-client bears the reputational consequence, not the Agentic AI provider. R6 argued that responsibility lies with users who fail to fully understand the capabilities and limitations of the AI they deploy, while R14 compared AI to a neutral tool akin to a knife, stating that its safe use ultimately depends on the user.

Algorithmic Bias and Discriminatory Outputs.

Given that LLMs function as black-box systems, underlying training biases can surface in enterprise production environments. R13 noted that 'We can't force AI to be unbiased, because AI is bound to be biased, as it depends on data. If the data is incomplete, bias will inevitably arise,' reflecting concerns that AI systems may generate bias decision. These risks raise both operational and ethical challenges, especially when AI outputs are used to support decisions affecting individuals as supported by existing literature (Pati, 2025).

Proprietary Data Leakage and Agency Risk.

A common fear surrounding Agentic AI adoption was that proprietary organisational data entered into third-party Agentic AI systems could be accessed, retained, or misused outside the organisation's control. This concern has led some companies to restrict or prohibit the use of external Agentic AI, reflecting broader corporate efforts to establish AI governance and risk management practices. Participants also expressed concerns about agency risk, particularly when AI systems are granted greater autonomy. R5 emphasised the importance of ensuring that Agentic AI systems are controlled appropriately, stating that such systems should be managed 'so that (Agentic AI) is not used by the wrong (third) party.' These findings suggest that concerns about data leakage, unauthorised access, and loss of control remain significant barriers to Agentic AI adoption, highlighting the importance of robust governance frameworks and safeguards for organisational AI use.

Personally Identifiable Information (PII) Exposure and Covert Data Collection.

Several informants expressed unease about the possibility that AI agents could access, collect, or infer information such as names, locations, and other personal data without users fully understanding how that information is being used. These concerns were linked to potential compliance risks under Indonesia's personal data protection regulations. Participants were especially cautious about deploying AI in high-stakes environments involving sensitive organisational data. R11 noted that 'all of that is a bit risky. For example, if we use it in crucial areas, especially security, especially those with data that's considered important,' adding that 'not many people are brave enough to integrate agents there yet.' This suggests that concerns over data privacy, security, and regulatory compliance remain significant barriers to AI adoption, particularly in contexts where the consequences of data exposure or misuse are perceived to be high.

Workforce Displacement Anxiety.

Participants expressed concerns that Agentic AI could reduce demand for certain software engineering roles, particularly at the junior level. One FGD participant commented that Agentic AI 'may be replacing' entire development teams because it can perform multiple role functions like designing, code development and testing by itself. A few participants have come to accept this with R13 noted that humans cannot realistically compete with systems that can operate continuously. He, therefore, concluded Agentic AI adoption as an inevitable response to technological change, stating, 'when something changes, I have to adapt.'

Practitioner's Awareness of Automation Potential

From the Gioia coding of P1 regarding performance expectancy, participants highlighted three valuable benefits from adopting Agentic AI. First, Agentic AI improves productivity and speeds up software development. Participants reported that coding tasks that previously took days can now be completed in minutes. This finding is consistent with previous research, which shows that software teams are motivated to use Agentic AI because of its ability to automate work, reduce manual effort, and enable developers to complete tasks more efficiently (Su et al., 2026). Second, elevated code quality was reported by multiple participants and that AI agents can apply best-practice recommendations and minimise the human error. Third, Agentic AI enables developers to focus on higher-value, innovative, and strategic work by automating

Table 3. Gioia Coding Analysis

Quote	1 st Order Concepts	2 nd Order Themes	Aggregate Dimensions
Basic (coding) function to be done within couple minutes, where I think people usually take probably like at least like a day-ish (FGD).	Rapid code completion vs. Manual effort	Productivity and speed enhancement	Enhancement of Job Performance and Software Quality
Agent AI can kind of like generate the documents for you within minutes (FGD).	Accelerated documentation generation		
Agentic AI is (implemented) to avoid human error (R2).	Early error detection and risk reduction	Elevated Code Quality	
Agentic AI can also help with recommendations based on best practices, because he already has the knowledge (R14).	Improved code consistency and standards		
But on the other hand, there's also an increase in interesting, innovative jobs. So, (Agentic AI) opens up the potential for innovation (R12).	Maintaining competitive advantage through innovation	Cognitive Resource Reallocation	
And we can delegate agents to focus on coding tasks, things that are repetitive like coding, and also manage the source code, so we can be hands-off (R9).	Delegating coding execution to AI agents		First-Mile / Last-Mile Adoption Gap
Integrating (Agentic AI) is easy, because we can just talk to the agents and have them create something. But ensuring they don't produce errors is a headache (R14).	Production error challenges	Integration–Deployment Gap	
Implementing agentic AI is certainly expensive ... Integration with legacy systems is still a problem (R2).	Legacy system integration challenges		
But the problem is that Agentic AI also has a lot of hallucinations (R14).	Hallucinations	Hallucination and Output Unreliability	
Then also (Agentic) AI itself, in terms of quality, may not be good yet (R9).	Poor task response quality		
Things that I've learned, especially using AI, is that it's very important to feed it like enough data first to make sure that it can work properly. (FGD)	Data dependency	Context Provision and Management	
The problem is, providing context to agentic AI is often difficult. They read code, usually focusing more on the "how." How we implement something. But the "why." That "why" is usually context that they can't read from the code. Why did we choose this feature? Why did we implement it? There are many other contexts. And if we don't provide that context and ask them to implement something, it often misses the mark (R10).	Providing context difficulty		
From my perspective, it definitely comes from the top. If they say the company is moving towards an AI-driven company, when it comes down to me, and I'll definitely try AI (R2).	Top-down AI mandate creating immediate adoption intent	Top-Down Leadership Mandate	
The biggest (Agentic AI) driver is direction ... once the direction comes from the top saying that everything has to be there and that (Agentic AI usage) will probably be reflected in the KPIs (R2).	KPI-linked AI mandate as binding behavioural driver		Leadership Directives and Testimony
My experience in the team was probably looking at success stories first. "Oh, I can use AI like this, and that makes people interested in using AI too" (R2).	Success stories drive adoption	Social Proof and Peer Validation	
But the internal drive plays a bigger role (R14).	Internal drive / self-efficacy		
If others are using AI and you are not, you risk being left behind (R13).	Individual Fear of Losing performance	Fear of Non-Adoption	
Maybe out there everything seems to move faster than it does here (without using external AI) (R9).	Competitive necessity in the market		

I'm worried I can't identify the bias (R5).	Difficulty identifying or detecting bias	The Training Deficit	Infrastructure Provision and Governance Gap
The most significant barrier ...is the fear of the unknown and the lack of knowledge (R7).	Fear-of-unknown		
The biggest reason is because labor costs in Indonesia are still much cheaper than the cost of implementing an agentic AI system (R1).	Cheap local labor undercuts ROI of agentic AI investment		
It's still difficult at the moment. Firstly, we might have to pay more to invest in AI agents. We might still have to spend a lot of money (R9).	Financial and Resource Constraints		
I've never heard of what the regulations are in Indonesia regarding the use of training data (R2)	Complete absence of Indonesian AI training-data regulation		
There aren't even regulations governing software. Let alone AI? (R14)	Absence of foundational software regulation as baseline concern		
As far as I know, it should be the one using it. Yes, the first person to use the AI agent is because they, as the user ... didn't really study the features of the AI they were using (R6).	Responsibility attributed to the user	Accountability Ambiguity and Distributed Liability	Ethical Concern
Like the knife earlier ... the level of safety is returned to the user (R14).	AI compared to neutral tools		
We can't force AI to be unbiased, because AI is bound to be biased, as it depends on data. If the data is incomplete, bias will inevitably arise (R13).	Agentic AI bias reflects incomplete or skewed datasets	Algorithmic Bias and Discriminatory Outputs	
If you have watched Coded Bias from Netflix, (racial bias) has happened. Unfortunately, it has happened before (FGD).	Documented real-world AI racial bias cases requiring data remediation		
As a company, we prohibit all our employees from using the free (AI tools) version so no one's data will be used for training. That is quite dangerous. We are very concerned about that (R1).	Company Trade Secret used for training data	Proprietary Data Leakage and Agency Risk	
Furthermore, the boundaries and sharing of sensitive data with AI also need to be clear. So that (Agentic AI) is not used by the wrong (third) party (R5).	Agency Risk and third party misused		
All of that is a bit risky. For example, if we use it in crucial areas, especially security, especially those with data that's considered important. So, it's possible to use it in areas like that. Not many people are brave enough to integrate agents there yet. (R11)	Personal Data Privacy and Security	Personally Identifiable Information (PII) Exposure Covert Data Collection	
I was surprised when AI showed my name without me writing it...it got it from login data (R13).	Implicit data harvesting from login metadata without user awareness		
I think the most concerning things is this AI may be replacing us...maybe in the future one company exists (as) just one person and the AI (FGD).	One-human-plus-AI replacing entire company.	Workforce Displacement Anxiety	
My principle is that when something changes, I have to adapt (R13).	Acceptance of technological evolution as inevitable		

Source: Primary Data

repetitive tasks. As AI takes over routine cognitive work, developers and other knowledge workers can devote more time to tasks that require human judgment, problem-solving, and creativity (Geninatti Cossatin et al., 2026).

Social Influence Mixed Effects on Agentic AI Adoption Intention

The findings suggest that social influence has a mixed effect on employees' intention to adopt Agentic AI. Social

influence appears to shape employees' perceptions and attitudes toward the technology, as several participants reported that executive mandates, peer success stories, and concerns about falling behind create normative pressures that encourage exploration and experimentation. This is supported by Cheng et al. (2023) who examine how community interactions contribute to developers' trust in AI tools.

Their study provides a comprehensive analysis of the

influence of online communities, including forums and discussion groups, on developers' perceptions of AI-assisted technologies. The findings suggest that shared experiences and collective discussions are key factors in fostering trust and shaping developers' confidence in AI-assisted development tools.

However, the decision to adopt and actively use Agentic AI also depends on individual factors. One participant highlighted the importance of personal motivation and self-efficacy in determining actual usage. This mixed view is consistent with previous studies on AI tool adoption among Indonesian workers, which found that social influence is not a primary determinant of behavioural intention (Anditama & Hidayanto, 2024; Samudra et al., 2024). This finding partially diverges from UTAUT, which identifies social influence as a significant predictor of technology adoption.

While awareness of performance potential and social pressure encourage interest and attitude towards Agentic AI adoption, three structural gaps explain further how Agentic AI Adoption paradox is sustained.

Gap 1 — The Capability-Deployment Verification Gap

Across all organisational tiers, participants demonstrated near-universal recognition that Agentic AI delivers genuine, measurable productivity benefits. This indicates that awareness of its value is not a barrier to adoption. However, its underlying output unreliability, caused by hallucinations and context rot, severely hinders broader production deployment. Practitioners reported integration challenges including context-management burdens, limited prompting skills, and the extensive effort required to verify autonomous outputs.

These findings align with prior literatures identifying a persistent gap between experimental agentic capabilities and their qualification for integration into software development workflows (Apostolou et al., 2026; Rasheed et al., 2026). Traditional certification frameworks are designed for deterministic systems and provide limited guidance for validating Agentic AI, whose outputs may vary across runs and model versions (Sun & Staron, 2026). As a result, Agentic AI cannot be reliably audited against user requirements or architectural design specifications.

Ultimately, these hurdles expose a stark asymmetry between the ease of initiating Agentic AI and the difficulty of sustaining reliable, enterprise-scale deployment. Viewed through UTAUT frameworks, this highlights a fundamental gap between the technology's Performance Expectancy and its Effort Expectancy, where strong recognition of productivity benefits is offset by substantial operational, verification, and governance challenges that constrain adoption in production environments.

Gap 2 — The ROI Gap

Although Agentic AI is expected to improve productivity and accelerate software development, its scalability may be constrained by the rising costs associated with operating and governing AI systems. Haque (2026) argues that Agentic AI significantly reduces software development costs by automating activities such as planning, coding, testing, and iteration. However, these savings do not eliminate costs; instead, they shift expenditure toward operational infrastructure, model inference, monitoring, cybersecurity, compliance, and governance. As development becomes cheaper, the costs of running and managing AI-enabled systems become increasingly important.

This cost shift presents a particular challenge in Indonesia. The economic justification for Agentic AI is

often based on the assumption that replacing or augmenting highly paid knowledge workers will generate substantial savings. However, because labour costs in Indonesia are relatively low as one participant mentioned so the financial benefits of substituting human workers with Agentic AI are less pronounced. While organisations may still gain productivity improvements, the return on investment may not be sufficient to offset the significant costs required to scale AI infrastructure and governance capabilities.

As organisations expand their use of Agentic AI, operational expenses can increase through greater demand for computing resources, AI orchestration tools, reliability engineering, and risk management mechanisms. Consequently, the marginal cost of scaling Agentic AI may rise faster than the economic value it creates. This creates a potential gap between Performance Expectancy and Facilitating Conditions. Executives may recognise Agentic AI's potential to improve software quality and productivity but remain hesitant to commit the resources needed for enterprise-wide deployment.

Gap 3 — Usefulness-as-Ethical-Risk-Amplifier

Ethical concerns can become a major barrier to the production deployment of Agentic AI because the capabilities that make these systems valuable also increase organisational risk. As Agentic AI gains access to organisational data and participates in more consequential decisions, concerns about accountability become more pronounced because of "black-box" nature of AI decision-making (Murugesan, 2025). The current research findings and existing literature also highlight concerns related to data quality, bias, privacy, and security. Since Agentic AI relies on large volumes of data, biased or inaccurate inputs can lead to unfair or flawed decisions (Jovanović et al., 2025). At the same time, granting AI agents access to sensitive organisational information increases the risk of data breaches, cyberattacks, and misuse by malicious actors. These risks become more significant as AI systems gain broader authority and access. Consequently, an adoption paradox emerges in which organisations recognise the value of Agentic AI yet hesitate to deploy it at scale because the realisation of its benefits simultaneously generates ethical, legal, and governance risks that hinder large-scale production deployment.

Beyond cost consideration, satisfying Agentic AI Adoption Paradox and its inherent ethical risk requires an effective AI governance and monitoring.

Governance and Monitoring Implications

An effective governance and monitoring of Agentic AI require satisfying three conditions. First, it requires embedding ethical guardrails at design time rather than retrofitting constraints post-deployment. Since ethical risks are inherent to the capabilities that drive value, governance must be co-designed into agent architecture through constitutional principles and capability-constrained design, preventing misalignment before it scales. Second, transparent and explainable audit logging enables accountability and reduces the risk of undetected ethical violations by providing traceable evidence of agent reasoning and decisions, facilitating root-cause identification and corrective iteration. Third, autonomy calibration must be tiered based on agent reliability and task risk exposure. Given LLM-based agent systems' inherent unreliability, human oversight remains necessary. However, human oversight could hardly keep up with the scaling potential of Agentic AI and therefore proven agents dealing with low-risk task should still be given autonomy while maintaining close supervision for unvalidated agent or high-risk tasks.

In practice, organisations can operationalise these governance mechanisms by deploying Agentic AI within a sandboxed environment—a controlled real-world testing setting established to evaluate AI agents under defined conditions before full-scale deployment. Within such an environment, design-time ethical constraints, comprehensive audit logs, and tiered autonomy can all be tested safely to see if they form a proportionate governance model that mitigates both technical and ethical risks without sacrificing operational efficiency.

Comparison with Previous Studies

The finding that productivity enhancement drives Agentic AI adoption is consistent with Li et al. (2024), who reported similar motivations among software practitioners using generative AI tools. However, this study documents a Capability-Deployment Verification Gap absent from Li et al.'s generative AI-focused framing, where output unreliability structurally obstructs enterprise deployment.

The present study differs from Akbar et al. (2025) and Su et al. (2026), who analysed Agentic AI adoption primarily at the SDLC-phase level and through architectural lens respectively. While their frameworks illuminate the technical dimensions of agent integration and coordination patterns, neither addresses the economic forces determining organisational deployment. One important contribution of this study is the identification of the ROI Gap. In Indonesia, relatively low labour costs reduce the financial benefits of replacing human work with AI, making organisations less willing to invest in the infrastructure needed for Agentic AI. This issue has no direct parallel in the other two literatures, which are predominantly anchored in higher-wage Western contexts where labour substitution economics favour Agentic AI adoption more readily.

Meanwhile, Rizinski & Trajanov (2026) conducted a systematic review of AI Agent-based systems in finance and fintech, emphasising their potential and the technical and ethical challenges. Similarly, Sayyad et al. (2025) reviewed the key technical and ethical issues that must be addressed to ensure Agentic AI is trustworthy and socially acceptable to deliver environmental sustainability. While our findings align with these studies in identifying common technical challenges and ethical concerns, their concentration on the financial services sector and environmental sectors renders their findings only tangentially applicable to a software engineering organisational context.

Regarding ethical concern, this study diverges from Jovanović et al. (2025) and Ferino et al. (2026), both of whom treat ethical concern as an external boundary condition or peripheral risk category that constrains adoption from outside the core adoption decision. In contrast, this study found that ethical concerns are closely connected to the core capabilities of Agentic AI. As AI systems become more autonomous and gain greater access to organisational data and decision-making processes, risks such as unclear accountability, algorithmic bias, data leakage, and job displacement also increase. This suggests that ethical considerations should be built into the design and governance of Agentic AI rather than being addressed after deployment. This is aligned with Tejaskumar Pujari et al. (2024), who proposed layered governance for multi-agent systems, yet lacked empirical grounding.

Limitations and Cautions

This study's findings should be interpreted in light of its qualitative case study design.

The dataset, comprising sixteen semi-structured

interviews and one focus group discussion, was purposively sampled to maximise variation across organisational roles, experience levels, and technical specialisations. The findings reflect contextual interpretations that may not be directly transferable to other organisational or national settings.

In addition, while the study identifies key constructs relevant to Agentic AI adoption, these remain conceptually derived and have not yet undergone quantitative validation, which limits the extent to which the proposed relationships can be inferred as generalisable or causally robust.

Recommendations for Future Research

Future research should extend this work in two complementary directions. First, longitudinal studies are needed to examine how AI governance investments influence the evolution of adoption gaps across organisational tiers over time, providing stronger evidence of temporal dynamics and potential causal pathways. Comparative cross-country research would further help determine whether Indonesia's observed structural conditions are context-specific or indicative of broader patterns across emerging digital economies. Second, the conceptual model developed in this study should be empirically validated through large-scale quantitative research using PLS-SEM, with sufficient sample size to ensure statistical power and reliable estimation of the constructs—Perceived Usefulness, Ease of Use, Social Influence, Facilitating Conditions, and Ethical Concerns—and their associated indicators derived from the qualitative findings.

Conclusion

This study set out to explain why Indonesian software organisations embrace Agentic AI in principle while struggling to deploy it at production grade. The findings support the proposed conceptual model: adoption intention and deployment capability diverge, and this divergence is not resolved by technical readiness alone but is actively shaped by weak facilitating conditions and by ethical concern that intensifies as its usefulness becomes more visible. This last point is the study's central theoretical contribution: it repositions Ethical Concern from a peripheral adoption barrier to an endogenous, value-scaling risk factor, and in doing so offers a structural account of why the adoption paradox persists rather than resolving itself as organisations mature.

The practical implication follows directly. Because the paradox is structural rather than a temporary capability gap, it will not close through incremental technical investment alone. Organisations require governance mechanisms embedded at the point of deployment to ensure its usefulness while addressing its ethical risk exposure. For software industry, this reframes Agentic AI governance as a precondition for scaling deployment, not a compliance overhead that follows it.

These conclusions should be read against the study's design limitations: a qualitative, single-country, non-probability sample without quantitative validation constrains generalisability. Longitudinal and cross-country replication, together with large-scale PLS-SEM testing of the proposed constructs, represent the logical next step for establishing the model's external validity and for translating it into actionable governance standards.

Author contributions

Ricky Tanudjaja contributed to the conceptualisation and execution of the study, including conducting background research, reviewing and comparing existing literature, designing the research framework, collecting data through participant interviews, transcribing the interviews, coding the qualitative data, and synthesising the findings and implications. Sahat Hutajulu contributed in a supervisory and advisory capacity, providing critical feedback, guidance, and suggestions to strengthen the research design and improve the completeness and clarity of the manuscript.

Funding

The authors declare that no funds, grants, or other support were received during the preparation of this manuscript.

Acknowledgements

References

- Akbar, M. A., Khan, A. A., Hamza, M., Ghaffar, A., & Hajikhani, A. (2025). *Agentic AI in Software Engineering: Practitioner Perspectives Across the Software Development Life Cycle*. <https://doi.org/10.2139/ssrn.5520159>
- Apostolou, S. A., Bosch, J., & Olsson, H. H. (2026). *Agentic AI in Industry: Adoption Level and Deployment Barriers*.
- Charlotte Relyea, Martin Harrysson, Matt Linderman, Jose Mario Pena, & Nandita Bothra. (2025, November 3). *Unlocking the value of AI in software development*. McKinsey. <https://www.mckinsey.com/industries/technology-media-and-telecommunications/our-insights/unlocking-the-value-of-ai-in-software-development#/>
- Cheng, R., Wang, R., Zimmermann, T., & Ford, D. (2023). *"It would work for me too": How Online Communities Shape Software Developers' Trust in AI-Powered Code Generation Tools*.
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly: Management Information Systems*, 13(3). <https://doi.org/10.2307/249008>
- Ferino, S., Hoda, R., Grundy, J., & Treude, C. (2026). *Walking the Tightrope of LLMs for Software Development: A Practitioners' Perspective*.
- Fishbein, M. (1976). A Behavior Theory Approach to the Relations between Beliefs about an Object and the Attitude Toward the Object. https://doi.org/10.1007/978-3-642-51565-1_25
- Gartner. (2025, June 25). *Gartner Predicts Over 40% of Agentic AI Projects Will Be Canceled by End of 2027*. Gartner. <https://www.gartner.com/en/newsroom/press-releases/2025-06-25-gartner-predicts-over-40-percent-of-agentic-ai-projects-will-be-canceled-by-end-of-2027>
- Geninatti Cossatin, A., Ferrero, F., Ardisson, L., & Mauro, N. (2026). The autonomy equation: How agentic AI reshapes trust and workload in routine productivity applications. *Information Processing & Management*, 63(5), 104681. <https://doi.org/10.1016/j.ipm.2026.104681>
- Gioia, D. A., Corley, K. G., & Hamilton, A. L. (2013). Seeking Qualitative Rigor in Inductive Research: Notes on the Gioia Methodology. *Organizational Research Methods*, 16(1). <https://doi.org/10.1177/1094428112452151>
- Glaser, B. G., & Strauss, A. L. (2017). Discovery of grounded theory: Strategies for qualitative research. In *Discovery of Grounded Theory: Strategies for Qualitative Research*. <https://doi.org/10.4324/9780203793206>
- Guest, G., Bunce, A., & Johnson, L. (2006). How Many Interviews Are Enough?: An Experiment with Data Saturation and Variability. *Field Methods*, 18(1). <https://doi.org/10.1177/1525822X05279903>
- Haque, W. (2026). The New Software Cost Curve under Agentic AI: Development Economics, Pricing, and Labor Impacts. *Journal of Software Engineering and Applications*, 19(01), 1–6. <https://doi.org/10.4236/jsea.2026.191001>
- IBM. (2025, February 9). *AI Alliance Announces Expansion to Indonesia*. IBM. https://asean.newsroom.ibm.com/AI-Alliance-Announces-Expansion-to-Indonesia?utm_id=IBM-and-Meta---AI-Alliance
- Iiris Aaltio, & Pia Heilmann. (2010). Case study as a methodological approach: From locality to understanding the essence. *Encyclopedia of Case Study Research*, 66–76.
- Jovanović, S., Dražeta, L., Petrović, A., Bačanin Džakula, N., & Živković, M. (2025). Ethical Concerns and Mass Agentic AI Adoption. *Proceedings of the International Scientific Conference - Sinteza 2025*, 369–375. <https://doi.org/10.15308/Sinteza-2025-369-375>
- Microsoft & LinkedIn. (2024, February 11). *Microsoft and LinkedIn Release the 2024 Work Trend Index on the State of AI at Work in Indonesia*. Microsoft. <https://news.microsoft.com/id-id/2024/06/11/microsoft-and-linkedin-release-the-2024-work-trend-index-on-the-state-of-ai-at-work-in-indonesia/>
- Muhammad Rizky Anditama, & Achmad Nizar Hidayanto. (2024). Analysis of Factors Influencing the Adoption of Artificial Intelligence Assistants among Software Developers. *The Indonesian Journal of Computer Science*, 13(4). <https://doi.org/10.33022/ijcs.v13i4.4239>
- Murugesan, S. (2025). The Rise of Agentic AI: Implications, Concerns, and the Path Forward. *IEEE Intelligent Systems*, 40(2). <https://doi.org/10.1109/MIS.2025.3544940>
- Noy, S., & Zhang, W. (2023). Experimental evidence on the productivity effects of generative artificial intelligence. *Science*, 381(6654), 187–192. <https://doi.org/10.1126/science.adh2586>
- Pati, A. K. (2025). Agentic AI: A Comprehensive Survey of Technologies, Applications, and Societal Implications. *IEEE Access*, 13. <https://doi.org/10.1109/ACCESS.2025.3585609>
- Patton, M. Q. (2015). Qualitative research and evaluation methods fourth edition. In *Qualitative Inquiry* (Vol. 3rd).
- Pierre Aubret, & Paul Florentin. (2025). <https://www.riskinsight-wavestone.com/en/2025/07/agentic-ai-typology-of-risks-and-security-measures/>. RiskInsight. <https://www.riskinsight-wavestone.com/en/2025/07/agentic-ai-typology-of-risks-and-security-measures/>
- Rasheed, Z., Waseem, M., Kemell, K.-K., Ahmad, A., Sami, M. A., Saari, M., Rasku, J., & Abrahamsson, P. (2026). *CodePori: Large-Scale System for Autonomous Software Development Using Multi-Agent Technology*.
- Rizinski, M., & Trajanov, D. (2026). AI Agents in Finance and Fintech: A Scientific Review of Agent-Based Systems, Applications, and Future Horizons. *Computers, Materials & Continua*, 86(1), 1–34. <https://doi.org/10.32604/cmc.2025.069678>
- Russo, D. (2024). Navigating the Complexity of Generative AI Adoption in Software Engineering. *ACM Transactions on Software Engineering and Methodology*, 33(5). <https://doi.org/10.1145/3652154>
- Samudra, B. S., Baihaqi, I., & Isnaini, F. (2024). The Use of Generative AI in Workplace: Driving Factors, Barriers, and Benefits. *TEMSCON-ASPAC 2024 - IEEE Technology and Engineering Management Conference - Asia Pacific*. <https://doi.org/10.1109/TEMSCON-ASPAC62480.2024.11025027>
- Sayyad, N., Dave, H., Kathane, M., & Patel, R. (2025). Agentic AI Systems: A Review of Architectures, Autonomy, and Ethical Implications. *International Journal of Environmental Sciences*. <https://doi.org/10.64252/xgrhd85>
- Su, R., Esposito, M., Capuano, R., Omar, R., Sallou, J., Muccini, H., & Taibi,

- D. (2026). *Making Sense of AI Agents Hype: Adoption, Architectures, and Takeaways from Practitioners*.
- Sun, S., & Staron, M. (2026). *Agentic Pipelines in Embedded Software Engineering: Emerging Practices and Challenges*.
- Tejaskumar Pujari, Anshul Goel, & Ashwin Sharma. (2024). Ethical and Responsible AI: Governance Frameworks and Policy Implications for Multi-Agent Systems. *International Journal Science and Technology*, 3(1), 72–89. <https://doi.org/10.56127/ijst.v3i1.1962>
- V, A., R., G. G., & Buyya, R. (2026). *Agentic Artificial Intelligence (AI): Architectures, Taxonomies, and Evaluation of Large Language Model Agents*.
- Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS Quarterly: Management Information Systems*, 27(3). <https://doi.org/10.2307/30036540>
- Yang, J., Yonack, N., Zyskowski, K., Yarats, D., Ho, J., & Ma, J. (2025). *The Adoption and Usage of AI Agents: Early Evidence from Perplexity*.
- Yin, R. K. (2003). Case Study Research . Design and Methods. In *SAGE Publications* (Vol. 26, Number 1). <https://doi.org/10.1097/FCH.0b013e31822dda9e>