



UNIVERSITÀ POLITECNICA DELLE MARCHE
FACOLTÀ DI ECONOMIA “GIORGIO FUÀ”

Master's or Specialist Degree Course in International Economics and Commerce
Specialisation: Business Organisation and Strategy

**Generative AI in Product Discovery: Adoption
Patterns, Perceived Value, Barriers, and Role
Evolution Among Product Managers**

Generative AI nel Product Discovery: Modelli di adozione, valore
percepito, barriere ed evoluzione del ruolo tra i Product Manager

Supervisor:

Prof. Marco Cucculleli

Degree Thesis of:

Yash Sudhir Todkari

Co-Supervisor:

Prof. Olgun Aydin

Academic Year 2025 – 2026

Acknowledgements

The completion of this thesis would not have been possible without the support, guidance, and encouragement of a number of people, to whom I owe my sincere gratitude.

First and foremost, I wish to express my deepest appreciation to my supervisor, Prof. Marco Cucculelli, for his invaluable guidance, scholarly insight, and patience throughout the development of this research. His thoughtful feedback and encouragement shaped both the direction and the rigour of this work. I am equally grateful to my co-supervisor, Prof. Olgun Aydin, whose expertise, constructive suggestions, and methodological guidance were instrumental at every stage of this study.

I would also like to extend my heartfelt thanks to the Product Manager (PM) and Product Owners who generously took the time to participate in the survey that forms the empirical backbone of this thesis. Their willingness to share their experiences and perspectives on Generative Artificial Intelligence (GenAI) in product discovery made this research possible, and their candour gave the findings their value.

Finally, I am profoundly grateful to my family for their unwavering love, patience, and support throughout my academic journey. Their belief in me has been a constant source of motivation, and this accomplishment is as much theirs as it is mine.

Abstract

Generative Artificial Intelligence (GenAI) has entered knowledge-work environments at an unprecedented pace, yet little academic research has examined how it is being adopted within the specific practice of product discovery. This thesis investigates how PMs are integrating GenAI into their discovery workflows, and explores the perceived value of these tools, the barriers limiting their effective use, and the implications of adoption for the evolving role of the PM.

Adopting a pragmatist research philosophy and a convergent parallel mixed-methods design, the study combines a quantitative survey of 82 practicing Product Managers (PMs) and Product Owners with a qualitative case study of four leading product companies (Figma, Notion, Miro, and Atlassian) drawn from publicly available documentation. The analysis is framed by three theoretical lenses: the Technology Acceptance Model and the Unified Theory of Acceptance and Use of Technology as adoption frameworks; Daugherty and Wilson's augmentation theory; and the AI Fluency Framework developed by Dakan and Feller, applied here for the first time in the academic literature to interpret professional GenAI adoption behaviour.

The findings indicate that GenAI has become a routine instrument of product discovery for most practitioners, valued chiefly for accelerating repetitive cognitive work such as research synthesis, drafting, and coordination. This value is, however, bounded: practitioners trust GenAI markedly less than they find it useful, with hallucination risk, missing organisational context, and governance constraints emerging as the principal barriers to deeper adoption. Trust is therefore shown to be conditional and situational, and human discernment remains indispensable to the discovery process.

The study concludes that the trajectory of GenAI in product discovery is one of augmentation rather than replacement, shifting the centre of gravity of the PM role away from manual production and toward judgement, interpretation, validation, and orchestration. As its principal original contribution, the thesis proposes a GenAI Adoption Maturity Model for product teams, positing that organisations advance not by using AI more, but by using it with greater discernment, calibrated trust, and governance maturity.

Keywords: Generative AI (GenAI); Product Discovery; Product Management; Product Manger; AI Adoption; Technology Acceptance Model; AI Fluency; Human-AI Collaboration; Augmentation; Maturity Model.

Statement on the Use of Artificial Intelligence

In the interest of transparency and in accordance with the academic-integrity principles of Università Politecnica delle Marche and prevailing European guidelines on the responsible use of Artificial Intelligence in research, the author discloses that GenAI tools were used in a supporting and supervised capacity during the preparation of this thesis. This statement is provided in keeping with the transparency commitments set out in Chapter 3 (Sections 3.8 and 3.9).

Specifically, AI-based tools were used for the following purposes: to assist in the identification, retrieval, and preliminary analysis of academic and practitioner literature relevant to the research questions; to provide writing assistance limited to language refinement, including grammar correction and paraphrasing for clarity and readability; to support the organisation and analysis of the anonymised survey data; and to assist in the generation of graphs and figures used to illustrate, more clearly, the quantitative data presented in tabular form.

The author affirms that the conception, research design, argumentation, interpretation of results, and conclusions of this thesis are entirely his own. All AI-generated or AI-assisted material was critically reviewed, verified, and edited by the author, who takes full and sole responsibility for the accuracy, integrity, and originality of the content presented. GenAI was not used to fabricate data, to produce original analytical conclusions, or to substitute for the author's own scholarly judgement. Where AI tools were employed, their outputs were treated as a starting point for human evaluation rather than as authoritative results, consistent with the principle of human discernment discussed within this thesis.

Table of Contents

Acknowledgements	i
Abstract	ii
Statement on the Use of Artificial Intelligence.....	iv
Table of Contents.....	v
Chapter 1: Introduction	1
1.1 Background and Context.....	1
1.2 Motivation and Relevance.....	2
1.3 Research Gap.....	2
1.4 Research Aim and Objectives.....	3
1.5 Research Questions.....	3
1.6 Scope and Delimitations.....	4
1.7 Significance of the Study.....	4
1.8 Structure of the Thesis.....	4
Chapter 2: Literature Review.....	6
2.1 Introduction.....	6
2.2 Product Discovery: Theory, Frameworks, and Practice	6
2.3 GenAI in Business: Adoption, Productivity, and Organisational Context.....	9
2.4 Human-AI Collaboration: Theoretical Frameworks.....	12
2.5 The Evolving Role of the PM.....	18
2.6 Synthesis: Towards a Conceptual Framework.....	20
2.7 Chapter Summary	23
Chapter 3: Research Methodology.....	24
3.1 Introduction.....	24
3.2 Research Philosophy.....	24
3.3 Research Design.....	25
3.4 Sampling Strategy and Target Population.....	26
3.5 Survey Instrument Design.....	27

3.6 Case Study Selection and Analysis.....	32
3.7 Data Analysis Plan.....	33
3.8 Ethical Considerations.....	35
3.9 Limitations of Research Design.....	36
3.10 Chapter Summary.....	37
Chapter 4: Case Study Analysis of GenAI in Leading Product Companies.....	38
4.1 Introduction and Purpose.....	38
4.2 Figma: AI-Augmented Visual Discovery and Prototyping.....	38
4.3 Notion: AI for PM Documentation, Research Synthesis, and Knowledge Management	41
4.4 Miro: GenAI in Visual Discovery Workshops and Continuous Innovation.....	44
4.5 Atlassian: Enterprise-Scale GenAI Governance and Product Lifecycle Integration...	46
4.6 Cross-Case Synthesis, Five Consistent Themes.....	49
4.7 Chapter Summary.....	53
Chapter 5: Survey Findings and Analysis.....	54
5.1 Introduction and Data Quality Review.....	54
5.2 Respondent Profile.....	54
5.3 GenAI Adoption and Usage (Q5-Q8, SRQ1).....	58
5.4 Perceived Value (Q9-Q11, SRQ2).....	62
5.5 Barriers and Trust (Q12-Q13, SRQ3).....	70
5.6 PM Role Evolution and Skills (Q14-Q15, SRQ4).....	70
5.7 Cross-Tabular Analysis.....	72
5.8 Thematic Analysis of Open-Ended Responses.....	83
5.9 Chapter Summary.....	89
Chapter 6: Results Discussion, Triangulation, and Conclusion.....	90
6.1 Introduction.....	90
6.2 SRQ1: Patterns of GenAI Adoption in Product Discovery.....	90
6.3 SRQ2: Perceived Value of GenAI.....	91

6.4 SRQ3: Barriers and the Nature of Trust.....	91
6.5 SRQ4: Evolution of the PM Role.....	92
6.6 Refined GenAI Adoption Maturity Model.....	93
6.7 Triangulation: Survey, Literature, and Case Studies.....	95
6.8 Contributions, Limitations, and Future Research.....	96
6.9 Chapter Conclusion.....	98
References.....	99

Chapter 1: Introduction

1.1 Background and Context

Generative Artificial Intelligence (GenAI) technologies are entering workplaces across almost all industries and disciplines, affecting how knowledge workers perform their jobs. In late 2022, OpenAI made ChatGPT publicly available, causing what was, for the tech world, unprecedented disruption. What was once an experimental research project experienced rapid adoption within weeks as a professional tool suddenly found itself in nearly every workplace. As of 2024, knowledge workers in legal, finance, marketing, and tech (among others) are beginning to integrate these technologies into their workflows, whether slowly and skeptically, or quickly and eagerly, with little to no formal guidance from their employers.

One such role these tools will impact significantly is the Product Manager (PM). PMs find themselves at a unique intersection of skill sets and responsibilities. They are tasked with aggregating and synthesizing both qualitative and quantitative information; making high-impact, judgment-based decisions with inherent uncertainty; and articulating the why and how behind these decisions to engineering, design, and business stakeholders. GenAI tools have already shown great promise in assisting with (if not outright replacing) many knowledge worker and language tasks.

At the core of what it means to be a PM is a practice known as product discovery. Product discovery is the process by which product teams identify what they should build prior to building it (validating user problems, understanding potential solutions, and testing assumptions) before jumping into development [1]. It includes whiteboarding sessions, synthesizing user interview notes, conducting competitive analysis, brainstorming product ideas, and writing requirements docs. For years, these activities have been done by-hand and have taken up significant portions of PMs' time. With the recent introduction of ChatGPT, Notion AI, Miro AI, Microsoft Copilot, Claude, and others, this is starting to change (at least, in theory).

The gap in knowledge that this research project addresses lies at the intersection of theory and practice. While there is significant discussion about GenAI at work, there is a lack of academic research exploring what is occurring in practice within product teams. How are PMs specifically using these tools for discovery? Are they finding them useful, or simply

trendy? What are the barriers? And most importantly, are these tools beginning to change the nature of the PM role itself?

1.2 Motivation and Relevance

The motivation for this research question is threefold. Firstly, it matters because product management itself seems to be at a crossroads. The McKinsey 2023 Global Survey on AI found that whilst one-third of respondents stated that their organisations were already using GenAI on a regular basis across at least one organisational function in 2023, that number is expected to grow to 79 percent by 2025 [2]. Product and service development featured as one of the primary functions impacted by this trend. Despite this widespread organisational adoption, knowledge worker communities (especially among product professionals) continue to echo a sentiment that individual PMs adopting these tools are doing so without much top-down coordination and are being motivated largely by personal interest.

There is also an ongoing discussion taking place at the individual talent level regarding what skills a PM needs to continue to do their jobs effectively when working alongside GenAI. The community is discussing what “traditional” PM skills will continue to hold value, which may become less important, and what totally new skills (how to ask AI the right questions, how to spot when it’s wrong, how to work with it responsibly) will be required. Lastly, this study was inspired by the author's direct experience with the AI Fluency Training programme developed by Dakan and Feller in collaboration with Anthropic (2025). The program outlined that there are four primary skillsets that are required to work effectively with GenAI: Delegation (understanding what to tell AI to do), Description (learning how to communicate your requests with specificity), Discernment (spotting when AI is wrong), and Diligence (using AI responsibly and keeping a record of your use). Completing this programme provided the author with a structured way of thinking about GenAI adoption in product discovery and shaped both the design of this study and the analytical framework applied in Chapter 5.

1.3 Research Gap

While there is research on GenAI in the workplace and research on product management (including discovery work, Agile methodologies, and the professional role of the PM), there does not yet exist (at the time of writing) a research publication analysing how PMs working in product discovery are adopting GenAI. Much of the research into GenAI at

work centres around questions of productivity impacts, potential for automation, and overall organisational readiness. Much of the research into product management does the same for that field. But what about where those fields intersect?

This thesis seeks to fill that void. Through conducting a survey amongst practicing PMs and observing GenAI use as reported by product companies publicly, this thesis collects data on current GenAI use in product discovery and interprets that data to understand what is enabling and limiting product teams from adopting GenAI.

1.4 Research Aim and Objectives

The purpose of this thesis is to explore the extent to which PMs are currently adopting GenAI into their product discovery workflows, as well as the perceived impacts, challenges, and implications for the role that adoption presents.

To accomplish this, the study has four goals:

- Describe the current state of GenAI adoption within product discovery
- Measure the value GenAI is perceived to bring to product discovery
- Identify challenges limiting effective GenAI adoption in product organisations
- Assess how GenAI is changing the required skillsets and practices of the PM role

1.5 Research Questions

To address the research aim of this study, one main research question and a set of supporting sub-research questions have been formulated as follows:

- ❖ Main Research Question (RQ):
How are PMs integrating GenAI into the product discovery process, and what are the perceived impacts on efficiency, output quality, and business outcomes?
- ❖ Sub-Research Questions (SRQs):
 - SRQ1: Which GenAI tools and use cases have been adopted within product discovery workflows, and at what stage of adoption maturity are product teams adopting these tools?
 - SRQ2: What value and business outcomes do PMs perceive GenAI integrations bring to their product discovery efforts?
 - SRQ3: What are the primary barriers (individual, organisational, or technological) preventing product teams from adopting GenAI effectively in their discovery processes?
 - SRQ4: How has the introduction of GenAI tools changed the skills needed to be an effective PM, and how is the role of the PM evolving?

1.6 Scope and Delimitations

This thesis is concerned with just one phase of the broader product management lifecycle. Product management covers a broad spectrum of activities (from roadmap planning to sprint delivery to stakeholder management to post-launch analysis etc.). The scope has been narrowed to product discovery because that is the phase which is most closely associated with learning - and by extension creating and validating new knowledge - about users and markets, and therefore one which presents the greatest opportunity for GenAI augmentation. PMs and Product Owners make up the target population of this study. Data was gathered through a structured online survey which was circulated on LinkedIn and other professional PM communities using convenience sampling. Participants included PMs and POs of various company sizes across industries. Distribution channels were limited to Europe, mostly because of the author's positioning on LinkedIn as an Erasmus double degree student between Italy and Poland, but respondents were otherwise global and survey questions included region of residence.

1.7 Significance of the Study

Academically, this research seeks to make an original contribution to a nascent area of interdisciplinary inquiry at the cross-section of Human-AI Interaction, Product Management, and Organisational Behaviour. It leverages the AI Fluency Framework [3] as a theoretical framework for making sense of PM adoption behaviour, the first study in the academic literature to do so, and proposes a GenAI Adoption Maturity Model for product teams, based on empirical data. For practitioners, the research will shed light on how product teams are actually using GenAI, what's working and what's not, and what skills and organisational conditions are correlated with success. Many organisations are only beginning to formulate their first GenAI strategies. Empirical research like this can have utility beyond academic contributions.

1.8 Structure of the Thesis

This thesis consists of six chapters. Chapter 1 situates the background, motivation, gap, aim, objectives, research questions, scope, and significance of the study. Chapter 2 synthesizes and critically reviews extant academic and practitioner literature related to the study's research questions, organised according to four thematic pillars. Chapter 3 presents the details and justifies the research methodology, including study design,

population/sample, survey instrument, case study selection criteria, data analysis techniques, ethical considerations, and limitations. Chapter 4 presents the case study analysis. Chapter 5 presents the survey findings and analysis, including cross-tabular and thematic analysis of the data. Chapter 6 discusses the results in light of the literature, presents the GenAI Adoption Maturity Model as the study's primary original contribution, and concludes the study.

Chapter 2: Literature Review

2.1 Introduction

This chapter covers a review of literature that bears upon the research questions identified for this thesis. It is broken into four sections that align with the four pillars presented in Chapter 1; each provides additional background and frames the ensuing research chapters. The first pillar examines background on product discovery theory and practice, and the recent shift from linear approaches to design towards more continuous discovery practices. The second pillar surveys extant literature on GenAI adoption in the business/professional context, with emphasis on work related to GenAI productivity and organisational impediments. The third pillar reviews theoretical approaches to human collaboration with AI systems, focusing on the Technology Acceptance Model [4] and AI Augmentation [5] literature streams and positioning the AI Fluency Framework [3] for introduction later in the chapter. The fourth pillar examines how work has described shifts in the role of the PM over time, as both their skills and responsibilities have been influenced by Agile practices and will undoubtedly continue to evolve with GenAI adoption. Taken together, these four background sections form the basis of theory upon which the survey research is interpreted in Chapters 4 and 5.

2.2 Product Discovery: Theory, Frameworks, and Practice

2.2.1 Defining Product Discovery

Product discovery refers to the processes by which product teams explore problems worth solving for users, brainstorm and evaluate solutions, and test key assumptions prior to development [1]. This work takes place primarily before delivery efforts begin, and seeks to answer questions surrounding what a team should build and why they should build it. Discovery versus delivery (learning what to build vs building it) has been treated by many theorists as a central schism in modern product management practice, with increasing focus on discovery as its own key discipline rather than merely a prerequisite to development work.

Recognition of the value of product discovery has appeared in both research settings as well as in professional circles. Ries [6], in his seminal work advocating for the Lean Startup methodology, famously observed that most product organisations faced less risk from not building software than from building the wrong software. Outlined was a framework of “validated learning”, or the process of rigorously testing product hypotheses

about customers and markets prior to beginning work. In doing so, Ries pushed back against the practice of building first, validating with (often costly) market failure second, and instead positioned discovery as an ongoing process.

2.2.2 The Double Diamond Framework

One of the most recognizable frameworks for product discovery introduced in 2005 by the British Design Council, the Double Diamond model for design frames product discovery as four phases: Discover, Define, Develop, and Deliver. These phases are split between two “diamonds” representing periods of divergence and convergence, respectively. During discovery, teams “open up” to learn as much as they can about the problem space through user research, gathering data, and challenging assumptions. During definition, teams “zoom in” on potential solutions before developing and delivering the final artefact [7].

Origins of the Double Diamond have informed its substantial impact on the field of PM practice. As described by authors at the Design Council [7], the initial framework was based on research into how pioneering, innovative companies around the world designed, and extracted the common structure that differentiated them from less-design-centric organisations (i.e. rushing to solutions without adequately understanding problems). An iteration of the Double Diamond was published in 2019 to account for more iterative work that takes place across discovery and delivery, as well as greater conformance to principles of Agile development [8]. While the line between learning and building may not be clearly drawn in many organisations’ processes today, the model highlights the critical need for teams to both understand problems before they define them, and define them before developing solutions.

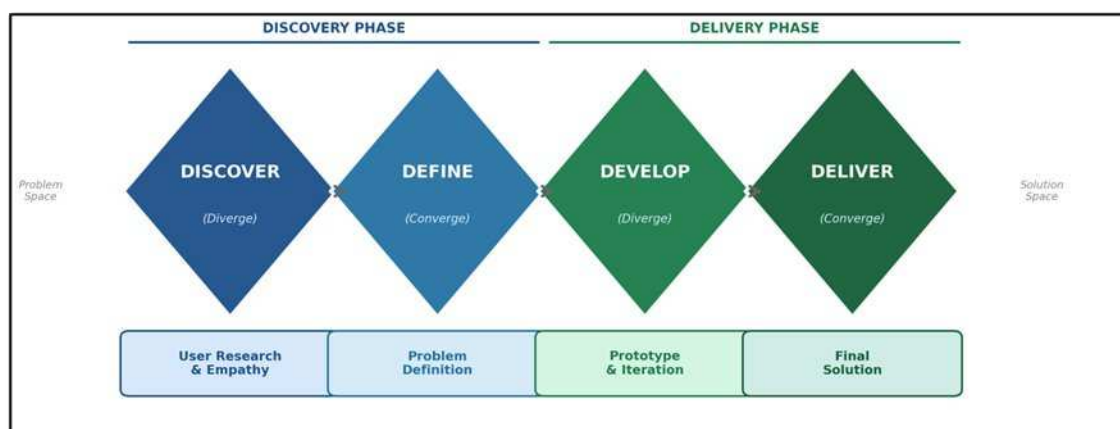


Figure 1: The Double Diamond Framework

Source: Design Council (2019) [8]. Four-phase design thinking model, Discover, Define, Develop, Deliver.

2.2.3 Continuous Discovery

A newer, and generally more influential, conceptualisation of product discovery has been forwarded by Torres [1] in Continuous Discovery Habits. In this framework, Torres characterises continuous discovery as weekly customer touchpoints conducted by the product team, consisting of lightweight research activities performed in service of a clearly defined desired outcome. The key here is continuity: a process of regular, lightweight engagement with research and discovery activities, rather than discrete, periodic bouts of research. This represents an evolution from earlier waterfall-centric thought in which discovery conceived as a discrete phase before development began. Torres makes the case that although Agile methods have done tremendous work to optimise how teams build products (speed and quality of delivery), they say little about what should be built. It is entirely possible for a team to excel at conducting Agile sprints and still build features that provide little or no value to users or business. Continuous discovery is meant to be the yin to Agile delivery's yang: an organised way of knowing that the team is always building the right thing. To that end, Torres presents the Opportunity Solution Tree as a visual model for mapping out the opportunity space (i.e., the problems, desires, and pain points of users that are worth solving) and subsequently linking those opportunities to product features and experiments. The Opportunity Solution Tree is of direct interest to this study, as it represents one example of a structured mental process that may be supported, augmented, or otherwise disrupted by GenAI tools.

More than that, however, Torres's framework matters for this thesis because it provides (1) a specific set of discovery activities against which GenAI tool usage can be mapped, and (2) proof that effective discovery is not simply a data/information problem but a judgment problem. If we accept Torres's framework as definitive, then successful discovery requires the judgment to evaluate which opportunities are most pressing, which pieces of data are most supportive of a given opportunity, and how to choose between feature ideas when tradeoffs must be made. This is the human judgment component that will be disrupted if GenAI tools completely replace these activities, or augmented if they instead exist to support human decision-making. See Section 2.4 for more.

2.2.4 Product Discovery in Agile Contexts

Drawing on the relationship between discovery activities and Agile development practices, some authors have described product discovery through the lens of dual-track Agile [9]. In this construct, discovery and delivery efforts occur in parallel, with the former responsible

for introducing validated ideas into the latter. In doing so, the risk of developing features which may be desirable from an engineering perspective but lack user or business value can be mitigated. Cagan [9] describes the optimal structure of product orgs as one where an empowered product team is responsible not just for delivery against a roadmap, but for continuously discovering user needs.

The preceding discussion is relevant to this study because it positions the PM's role in discovery work as one requiring strategy and business judgement. The work described by Cagan [9] as characteristic of discovery (user interviews, research synthesis, working with teams to ideate solutions, defining/testing hypotheses) aligns strongly with the types of activities GenAI tools have been reported by practitioners to support or "augment". As such, understanding how PMs are adopting GenAI tools into their workflow will require first understanding what effective discovery workflows look like.

2.3 GenAI in Business: Adoption, Productivity, and Organisational Context

2.3.1 The Rise of GenAI

GenAI is a subset of AI technologies that produce new content (text, images, audio, code) after learning from huge amounts of data they were trained on. When OpenAI's ChatGPT was released to the public in November 2022, GenAI tools like it saw massive adoption across every industry in just a few months. What was considered esoteric technology used only in academic and developer circles is now used daily by knowledge workers everywhere. Adoption has been rapid. According to McKinsey's 2023 AI Adoption Global Survey, one-third of survey takers answered that their organisation was already using GenAI regularly for at least one business function. This represents 60 percent of the survey's respondents that had adopted some form of AI technology [10]. Of the business functions where survey respondents noted using GenAI, developing products and services was one of the top reported uses (along with Marketing & Sales and Service Operations). By mid-2025, 79 percent of companies were using GenAI for at least one business function, up from 65 percent in early 2024. Adoption continues to grow despite a rocky start to GenAI's public availability, with some executives voicing concerns around governance, trust in outputs, and other issues [2].

2.3.2 Economic and Productivity Implications

To get a sense of scale for GenAI's economic opportunity, many organisations have conducted analysis into GenAI's potential value creation. In its report *GenAI Creates Economic Opportunity*, McKinsey calculates between \$2.6 trillion to \$4.4 trillion of new value generated across industries and use cases each year, with knowledge work sectors like technology, finance, and professional services being among those most affected [11]. The report goes on to highlight how GenAI has been applied across a variety of business functions, particularly noting the impact of GenAI on white-collar knowledge work. Because GenAI is so effective at automating language, it is more impactful for information workers than physical workers who were most affected by previous technology advances.

Research from Goldman Sachs comes to a similar conclusion regarding GenAI's productivity impacts, estimating its overall effect on global GDP to be around seven percent [12]. While there are varying forecasts about the size of GenAI's impact, for the purposes of this paper it is enough to understand that productivity effects from GenAI will be significant.

These large-scale estimates place GenAI's impact in context: while they are not forecasts per se, if productivity increases of this magnitude are to be realised it will be due in large part to changes at the level of the worker. Product management is among the professions identified by these reports as most affected by GenAI technology, and this study aims to understand how and why GenAI tools get used by PMs. Zooming in on the enterprise level, Bahoo et al. [13] performed a hybrid systematic literature review of 364 AI-related academic articles, from which they identified eight categories of AI-driven corporate innovation literature. Of special interest to this paper are Streams 2 (AI and Product Innovation), 4 (AI and the Innovation Process), and 6 (AI and Knowledge Management), which together can be understood to define what product discovery is. While the authors of this paper study innovation at the corporate level, this paper studies innovation at the practitioner level: how individual PMs use GenAI within Streams 2, 4, and 6 as a part of their discovery work. Finally, at the level of the individual knowledge worker, a body of literature has already developed studying AI's effect on worker productivity. The McKinsey Global Institute revised its previous forecast for work hour automation based on the capabilities of GenAI [10]. Their findings showed that the proportion of work hours that could theoretically be automated by current technologies increased from ~50 percent to between 60-70 percent. Perhaps more importantly, their forecast for when machines will

reliably match median human performance for natural-language understanding tasks previously projected to reach by 2027 was bumped forward to 2023. While previous generations of automation technology would take decades to penetrate markets, GenAI will spread more rapidly. In work directly related to this study, Yun et al. [15] conducted a user study with 16 PMs and knowledge workers at CHI 2025. They found GenAI's main benefit to knowledge work is its ability to help synthesise unstructured data to facilitate decision making, which agrees with qualitative findings from this study regarding use of GenAI for product discovery.

2.3.3 GenAI Adoption in Organisations: Patterns and Barriers

While the GenAI opportunity appears massive on macroeconomic scales, studies have shown that GenAI tool usage within organisations is currently disjointed, unsystematic, and faces significant barriers to entry. Deloitte's recent State of GenAI in the Enterprise (2024) report, which surveyed over 2,800 director-to-C-suite respondents in 16 countries, demonstrates high levels of optimism for GenAI early on in its organisational lifecycle but many leaders felt significant pressure to realise value quickly while managing substantial risks. Governance, talent, and exacerbation of economic inequality were seen as the top risks for organisations (research question around barriers, which this study directly examines, are highlighted in this finding [16]). Furthering this research into small and medium businesses, Pradhan, Dash, Sharma, and Ullah [17] found that barriers to GenAI adoption were significantly more complex for SMEs due to limited resources and lower team member familiarity with AI concepts, which highlights that while similar barriers exist at smaller scales, they do not always present the same way. Intrapersonal individual usage versus enterprise policy mismatches also appear to be commonplace. The same McKinsey [10] study showed that only 21% of respondents working at AI adopting companies had created governance guidelines around use by employees. In practice, this has meant that many PMs using GenAI tools have done so without corporate guidance; they know what they want the tool to do, but not how it works. These findings parallel anecdotal evidence from PMs that GenAI tools are being adopted first on an individual level (particularly by individual PMs) before being adopted at the team or organisational level, which the current survey seeks to quantify.

2.3.4 Trust and the Hallucination Problem

Trust, especially in relation to the so-called hallucination problem inherent to GenAI solutions, was observed as a key barrier to consideration and adoption across professional use cases. Hallucinations are defined as ‘the generation of factually incorrect, fabricated, or inconsistent information by large language models’ that are often presented with high confidence. The National Institute of Standards and Technology (NIST) in the United States defines hallucinations ‘from an AI system as confidently stated but false content’. Where false content might be inconsequential or detected quickly in consumer use cases, hallucinations in an enterprise setting have the potential to lead to violations, financial losses, and reputational damage [18]. Hallucination can directly affect trust in GenAI outputs by professional users. Study [14] identified perceived usefulness as the strongest predictor of AI usage attitude among professional users, however found trust in AI outputs to mediate this relationship. Users encountering inconsistent or incorrect outputs frequently will tend to form negative associations with the tool that override its perceived usefulness and inhibit adoption. Applied to product management, decision-making around product investments supported by user research findings is an outcome-focused exercise. A PM who inputs GenAI with user interview transcripts to summarise findings may accidentally relay mischaracterized user sentiment to engineers, stakeholders, or business leaders by acting on a hallucination.

Hallucination has institutional trust implications as well. Global management consulting firm Deloitte reported in a GenAI study [16] that ‘almost 70 percent of respondents said 30 percent or fewer of their GenAI pilots have moved into production’. This observation suggests that many enterprise customers do not yet trust GenAI outputs as reliable or safe to business objectives at scale. Regarding the present study, this pattern indicates that the barriers category may elicit concerns about hallucinations risk and trust from PMs, particularly those who operate in larger organisations where output accuracy is more likely to be rigorously evaluated.

2.4 Human-AI Collaboration: Theoretical Frameworks

2.4.1 The Technology Acceptance Model

Originally proposed by Davis in 1989 [4], the Technology Acceptance Model (TAM) is the most ubiquitous theory of technology adoption. In TAM, Davis posits that there are two fundamental cognitive factors that influence an individual’s decision to use a technology:

Perceived Usefulness (PU) and Perceived Ease of Use (PEOU). PU is defined as ‘the degree to which a person believes that using a technology will help him or her to perform job-related tasks’, while PEOU is defined as ‘the degree to which a person believes that using a technology will be free of effort’ [4]. PU and PEOU exert influence on user attitude toward using the target technology, which in turn determines their behavioural intention to use it, which ultimately manifests as actual usage behaviour [4]. TAM has since been applied to virtually every professional technology adoption context since its inception, and has been expanded and modified by other scholars in the decades since. For example, Venkatesh and Davis [19] introduced TAM2 to address criticisms of the original model’s exclusion of social factors, adding subjective norm to the model. In a followup paper, Venkatesh et al. (2003) reviewed eight technology adoption models that came before them (including TAM and TAM2) and consolidated them into the Unified Theory of Acceptance and Use of Technology (UTAUT), identifying Performance Expectancy, Effort Expectancy, Social Influence, and Facilitating Conditions as the four core determinants of technology adoption intention, moderated by age, gender, experience, and voluntariness of use [14]. UTAUT has been shown to predict up to 70% of the variance in intention to adopt technology at work. Recently, researchers applied TAM to understand technology adoption specific to AI ActTools and found that consistent with prior research, PU was the best predictor of behavioural intention to use AI ActTools. They also found that AI mindset – an individual’s cognitive orientation toward the promise of AI technology – moderated the relationship between PU and AI adoption intention [20]. This study relies on TAM as its primary adoption lens. The survey questions measuring respondents’ beliefs about the extent to which GenAI will make them more efficient at work (PU) and integrate into their workflow (PEOU) map directly to the original constructs of TAM. Additionally, this study will cross-tab responses to these and other survey questions by both role seniority and company size. This is done in consideration of UTAUT’s moderating variables of subjective norms and facilitating conditions, which among professionals can include an organisations’s governance frameworks surrounding tools as well as usage behaviour by peers. If senior PMs report lower trust in GenAI despite believing it will be useful at work, TAM would predict a decrease in adoption intention, a hypothesis that the survey data will enable this study to evaluate.

Table 2.2: TAM and UTAUT, Theoretical Constructs and Application to This Study

Construct	TAM (Davis, 1989) [4]	UTAUT (Venkatesh et al., 2003) [14]	Application in This Study
Core adoption predictor	Perceived Usefulness (PU)	Performance Expectancy (PE)	Q9–Q11 Likert scale: efficiency and output quality
Effort dimension	Perceived Ease of Use (PEOU)	Effort Expectancy (EE)	Q8: ease of GenAI workflow integration
Social influence	Subjective Norm (TAM2) [19]	Social Influence (SI)	Peer adoption in PM communities
Organisational support	Not explicitly modelled	Facilitating Conditions (FC)	Tool governance, plan-tier access constraints
Moderating variables	Experience, voluntariness	Age, gender, experience, voluntariness	Role seniority and company size (Q1, Q3)
Explained variance	~40% of adoption intention	Up to 70% of adoption intention	Cross-tabulation analysis (Chapter 5)

2.4.2 Augmentation versus Automation: Daugherty and Wilson

Second to TAM is the augmentation framework from Daugherty and Wilson, authors of *Human + Machine: Reimagining Work in the Age of AI* (2018). Citing research from 1,500 organisations, Daugherty and Wilson claimed that AI replacing human workers was the wrong narrative. From their research they distilled that AI implemented correctly does not replace humans but augments them [5].

Specifically, they coined the term 'missing middle' to account for the human-machine relationship. Humans and AI are not at opposite ends of a spectrum. Between fully human and fully automated tasks lies a hybrid grey area. It is in this area that AI provides large data analysis capabilities and humans can provide inputs that require emotional, societal context. The authors identified three categories of human-machine relationship in this middle grey area: amplification (AI augments what humans can do by analysing large quantities of data to provide insights), interaction (AI allows humans to interface with complex systems using more natural forms of communication) and embodiment (AI allows humans to extend their physical and cognitive capabilities) [5]. Daugherty & Wilson's subsequent 2024 book on leading in the age of AI reinforces these ideas and introduced six new hybrid human-machine jobs as well as a leader's guide for the AI age. The authors reinforce their central argument that human skill and AI capability are most valuable when combined rather than competed [21].

This theory will help answer SRQ4 which asks PM practitioners how they view GenAI, as augmentation to their role or as a replacement of parts of their role. Question 12 ('GenAI augments my work as a PM rather than replacing parts of my role') on the survey is based

on Daugherty and Wilson's conceptualisation of AI and humans working together vs separately. Findings relating to this question will be discussed in the context of augmentation theory in Chapter 5. Rai and Rai's [22] research on AI in software product teams supports the decision to frame the human-AI relationship as one of collaborator vs replacement. They found that teams are in many cases choosing to divide cognitive labour between AI and humans such that AI is responsible for task decomposition while humans focus on interpersonal coordination.

2.4.3 The AI Fluency Framework: Dakan and Feller

Professor Rick Dakan (Ringling College of Art and Design) and Professor Joseph Feller (University College Cork) have developed a framework called AI Fluency in ongoing research partnership investigating the convergence of AI, creativity, innovation, and learning. The framework has been developed in collaboration with Anthropic and can be accessed for free (via a Creative Commons licence) via Anthropic Academy and the AI Fluency Framework site [3].

AI Fluency has been described as 'meeting a standard of working effectively, efficiently, ethically, and safely in the emerging modalities of Human-AI collaboration' [3]. These modalities (effectiveness, efficiency, ethical behavior, and safety) are framed not as four goals that may or may not be met, but four dimensions of a unified standard of excellent practice in human-AI collaboration. Towards that standard of excellence, the AI Fluency Framework defines four core competencies, known as the 4Ds, which span the breadth of skills required for fluent human-AI collaboration.

The first core competency is Delegation, which centres on the question of how and when to engage with AI. Delegation requires practitioners to understand the three modes of Human-AI interaction (Automation, Augmentation, and Agency) and to choose which mode best suits the task at hand [3]. Applied to the work of a PM, Delegation skill would answer questions such as: Which discovery activities can be delegated to GenAI (e.g., initial literature review, first draft synthesis)? Which should be performed in an Augmented partnership with AI (e.g., opportunity mapping, idea generation)? Which must be performed without AI at all (e.g., empathic user interviews)?

The second core competency, Description, centres on the mechanics of communicating intent to AI systems. Description includes the practitioner's ability to translate their desired output into a clear prompt, provide appropriate context, and indicate an appropriate collaboration style. Feller [3] distinguished cleanly between prompt 'engineering' as a

tactical discipline with shallow technical demands and the deeper fluency that accrues from mastering Description: the ability to translate professional objectives into terms that an AI system can use to generate valuable output. Considered in the context of product discovery, Description competency would be the difference between prompting a GenAI system for insights that take your understanding of users to the next level, and obtaining a generic, one-size-fits-all summary that misses key facets of your unique product context.

Discernment, the third core competency in the AI Fluency Framework, concerns the evaluation of AI outputs. On whom does the responsibility lie for judging the quality, accuracy, safety, bias, and need for improvement of AI-generated artifacts? For Discernment practitioners, it lies first and always with the human. High Discernment is required to recognize the occasions where AI has failed to provide pertinent context, where generated ‘facts’ should be double-checked, and where the AI ‘hallucinated’ output that may be nonsensical or outright false (Section 2.3.4). The AI Fluency Index (Swanson et al., 2026) [20] found that in conversations where artefacts were created, users demonstrated markedly lower rates of fact-checking behaviours, suggesting that the apparent coherence of AI-generated outputs may actively suppress the critical evaluation that Discernment requires. The neatness and coherence of AI outputs may actually discourage the sorts of critical human oversight that Discernment entails.

Last, Diligence refers to the ethical choices made around which AI tools are used, whether and how those tools are disclosed to others, and whether humans take responsibility for AI-generated outputs created during the course of their work [3]. In the practice of product management, Diligence comes into play whenever a PM passes an AI-generated artifact along to someone else: sharing a synthesis of user research reviews with a product team, presenting competitive analysis prepared with the help of GenAI to stakeholders, or adopting GenAI-written docs as the requirements against which engineering will be measured. Does the PM validate the accuracy of these materials? Do they disclose the use of GenAI tools? Ultimately, do they take full accountability for the work they present as their own? The level of Diligence exercised by the PM will be apparent in the answer to each of these questions.

The AI Fluency Framework offers this study a complementary theoretical lens through which to consider GenAI use by PMs. Whereas TAM provides the study with a theory of Practical Intention and the Daugherty & Wilson literature provides a theory of augmentation versus substitution, AI Fluency provides a framework for understanding how well PMs are working with GenAI. It is one thing to adopt ChatGPT as a tool for creating

user personas; it is another to paste those personas unmodified into a deliverable. A PM who prompts ChatGPT to produce user personas and accepts the results with no human oversight or modification has made a Delegation decision that suggests low Discernment. But a PM who crafts context-rich prompts (Description), vets model outputs against their own understanding of real user data (Discernment), and notes “drafted with ChatGPT’s assistance” on every piece of AI-informed work shared with their team (Diligence) is interacting with AI at a different level of sophistication. The 4D framework makes these distinctions analytically tractable. The author completed Anthropic's AI Fluency: Framework and Foundations certification (2025), providing direct experience of the framework as applied to professional workflows.

Table 2.1: AI Fluency Framework, Four Core Competencies (Dakan & Feller, 2025)
[3]

Competency	Core Question	Application in Product Discovery
Delegation	"What should I give to AI?"	Identifying which discovery tasks are suitable for AI vs. human judgment
Description	"How do I instruct AI clearly?"	Writing precise, context-rich prompts for research synthesis and PRD drafting
Discernment	"How do I evaluate AI outputs?"	Critically assessing accuracy, bias, and depth of AI-generated insights
Diligence	"How do I use AI responsibly?"	Ensuring data privacy, transparency, and ethical use in customer research

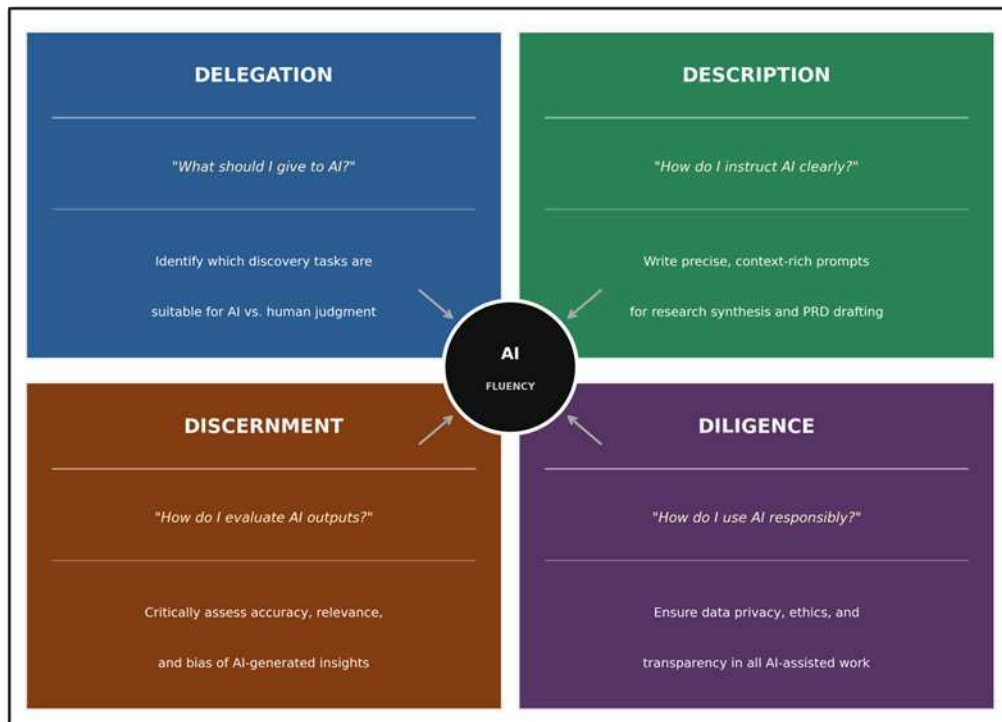


Figure 2: The AI Fluency Framework, Four Core Competencies

Source: Dakan & Feller (2025) [3]. The four competencies (Delegation, Description, Discernment, Diligence) define AI fluency for product practitioners.

2.5 The Evolving Role of the PM

2.5.1 The PM's Core Responsibilities

Increasing focus has been placed on the role of the PM as the profession has grown and evolved. Originating in brand management for consumer goods products, the PM role has had its definitive responsibilities outlined by many. Cagan [9], author of *Inspired: How to Create Tech Products Customers Love*, describes the PM as being responsible for the product team delivering a product that is both valuable (users actually choose to buy or use it) and viable, the solution works. Accountability for both value and viability is the PMs cross to bear and distinguishes their role from other roles like project manager, business analyst, or product owner.

Springer and Miler [23] analysed the role of PM and identified that much like this study, there is not one singular way to be a PM. At larger companies, PMs often have separate strategists, discoverers and delivery partners. At smaller companies, especially startups and scale-ups, the PM will likely be responsible for interviews, user research, moderating design workshops, writing engineering tickets, stakeholder management etc. This finding

is significant to this research as the implementation of GenAI will have different effects on a startup PM who may do majority of the discovery work compared to an enterprise PM who has limited involvement in this stage.

2.5.2 The Impact of Agile on Product Management

Product Management as a professional discipline has been significantly affected by Agile adoption across the software industry. Adoption of Scrum's Product Owner role (which places primary responsibility for backlog management, sprint planning, and stakeholder prioritisation on a single individual) left some ambiguity about the overlap between product management and Agile development [9]. For years, Marty Cagan has argued that Scrum's Product Owner is only a subset of what product management should be - namely, that governance of the Agile process is an important component of PM work, but not sufficient to produce long-term product value. Performing the function of PM without deeply engaging in discovery, user research, and opportunity analysis is, in Cagan's framework, manufacturing features rather than creating value.

Agile adoption and associated cultural/contextual pressures are key considerations in predicting GenAI adoption by PMs. Scrum, as described by the Scrum Alliance in their State of Scrum Report [24], is the dominant Agile framework in use today with most organisations adopting it as their sole Agile method or in combination with others. PMs who work in Scrum environments are often spread thin managing ceremony (daily standups, sprint reviews, backlog refinement sessions, sprint planning) as well as the independent work that discovery entails. If there is a class of tool that can meaningfully reduce the time required to synthesise research, conduct competitive analyses, or generate documentation, it seems natural that PMs would want to adopt it. Does utilising these tools to trade hours reduce quality? Increase speed at the cost of accuracy? These are some of the empirical questions this study seeks to answer.

2.5.3 GenAI and the Emerging PM Skill Set

The introduction of GenAI technologies has also triggered discussion among PMs about how skill requirements might change for product professionals. Emerging consensus suggests there is a set of emerging competencies (sometimes described as AI-adjacent skills) that distinguish PMs who use GenAI effectively. Several recent studies of PMs' working with GenAI support the need to answer this question: Yousif [25] and Singh [26] explore how GenAI is changing PM workflows, while Jon-Onwubuoya's [27] research

explains how localised factors influence whether PMs approach their work with AI critically or naively.

Skill with constructing Prompts (Prompt engineering: technical instructions that tell an AI system what we want it to do) has been one of the most commonly discussed skills required for effective GenAI use. McKinsey's 2023 AI report [10], highlights that 7 percent of respondents working at organisations that had adopted AI reported hiring for prompt engineers at some point in the last year, implying a desire to staff this skill that was not present before generative models. The AI Fluency Framework [3], however, paints a more complex picture. Prompt engineering, Feller [3] explains, is not the same skill as Description. One can write great prompts without having a firm grasp of how to describe the information needs of the user to a computer. Conversely, someone who understands how to describe detailed requirements in a usable form for AI will inherently know how to craft good prompts.

Evaluating AI outputs (deciding if the information an AI system gives you is correct, relevant, biased, etc.) has been cited by experts and empirically shown to be one of the skills most lacking in today's PM community. Elevating this skill is directly related to the Fluency Framework's dimension of Discernment and ties back to the risk of hallucinations described in Section 2.3.4. Preceding research featured in this literature review has found a strong correlation between AI Fluency (measured as a high-level understanding of what AI can do across all categories of the Framework) and effective Discernment: practitioners who underestimate the generating capabilities of LLMs were far more likely to overtrust the responses they provide [3].

Analysis and judgment (skills which have traditionally been at the core of PM work) are anticipated to become more important as GenAI begins to automate portions of the information gathering/analysis process. This sentiment mirrors predictions by Daugherty and Wilson (2018) that long-term AI adoption will make "advanced human skills" more valuable as AI systems take over more of the lower-level cognitive work.

2.6 Synthesis: Towards a Conceptual Framework

Recall the four pillars supporting this literature review. Each has introduced a set of conceptual links that apply to the research questions guiding this thesis.

1. Product discovery is framed as a series of information work practices involving high cognitive load and uncertainty; furthermore, prior literature in this space converges around a shared set of discrete, repeatable activities (user research synthesis/opportunity mapping,

idea generation/ideation, hypothesis testing, output documentation) that can be treated as individual opportunities for augmentation through GenAI.

2. GenAI technology and its relevant applications are rapidly being adopted across industries and professions due to gains in perceived productivity; however, enterprise adoption is limited by trust and hallucination issues as well as governance overhead.

3. Prior theories and models of human-AI collaboration (ranging from TAM and UTAUT's focus on factors like perceived usefulness and ease of use, social influence, and facilitating conditions, to Daugherty and Wilson's reframing of the discussion around "augmentation," to the proposed 4D competency framework from Dakan and Feller's AI Fluency Framework) provide a multidimensional lens for interpreting not just whether PMs are using GenAI, but how well and how responsibly they are doing so.

4. The PM role has been identified to involve core competencies associated with strategic visioning/user empathy, but is now shifting to also emphasise more AI-related skills, including Discovery skills like those articulated in the Description/Discernment/Diligence components of the 4D framework above.

A key insight from mapping this literature is that no studies, to the author's knowledge, have been conducted at the practitioner level to measure how and why GenAI is being used (or not used) specifically for product discovery activities. There are broader-scale surveys around AI adoption at the enterprise level [10][16], as well as theoretical treatises around human-AI collaboration [4][5][3]. However, these works fail to consider the specific question of how PM professionals are working with GenAI tools during the specific types of information work described above. The primary exception is Ghazi and Mariam [28], who conceptually defined how GenAI tools could fit within the overall PM lifecycle; however, their framework is again at the process level and was not validated against responses from practising PMs. This gap in the literature is the area of originality that this thesis seeks to fill with empirically-derived insights from a practitioner survey around how PMs are currently using GenAI.

Finally, taken as a whole the four pillars present here also seem to imply that there will be variance in how practitioners use GenAI tools, and that more use will not necessarily mean better discovery outcomes. Adoption as characterised by exercise of the 4Ds of Delegation, Description, Discernment, and Diligence will likely play a role in practitioner success using these tools just as much as adoption rate itself. This hypothesis will inform the GenAI Adoption Maturity Model introduced in Chapter 6, which posits that teams evolve

not simply by using more AI tools but by using them with increasing sophistication, criticality, and accountability.

Table 2.3: Literature Review, Research Gaps Addressed by This Study

Research Area	Existing Literature Coverage	Gap Addressed by This Study
GenAI in the Workplace	Broad productivity effects, automation risk [10][11]	PM-specific discovery workflows, not studied empirically
Product Discovery Frameworks	Torres [1], Double Diamond [7][8], Dual-track Agile [9]	How GenAI maps onto each discovery activity stage
Technology Acceptance (TAM/UTAUT)	General IT adoption [4][14]; some AI-TAM studies	TAM applied specifically to GenAI in PM practice
Human-AI Collaboration	Augmentation theory [5][21]; AI Fluency Framework [3]	4D fluency model applied to product discovery context
PM Role Evolution	Agile impact on PM [9]; competency debates [54]	Empirical evidence of skill shifts driven by GenAI
Organisational GenAI Adoption	Enterprise barriers and governance [16][10]	Barriers specifically within product team settings



Figure 3: GenAI Adoption Maturity Model for Product Teams

Conceptual model derived from literature synthesis. Five levels from No Adoption (Level 1) to Strategic capability (Level 5), with estimated distribution of product teams across levels.

2.7 Chapter Summary

Chapter 2 reviewed foundational literature across four key thematic areas informing the research questions of this thesis. The product discovery literature review, built around Torres [1], Ries [6], and Cagan [9], defined product discovery both as a practice and as a categorisation of specific activities to which AI tools may be applied. The GenAI in business literature review synthesised McKinsey [10] and Deloitte [16] research on enterprise AI adoption rates, uses cases, and adoption challenges. The human-AI collaboration literature review highlighted academic understanding of why and how people adopt new technology (1989 Technology Acceptance Model from Davis [4] and 2003 Unified Theory of Acceptance and Use of Technology from Venkatesh et al. [14]) evolved towards human-AI teamwork research (2018 perspective on AI augmentation from Daugherty and Wilson [5]) and eventually founded skill frameworks for AI collaboration (2025 AI Fluency Framework from Dakan and Feller [3]). Finally, the product management role evolution literature review provides evidence that PM practitioners today find themselves negotiating a shift in skill expectations that include both new AI fluencies and their traditional expertise.

Taken together, the four pillars presented in this chapter support the framing of the empirical research study detailed in Chapter 3. The methodology for which will be introduced along with explanations for the design of the survey and case study literature review processes used to gather data informing the answers to the research questions established in Chapter 1.

Chapter 3: Research Methodology

3.1 Introduction

This chapter describes how this study was planned and explains why the various methodological decisions described are well suited to answering the research questions outlined in Chapter 1. The purpose of a methodology chapter is not to merely enumerate decisions, but to explain them so that the reader can judge whether the evidence presented in subsequent chapters will be fit-for-purpose – that is trustworthy, relevant, and adequate to support the conclusions presented. All of the decisions described below (everything from research philosophy to survey question phrasing to data analysis techniques) were made with a purpose in mind and can be defended on logical and/or practical grounds.

This chapter is divided into nine main sections. Section 3.2 describes the research philosophy. Section 3.3 outlines the research design. Section 3.4 explains the sampling strategy. Section 3.5 presents the survey instrument in full. Section 3.6 outlines the case study analysis. Section 3.7 describes the data analysis approach. Section 3.8 discusses ethical considerations. Section 3.9 acknowledges limitations. Section 3.10 summarises the chapter.

3.2 Research Philosophy

3.2.1 Pragmatism

Research philosophy encompasses the beliefs that a researcher has about the nature of knowledge (realism vs. interpretivism) and how best to acquire it (induction vs. deduction). A researcher's philosophy will invariably influence each of the methodological decisions they take. This research adopts something called a pragmatist philosophical orientation. Pragmatism as a philosophy asserts that there is no single correct approach to research and that the main criterion for selecting a research strategy should be fitness for purpose [29]. Simply put: whatever research methods allow you to answer your research questions most completely and accurately are the right methods to use, regardless of whether they are traditionally quantitative or qualitative.

There are three reasons why pragmatism is suitable for this study. The first reason is that the research questions are aimed at exploration. In other words, we know relatively little about how PMs in particular currently integrate GenAI into product discovery processes. Therefore the study cannot begin with well-developed hypotheses. It must both

establish descriptive facts about how things work now and analyse people's understanding of and interpretations of those facts. Asking research participants to provide their thoughts directly requires qualitative data, but the study also clearly benefits from quantitative analysis. Fulfilling both of these objectives requires a mixed-method approach. The second reason to choose pragmatism is that it allows this researcher to draw on whichever theoretical frameworks are most illuminating, rather than being constrained by a single paradigm. Finally, pragmatism is conducive to application: the intent of this research is to produce findings that will be valuable to practitioners and their organisations, not just add fodder for academic debate.

3.2.2 Mixed-Method Design Rationale

Mixed-method research designs 'merge' quantitative and qualitative data within the same study, enabling the researcher to take advantage of both [29]. The quantitative component of a mixed-method study (generally obtained via survey questions with closed responses such as rating scales and forced-choice selections) is used to identify patterns, measure behaviours and attitudes, and compare subgroups within the sample. The qualitative component (generally obtained from open-text survey responses and interviews/case studies) seeks to understand respondents' experiences, reasoning, and circumstances which cause or explain those patterns.

The two types of data complement each other and will allow this study to address the research questions more completely. Quantitative data showing that e.g. 70% of respondents report using ChatGPT as part of their discovery activities is more meaningful when paired with qualitative data describing exactly how practitioners use that tool, what they like about it, and where it falls short of their needs. Data from each method will validate and enrich the other. This approach to data collection and analysis is known as convergent parallel mixed-method design [29]. Quant and qualitative data will be collected simultaneously (through the survey), analysed separately, and only brought together in discussion of the results.

3.3 Research Design

The overall research design comprises two components. The first and primary component is a structured survey of practicing PMs and Product Owners. The second component is a review and synthesis of existing case study materials from companies who have publicly

documented their use of GenAI in product workflows such as Figma, Notion, Miro, and Atlassian. Both of these components will make use of mixed-method data collection.

Both components of this study were designed to answer the research questions as fully as possible. The survey does so by gathering uniform data from each respondent. The case study review adds breadth to the data available by performing a triangulation step – that is, by collecting an independent dataset from a different level of analysis (organisation rather than individual) against which to compare and validate the survey results.

3.4 Sampling Strategy and Target Population

3.4.1 Target Population

PMs and Product Owners working in technology companies. PMs and Owners were selected as the population of interest for this study because they have primary responsibility for leading and delivering on discovery activities (conducting user research and synthesising findings, identifying business opportunities, brainstorming potential solutions, and defining product requirements). As such, they are the professionals likeliest to have first-hand experience with using GenAI tools during discovery. While this study will focus on PMs and Owners who work in technology companies of all sizes across all industries related tech and at any level of seniority, it will record information on geographic region to allow for analysis of any regional differences in adoption.

3.4.2 Convenience Sampling

The survey was sent to potential participants using convenience sampling methodology: inviting those to whom the researcher had access via professional networks rather than selecting from a randomized or stratified sample [30]. The survey link was posted on the researcher's LinkedIn network and shared amongst PM communities and discussion groups, and directly with practitioners whom the researcher had connections with through their Erasmus academic network in Italy and Poland. The survey was also promoted by the student's thesis supervisors at UNIVPM and Gdańsk University of Technology.

Convenience sampling is standard methodology at the master's thesis level of inquiry, especially for investigations into working professionals where membership registries are not generally publicly maintained. Convenience sampling also adheres to the pragmatist paradigm upon which this study is based: instead of striving towards a theoretically ideal methodology that is impossible given resource constraints, the researcher pragmatically

does what is possible and clearly outlines the methods they used. It is acknowledged that the sample is not probabilistic, and that the findings cannot be treated as statistically representative of the global PM population. The goal of this study is analytical generalisation (finding concepts and meaning that are useful outside of the sample analyzed) rather than statistical generalisation [31].

3.4.3 Sample Size

This study anticipated gathering at least 80 responses from participants. This sample size is large enough to support the types of analysis this study requires in this mixed-method master's thesis: both quantitative descriptive statistics and qualitative thematic analysis. Quantitatively, a sample of $N > 80$ allows meaningful cross-tabulations by subgroup while acknowledging indicative rather than statistically generalisable results. Qualitatively, research on thematic saturation indicates that themes typically emerge clearly within 80 to 82 open-text responses, making a target of 80 or more respondents comfortably sufficient [32]. Outreached over 300+ professionals to allow for attrition due to non-response.

3.4.4 Inclusion Criteria

Participants were required to hold (or most recently hold) a position where the primary responsibility was product management or product ownership. Participants were also required to complete all mandatory questions through Section D of the survey. Only responses that returned at least demographic information (Section A) and one additional section of the survey were kept. Such partial responses were ineligible for analysis segments that required fully completed surveys, such as cross-tabulations including free-text questions.

3.5 Survey Instrument Design

3.5.1 Instrument Development and Platform

This survey was developed in stages based on both the literature presented in Chapter 2 as well as the information needed to answer each sub-question identified in Section 3. The researcher tested the draft survey questions informally on colleagues who provided professional insight into product management. Minor adjustments were made to formal question text and answer choice labels based on this feedback. The final survey was created using Microsoft Forms and deployed as a secure online questionnaire distributed

through a shareable link. Microsoft Forms was chosen from alternatives like Google Forms or Typeform for its clean interface, native Office 365 compatibility, and robust export features. This study took place in a European university setting and the survey author works primarily with Microsoft 365, making Forms a natural fit.

Respondents were expected to spend 10–12 minutes. Survey length was determined by aiming for minimal time necessary to collect meaningful data while not unduly taxing the time of professional respondents. This study used non-random recruitment methods (chiefly direct messages on LinkedIn) to invite participants without offering incentivization, creating extra friction in the decision to participate. The survey introduction also stated that completion would not result in requests for follow-up calls, which can frequently dissuade time-constrained practitioners from participating.

The survey contains 19 questions presented in six distinct sections. Question types include single-choice and multiple-choice, five-point Likert scale rating questions, and open-text responses. One question (Q4) was designated as an optional question and was used purely for segmenting results should sufficient sample size be reached.

3.5.2 Survey Structure Overview

The table below breaks down the survey questionnaire by section and cross-references each question to its corresponding sub-question(s).

Table 3.1: Survey Structure and Alignment with Research Sub-questions

Section	Title	Questions	Research Sub-question Addressed
A	Respondent Profile	Q1–Q4	Segmentation variables (role, experience, company size, region)
B	GenAI Adoption & Usage	Q5–Q8	SRQ1, Which GenAI tools and use cases have been adopted?
C	Perceived Value	Q9–Q11	SRQ2, What value and business outcomes do PMs perceive?
D	Barriers & Trust	Q12–Q13	SRQ3, What barriers limit effective adoption?
E	Impact on PM Role	Q14–Q15	SRQ4, How is the PM role and skill set evolving?
F	Open-Ended Questions	Q16–Q19	Qualitative depth across all four research sub-questions

3.5.3 Section A, Respondent Profile (Q1-Q4)

Section A consisted of four questions(Q1-Q4) designed to collect basic demographic data about each respondent.

- Q1 asks about the current position title of the respondent. Options include PMs, Senior PMs, Product Owner, Head of Product / VP of Product, Product Lead, and Other.
- Q2 asks how many years of experience the respondent has working in product management.
- Q3 asks how many employees their current company has. Responses were bucketed into four categories: Startup (up to 50 employees), Scale-up (51-250 employees), Mid-size (251-1,000 employees), and Enterprise (more than 1,000 employees).
- Q4 asks respondents to select their region. Q4 is optional.

Region was collected because the survey linked was distributed via personal networks on LinkedIn as well as global PM communities. Despite purposive sampling efforts to focus respondents from Europe, the survey attracted responses from respondents located across North America (US & Canada), Latin America & Caribbean (LATAM), Europe, Asia-Pacific (APAC), Middle East & Africa (MEA) and Australia & Oceania.

3.5.4 Section B, GenAI Adoption and Usage (Q5–Q8)

Section B tackles the first research sub-question (Which GenAI tools and use cases have been adopted in discovery workflows?) with four questions.

- Q5 asks respondents to pick from a multiple-choice list of ten popular professional GenAI tools as of 2025–2026 (ChatGPT, Google Gemini, Microsoft Copilot, Notion AI, Miro AI, Claude, Perplexity AI, Otter.ai, Productboard AI, None yet, Other) which tools they currently use in their product work.
- Q6 then asks respondents to pick from a list of ten product discovery activities (synthesising user research, building personas or segments, ideation/brainstorming, researching/comparing competitors, writing PRDs/specs, writing user stories, preparing for user interviews, analysing quantitative data, summarising meeting notes, Other) which activities they currently use GenAI tools for.

Note: Respondents who answered Q5 with 'None yet' skip to the subsequent section. Q6 includes an explicit option for respondents who use GenAI in discovery but not for any of the listed activities. By mapping reported GenAI usage onto the concrete activities that make up the product discovery process according to the literature review, this question allows the study to drill down beyond the binary question of whether PMs are using GenAI to understand precisely where in the discovery workflow they are applying it.

- Q7 asks about usage frequency across five options (Daily, Several times per week, Weekly, Occasionally, Rarely or never)

- Q8 is a five-point Likert-type agreement statement ('Using GenAI tools has improved my efficiency in product discovery') measuring perceived efficiency impact. Q8 operationalises the Perceived Usefulness construct from TAM and the Performance Expectancy construct from UTAUT[4][14], both of which theorise job performance improvements as the key driver of technology adoption.

3.5.5 Section C, Perceived Value (Q9–Q11)

Section C builds on the perceived value questions above with three additional questions.

- Q9 is a Likert agreement statement measuring perceived impact on output quality: 'GenAI improves the quality of my product discovery outputs (e.g. insights depth or idea quality)'. This question distinguishes between the efficiency dimension covered by Q8 and a quality dimension, which literature from the review suggests practitioners may not perceive GenAI as addressing despite acknowledging that it saves them time.
- Q10 is a standard Likert statement measuring perceived ease of use ('I find it easy to integrate GenAI tools into my workflow'). This question operationalises the Perceived Ease of Use (PEOU) construct from TAM and the Effort Expectancy construct of UTAUT [4] [14] and is expected to see some divergence by company size and seniority bracket, with enterprise respondents potentially indicating lower ease of integration due to governance hurdles and tool approval processes.
- Q11 then asks respondents to pick up to three specific benefits they have experienced from the multiple-choice options provided. The option list faster synthesis of research data, more or faster idea generation, better-quality documentation, improved decision-making, time saved on manual tasks, and help with organising or structuring thinking. By instructing respondents to choose up to three rather than allowing them to select all that apply, this question pushes them to identify the benefits they feel are most salient rather than selecting every option they can conceive, helping to produce a cleaner signal of highest-value use cases.

3.5.6 Section D, Barriers and Trust (Q12–Q13)

Section D addresses the final research sub-question (What challenges and barriers (individual, organisational, or technological) are preventing product teams from adopting GenAI effectively in their discovery processes?) with two questions:

- Q12 asks respondents to pick up to three barriers to using GenAI for their product work from a list of nine options. The options include: GenAI tools produce inaccurate or hallucinated outputs; I have data privacy and security concerns when using GenAI tools; Company policy prevents or limits my use of GenAI tools; I lack the skills or training to utilise GenAI tools effectively; GenAI tools do not have knowledge of my product's unique context; It's difficult to integrate GenAI tools into my workflow; I haven't seen a clear ROI or impact from using GenAI tools; My team members are reluctant to use GenAI tools; and None of the

above. The list of barriers was taken directly from the findings from the literature review Sections 2.3.3 and 2.3.4 to ensure that each option is firmly rooted in observed barriers to enterprise GenAI adoption, rather than proposing speculative barrier categories that may not be top-of-mind for practitioners.

- Q13 is a Likert agreement statement ('I trust the outputs from GenAI tools in my product discovery work') which captures the trust dimension of GenAI usage. The question was positioned in Section D rather than Section C because while trust does relate to perceived value, it also relates to willingness to take action on AI-generated outputs without verbose checking or consideration. As such, it falls more naturally into a section about barriers to adoption than a section about the benefits of said adoption. Cross-tabulating the results of Q13 with Q3 (company size) and Q7 (usage frequency) should produce particularly instructive results on the adoption-trust relationship.

3.5.7 Section E, Impact on the PM Role (Q14–Q15)

Section E covers the fourth research sub-question with two closed-response items.

- Q14 is a Likert agree/disagree statement ('I see GenAI as augmenting my work as a PM rather than replacing parts of my role') which directly operationalises the augmentation vs automation dimension of Daugherty and Wilson's (2018) framework and relates back to the Delegation competency of the AI Fluency Framework [3]. The average score on this item, and its distribution within role and seniority subgroups, will allow the researcher to substantiate or refute the assertion that GenAI is viewed by practitioners as augmenting rather than displacing the PM role.
- Q15 is a multiple-choice question asking respondents which PM skills have become more important since GenAI became widely available, with options including critical thinking and judgment, prompt engineering, AI and ML literacy, strategic decision-making, stakeholder communication, interpreting data and AI outputs, empathy and user insight, and no significant change. This question connects to the PM role evolution literature reviewed in Section 2.5.3 and to the Description and Discernment dimensions of the AI Fluency Framework.

3.5.8 Section F, Open-Ended Questions (Q16–Q19)

Section F comprises four open-text responses that will lend the qualitative depth necessary to contextualise and augment quantitative insights gained from Sections A-E.

- Q16 asks the respondent to provide a concise description of one way they have used a GenAI tool during product discovery (what went well / did not for you.). This question was written to elicit the most analytically useful data for qualitative analysis: specific, anecdotal examples of respondent GenAI use generate richer, more informative qualitative data than general, unstructured opinions. The specific-example prompt framing, rather than general 'evaluation', is intentional and in line with Braun and Clarke's (2006) advice that increased depth and usefulness of open-text survey answers can be achieved through careful prompt wording.

- Q17 asks for the one limitation of GenAI that has had the biggest impact on the respondents product discovery work. By asking about the single biggest limitation, rather than however many come to mind, the qualitative barrier data gathered here will directly complement and elaborate upon the categorical barriers data collected in Q12 with qualitative explanation of how specific barriers play out in practice.
- Q18 requests, in the respondents own words, how GenAI has changed product management for them. This question is designed to elicit behavioural-change evidence, not opinion.
- Q19 gives the respondent an opportunity to share anything about GenAI + product discovery that they feel was not covered in this survey. Including an open-ended question inviting unstructured final comments at the conclusion of qualitative surveys is standard practice and allows the researcher to collect unexpected but often illuminating insights from respondents with strong opinions on the survey topic [30].

3.6 Case Study Selection and Analysis

3.6.1 Rationale

Case study analysis was included in this research design for two reasons. Firstly, data triangulation: viewing GenAI adoption from the perspective of entire organisations, as documented through their public-facing commentary on the experience, allows the researcher to compare and contrast individual responses with how teams are actually using GenAI in practice at leading companies. Positive survey responses about GenAI's capabilities are given more weight if high-profile, successful product companies are publicly reporting similar experiences. Conversely, points of divergence between company-reported challenges or use cases and the aggregate survey data are themselves interesting from an analytic perspective. Secondly, theory-building / enrichment: Company narratives discussing their implementation and adoption of GenAI often include detail on aspects of the product management technology interface which are impossible to collect via a survey distributed to individuals.

3.6.2 Company Selection

The decision to include Figma, Notion, Miro and Atlassian was motivated by three criteria. Firstly, each company had publicly and prolifically documented their GenAI integration efforts, ensuring a robust dataset of secondary source material was available for the researcher. Secondly, each company's primary product offering maps onto or enhances workflows commonly facilitated by PMs. Using categories relevant to the four key PM workflows identified in the literature review, these are Figma for collaborative design /

discovery whiteboarding; Notion for PM documentation / knowledge management; Miro for virtual discovery workshops / opportunity mapping / brainstorming; and Atlassian (specifically Jira and Confluence) as the most widely adopted project management / documentation platforms for Agile product teams. Thirdly, the four companies vary in size, scope, and market; this selection provides useful diversity. Company information was gathered from sources publicly posted by the companies themselves. This includes official blog posts, product announcements, engineering blogs, and press materials. No individuals from these companies were contacted and no proprietary information was used.

3.6.3 Data Collection Procedure

Secondary data was gathered by aggregating content published by each company on its official blog, product updates page and press releases pages that mentioned AI or GenAI functionality, alongside technology journalism and practitioner media articles covering these releases. A summary was then compiled for each company of: GenAI features/capabilities released; the discovery/product PM workflows said to benefit from them; the stated purpose for releasing them; and any user response/outcome data publicly available. These tables are included in Chapter 4 with the survey results and will be used in Chapter 5 to support the triangulation discussion.

3.7 Data Analysis Plan

3.7.1 Quantitative Analysis: Descriptive Statistics

Descriptive statistics allow a dataset to be summarised, understood and described numerically in a way that is clear and directly interpretable by the reader [33]. For all categorical questions (single-choice & multiple-choice), frequency is the primary statistic; calculated by determining how many respondents selected each response option and what percentage of respondents this represents. Likert-scale questions which were answered on a scale from one to five will be summarised using both the mean response score across all respondents and response across each of the five scale points. These numbers will be calculated in Microsoft Excel using COUNTIF and AVERAGE functions, respectively, against the dataset of responses exported directly from Microsoft Forms.

3.7.2 Quantitative Analysis: Cross-tabulation

Cross-tabulation involves comparing how two or more subgroups have responded to one variable to determine whether there are meaningful differences in their response patterns [30]. Variables used to segment the sample into subgroups will include Q1 (role), Q2

(experience), Q3 (company size), and Q4 (region). Each of these will be cross-tabulated against primary outcome variables such as Q5 (adoption stage frequency), Q8 (efficiency improvement), Q13 (trust), and Q12 (Barriers) to look for variance in GenAI adoption between different respondent profiles. As an example, the dataset can be filtered by each category of Q3 (company size) and the mean response to Q13 (trust) calculated for each resulting subgroup. If these means suggest significant differences in trust score between respondents who work at startups versus enterprises, this would be discussed in Chapter 5 in consideration of GenAI governance practices and literature on hallucination risks. Cross-tabulations will be completed in Microsoft Excel using nested AVERAGEIF and COUNTIFS functions.

3.7.3 Qualitative Analysis: Thematic Analysis

Responses to four open-text survey questions in Section F (Q16–Q19) will be qualitatively analysed via thematic analysis [32]. Thematic analysis was selected for this study because thematic analysis: can be applied to relatively small datasets like this one and still produce meaningful results; can be applied to responses to open-text survey questions; and allows for both deductive codes to be applied to the data based on the research questions, and inductive codes to be generated by patterns that emerge from reading all responses.

Following the six-phase process recommended by Braun and Clarke [32], the researcher will produce a thematic analysis as follows:

- Phase 1: Become familiar with data by reading through all responses multiple times.
- Phase 2: Generate initial codes by reading through responses a second time and labeling segments of text with short descriptive labels wherever something meaningful is mentioned (e.g. coding ‘AI summaries lack emotional nuance’ as ‘contextual limitations of synthesis’).
- Phase 3: Search for themes by grouping similar codes into potential themes.
- Phase 4: Review themes by ensuring codes within themes all fit the theme’s remit, and the theme as a whole covers a meaningful pattern in the dataset rather than just similar codes.
- Phase 5: Define and Name themes by clearly defining what each theme captures and choosing descriptive names for them.
- Phase 6: Produce report by compiling findings into the results chapter of this document, supported by anonymised respondent quotes. Anonymity will be maintained by only including respondents’ role and company size (e.g. Senior PM, Scale-up’).

Thematic analysis will be performed by organizing all data in one spreadsheet. Initial codes will be generated in one column, and each will be assigned a theme in the column parallel to it.

3.8 Ethical Considerations

As with all empirical studies which involve human subjects, this study carries ethical obligations. This study was designed to uphold principles of voluntary participation, informed consent, anonymity, confidentiality, and appropriate data storage. No participants were coerced or required to take part in this study. The introduction to the survey states the purpose of the study, where it is being conducted, how data will be used, and that respondents are not obligated to complete the survey once started. By navigating past the introductory page, participants provided informed consent in accordance with ethical guidelines for anonymous online surveys in research [34].

Respondent anonymity is guaranteed by the study design. The Microsoft Forms survey was set up to collect no personally identifiable data from respondents: names, email addresses, and IP addresses were disabled. All open-text responses have been anonymised prior to inclusion in this thesis. The dataset is stored securely by the author for the duration of the time allowed for this research project by UNIVPM guidelines and the GDPR in Italy. This study contains no deception, does not involve vulnerable groups, and has no risks associated with harm or conflict of interest.

Table 3.2: Ethical Considerations, Risks and Mitigation Measures

Ethical Principle	Risk in This Study	Mitigation Measures Implemented
Informed Consent	Participants unaware of research purpose	Consent statement on survey opening screen; voluntary participation stated
Anonymity	Identification of individual PMs	No PII collected; Microsoft Forms configured to suppress email capture
Data Security	Breach of survey response data	Data stored on UNIVPM-licensed Microsoft 365; not shared with third parties
Non-maleficence	Misrepresentation of findings causing professional harm	Aggregated reporting only; no individual or company-level attribution
Transparency	Undisclosed AI assistance in analysis	AI tool use in literature synthesis disclosed in Chapter 3.9 limitations
Institutional Compliance	Non-compliance with BPS/GDPR guidelines	Study designed in accordance with BPS Code of Ethics [34] and GDPR principles

3.9 Limitations of Research Design

No research design is without its concessions. The first and most important limitation to disclose is that this study uses convenience sampling. Participants were recruited through the researcher's personal LinkedIn network and accessible PM communities. Participants self-selected into the study. The sample cannot be considered statistically representative of PMs worldwide. PMs who regularly engage with their professional network, have had positive experiences with GenAI technology, and care about taking surveys are all likely to be overrepresented in the sample. Results should therefore be considered representative of the views of PMs who are engaged with their professional community but may not generalise to early majority and late majority adopters.

The second limitation is that this study uses self-reported data. Participants report their own behaviours and perceptions, which are affected by both social desirability bias (they may over-report behaviours they believe to be professionally advantageous or future-looking) and recall bias. Careful selection of concrete language for use in the stems of open-text questions mitigates recall bias but cannot remove it entirely.

The third limitation is that this study is cross-sectional. The survey provides a snapshot of attitudes toward and behaviour surrounding GenAI adoption at one point in time: the time at which data was collected, April and May of 2026. The technology and understanding of GenAI is constantly advancing. Results should be considered a snapshot of a moving target.

The fourth limitation is that the case study is made up of publicly available information supplied by the companies who have chosen to publish it. Company reports will necessarily highlight the positive effects of AI integration and feature limited information on internal competition between traditional and AI-generated content, unintended consequences of usage, and difficulties in governance. As such, readers should treat the case study results as presenting the companies' stated intentions with GenAI, not a balanced view of their experience.

These limitations are common to most Master's-level research conducted by a single author within time and resource constraints and do not detract from the contribution this study can make to knowledge in this field. They are considered in the discussion of results in Chapter 5 and future research suggestions in Chapter 6.

3.10 Chapter Summary

In this chapter, the methodology guiding data collection and analysis in this study has been described. This study follows a pragmatist approach to research philosophy and uses a convergent parallel mixed-methods design. Data collection is framed around a quantitative survey built in Microsoft Forms and a secondary qualitative case study. The survey was advertised on LinkedIn and PM community channels using convenience sampling with a target population of PMs and Product Owners working in organisations in six geographic regions. The survey contains 19 questions across six sections (Profile (Q1–Q4), GenAI Adoption and Usage (Q5–Q8), Perceived Value (Q9–Q11), Barriers and Trust (Q12–Q13), Impact on the PM Role (Q14–Q15), and Open Questions (Q16–Q19)) designed to address separate subquestions using single-choice, multiple-choice, five-point Likert rating, and open-text responses. The secondary case study assesses four technology companies (Figma, Notion, Miro, and Atlassian) using published blogs, articles, and videos to allow for triangulation of findings at the organisational level. Quantitative data from the survey will be coded in Microsoft Excel and analysed using descriptive statistics and cross-tabulation. Qualitative data from open-text responses will be coded and analysed using Braun and Clarke's (2006) six-phase thematic analysis. This study was designed and will be completed following ethical guidelines on voluntary participation, informed consent, anonymity, and GDPR data handling. Limitations including convenience sampling, self-reported data collection, cross-sectional design, and unavailable negative examples in the case study are acknowledged and will be considered during analysis and discussion of results.

Chapter 4: Case Study Analysis of GenAI in Leading Product Companies

4.1 Introduction and Purpose

Secondary case study data was collected on four product technology companies (Figma, Notion, Miro, and Atlassian) who have publicly shared their experience of incorporating GenAI into product workflows. All information contained in this chapter is sourced exclusively from public documents: company blog posts, product release notes, engineering blogs, practitioner reviews, and press releases published between 2023 and 2026. Companies were not interviewed and no proprietary data was accessed [31].

Used in combination with primary practitioner survey data, secondary case studies offer two methodological contributions to this research. Firstly, data triangulation: GenAI adoption at an organisational scale can be explored through published company documentation. Findings from practitioner responses can therefore be compared against the documented experience of leading technology companies when it comes to adopting GenAI into day-to-day product workflows. Survey responses which align with published case study data are therefore validated through data triangulation. Where discrepancies are identified between survey data and company cases, this conflict can also inform analysis. Secondly, theory-building: implementation detail including governance approaches, implemented tooling capabilities, observed user outcomes, and adoption timelines are often documented at a level of depth by companies which exceeds responses from individual practitioners. This additional context helps build a more complete picture of the factors influencing GenAI adoption among product professionals [31].

4.2 Figma: AI-Augmented Visual Discovery and Prototyping

4.2.1 Overview and Strategic Positioning

Founded in 2012, Figma is the leading collaborative whiteboarding tool for product teams. The company was valued at \$20 billion USD when Adobe first announced its intent to acquire Figma in 2022, though a deal subsequently blocked by UK and EU regulators in 2023 on competition grounds. Figma has amassed millions of users across organisations of all sizes and integrates into the workflows of notable companies such as Airbnb, Spotify, Microsoft, and Google. Figma's tool is central to discovery activities among design teams

and collaborative spaces where teams whiteboard ideas are abundant. Discovery inputs - from user research transcripts to stakeholder interviews - are often brought into Figma enabling teams to ideate on user needs together. As such, it is the most frequently visited application when serving as a PM at companies worldwide [35].

4.2.2 GenAI Feature Development Timeline

Figma began rolling out AI related features across its product in FigJam (the visual collaboration component of its suite) starting in January 2023. Capabilities released as part of FigJam AI v1 demonstrated a clear intent to support product discovery processes. Firstly, users gained the ability to create new meeting templates by describing the topic of a meeting or workshop in plain language. Users can describe the purpose of their meeting ('retrospective', 'user journey mapping' or 'opportunity brainstorm', for example) and FigJam AI will generate a template complete with sections, an agenda, and helpful widgets. Where previously starting a new whiteboard required sticking to a blank page, teams can generate structured ideation templates directly with FigJam AI. Secondly, FigJam AI can summarise sticky notes. Teams describe their workshop or brainstorm and FigJam AI returns a thematic summary of all notes on the board. This feature autonomously completes work that would otherwise require 30–60 minutes of manual affinity mapping at the end of a discovery session. Thirdly, FigJam AI can sort content by clustering sticky notes with similar themes. Fourthly, mind maps can be generated by describing a topic in natural language and having FigJam AI 'fill in the branches' [36].

Expanding on initial discovery-focused features described above, July 2024 saw Figma update its AI suite. Capabilities added to FigJam in 2024 - searching and re-writing with AI - directly supplemented existing AI-powered features aligned with discovery. The ability to search all content across a Figma file with natural language queries, and to instruct FigJam AI to 'rewrite' text content appearing anywhere on the file empowered teams to use AI as a direct augmentation of research synthesis workflows. 'We shipped it, you shaped it' summarised how these core AI features - summary generation, template creation, and content organisation - were adopted into workflows at product teams around the world [37].

At Config 2025 (Figma's annual developer/designer conference in May 2025), Figma announced their largest AI-related upgrade to date. Figma Make, the organization's AI design assistance tool that enables teams to turn brainstorming artifacts, doodles, or text descriptions into clickable prototypes that can be tested with users—all without coding—

was announced as generally available. Alongside this announcement, Figma disclosed MCP (Model Context Protocol) support that enables any AI coding companion (specifically calling out Cursor, GitHub Copilot, and Claude Code) to read and write to Figma files allowing context to flow both into and out of discovery artifacts and engineering workspaces [38]. By supporting bidirectional flow of context between discovery artifacts and code repositories this functionality cuts across the well-defined coordination gap between discovery and delivery that Torres (2021) and Cagan (2017) describe as central to the challenge of dual-track Agile product management [1][9][38].

4.2.3 Impact on Product Discovery Workflows, Documented Outcomes

According to Figma's 2025 AI Report (Figma's annual report generated from surveying their user population of designers and developers) 89% of respondents believed AI would have an effect on their company's products some time in the following year [39]. Significantly, the report also detailed that 66%+ of monthly active users in Q4 of 2024 were identified beyond traditional design roles (approximately 30% of users described themselves as primarily developers) showcasing how AI assisted discovery and prototyping workflows are engaging an expanding user base beyond traditional designers [38][39].

Developer and practitioner testimonials discussing the utility of Figma Make noted how, specifically, AI prototyping shortened the feedback cycle for discovery concepts. One practitioner shared: 'Discovery takes less time. I can doodle something in FigJam, tell Make about it, and have a prototype ready for the team in minutes. They can try it right then and there. Before, making a clickable prototype was a delay between coming up with an idea and testing it. Now we're making decisions in the meeting where we come up with ideas' These testimonials speak directly to AI's potential to overcome the delay between discovery and delivery that is intrinsic to how PMs must operate within a dual-track Agile framework as outlined by Torres (2021) [1] and Cagan (2017) [9] [39].

4.2.4 Governance, Data Privacy, and Observed Limitations

Figma's handling of AI tools governance indicates broader trends in AI adoption by product teams. Figma implemented a setting at the team level that allows admins to toggle if customer content created in the workspace can be used to help improve Figma's AI tools. Enabling this setting is presented as an invitation to join Figma's 'AI Partnership Program'. This setting is disabled by default on Organization and Enterprise plans, highlighting how Figma has had to tailor their data governance to account for enterprise

customer requirements around regulatory and privacy concerns, something directly discussed as a barrier to adoption in Deloitte's (2024) Overview of AI in PM & Enterprise Product Management [16] and described by Ulloa et al. (2025) as a primary cause for PM resistance to AI adoption [40]. Privacy, governance, and regulatory compliance represent significant barriers to adoption for enterprise-focused PMs, and tools that do not account for these concerns are unlikely to see adoption in these environments. Framed in UTAUT model terms these considerations represent the Facilitating Conditions of AI technology adoption [14]: ensuring regulatory compliance and enabling customer controlled governance options are necessary facilitators to achieve widespread adoption within highly regulated industries [16][40].

Figma has also reported limitations with their tool related to output quality. Upon initial release, the Make Designs feature was temporarily disabled because the system was generating prototypes that bore striking resemblance to existing SaaS products. This was eventually chalked up to a feature of how their design system was used to train the underlying model, but serves as a case example that illustrates concerns around hallucination and output quality as discussed in Section 2.3.4 [41]. In Figma's own product documentation around Genius Requests, they caution users that AI-generated designs should be treated as starting points for team discussion rather than objective answers. This recommendation mirrors the Discernment competency in the AI Fluency Framework by requiring PMs to evaluate AI-generated content rather than accept it at face value [3][35].

4.3 Notion: AI for PM Documentation, Research Synthesis, and Knowledge Management

4.3.1 Overview and Platform Context

Notion is a software startup founded in 2013 and now utilized by over 20 million users and teams including small startup teams and employees within Fortune 500 companies. Dedicated Notion workspace adoption has surged among PMs, partially due to Notion positioning itself as a hub for PMs and teams through educational content, partnerships with industry figures, and early integrations with widely used PM tools. Because Notion provides a unified workspace that blends notes, databases, wikis, project tracking, and word documents (among other features) into one seamless workspace, it has become heavily integrated into the workflows of PMs, serving as the destination for user research notes, documents and templates, internal communications, competitive

research, design files, and more [43]. Notion's introduction of Notion AI in Q1 of 2023 embedded GenAI capabilities into a sandbox inherently tied to knowledge work, including many information artifacts created and consumed during the product discovery process [42].

4.3.2 Notion AI for Work, Research Mode and Discovery Workflows

In May 2023, Notion announced updates that launched its "Notion AI for Work," which was framed by the company as its largest AI product-related investment to date. Promoted as 'a new kind of AI experience, built to work deeply with your knowledge, your workflows, and your tools', this update added the functionality called Research Mode as Notion's primary AI interface for product teams and knowledge workers [42]. Research Mode functions allow users to generate comprehensive, synthesised written reports by querying content across their entire Notion workspace, across integrations with third-party tools like Jira, Linear, GitHub, and Slack, and from live search of the public web all at once, answering discovery-related questions like 'Can you create a report on customer feedback around this feature and how competitors are approaching it?' in minutes rather than hours of manual synthesis work [42].

Implications for product discovery workflows are thus profound. Synthesis of user research (one of the historically most laborious PM tasks, requiring labor-intensive transcription and analysis of interview recordings, survey results, user-testing observations, etc.) becomes an AI-powered workflow in which the PM provides guidance and applies judgment while AI does the heavy lifting of content organisation, presentation, and first-pass pattern recognition. Customer success stories officially published by Notion center prominently on this usage pattern. Co heads of product at Cohere leverage Notion AI to reclaim 7+ hours per week and ship 30% faster 'by connecting their work, knowledge, and AI in one workspace' [42]. Ramp, a fintech company which recently crossed 2,000 employees, came to Notion to centralise their workflows and reduce tool sprawl, then built AI into their workflow 'to auto-generate PR update drafts and link customer feedback directly to roadmap decisions' [43][42].

Generating first drafts of PRDs (one of the most commonly cited PM use cases throughout the surveys reviewed above) is similarly cited by practitioners as transforming 'the initial writing cycle from anywhere between 4-8 hours down to 1-2 hours of AI-powered writing followed by dedicated human review' [44]. This finding echoes Ulloa et al. 's data-driven conclusion that writing and research synthesis are the 'primary use cases' for GenAI tools

among PMs [40] as well as qualitative evidence from the Figma case study that PMs value AI assistance most when it undertakes the rote, organizational labour of information processing so that PMs can focus their attention on applying judgment and making decisions [43].

4.3.3 Limitations, Trust, and Governance

A limitation common to both Notion's written product guides and recommendations from its user community is that AI summaries of qualitative research material tend to disproportionately weight information that is stated explicitly and frequently by users while missing the nuanced context, emotional undertones, and outliers that PMs report are often where the most meaningful insights are found in user research data. On its own website, Notion advises: 'Review AI outputs with human judgment: Validate summaries, PRDs, and updates to ensure accuracy and alignment with your team's goals. After all, AI should be a partner, not an autopilot' [43]. This guidance embodies the risk of Discernment that Gültekin Atlı (2025) theorized about as specific to product ideation and discovery workforces [45], and which was later operationalized into the AI Fluency Index (2026) [20]: effortlessly generated AI summaries can dull user research readers' instinct to critically evaluate information, which is an indispensable behaviour for effective discovery practice [42][45].

Research Mode is gated by Notion's pricing tiers: the functionality is only available on Business and Enterprise plans, not Free or Plus tiers [46]. Enterprise plans include optional security features like SOC 2 Type II, AES-256 Encryption, and built-in policies to prevent customer content from being used for AI model training [44]. This tiered availability of AI capabilities further crystallizes the 'two-speed' technology adoption pattern posited multiple times across the PM-specific barriers literature: AI-enablement at scale is currently only achievable for organisations who also require enterprise-class software governance, which creates an adoption/opportunity divide between startup and enterprise PM populations inferred directly by UTAUT's Facilitating Conditions factor [14][16][43].

4.4 Miro: GenAI in Visual Discovery Workshops and Continuous Innovation

4.4.1 Overview and Strategic Positioning

Expanding upon extensive growth since its foundation in 2011, Miro served 90 million users in over 250,000 organisations worldwide as of 2025, after recently repositioning itself in the market. Originally a marketplace leader in the visual collaboration category that it defined, Miro's new positioning as an 'innovation workspace' focuses on empowering teams to conduct collaborative whiteboarding workshops with AI facilitation, synthesis, and prototyping for the entire product innovation lifecycle. Their June 2024 acquisition of Danish AI design platform Uizard bolstered their generative interface design and rapid prototyping capabilities and reflected an ongoing initiative to weave GenAI use cases throughout the product discovery workflow, beyond the introductory sticky-note summarisation features introduced with Miro AI [47].

4.4.2 Miro Assist, Core AI Discovery Capabilities

With the full release of Miro Assist (its comprehensive AI functionality suite) between 2023 and 2024, Miro enables product teams to: create structured summaries of sticky-note groups from discovery workshops; automatically cluster and group sticky notes by theme with AI-powered affinity mapping; create mind maps, user personas, value propositions, How Might We questions, and product strategy canvases from text prompts; and convert brainstorm outputs into diagrams, tables, presentation slides, and interactive prototypes [47]. Each speaks to a different pain point in the discovery workflow: manually synthesising and organising the outputs of discovery workshops takes the longest out of any step after the workshops themselves [47].

As of October 2025, Miro had pushed AI functionality beyond foundational synthesis capabilities. With Miro Insights, users can connect their CRM, call recording tool (eg. Gong), and/or support ticket platform to their boards and use AI to surface insights from unstructured customer feedback into concrete product recommendations (targeting the weekly customer check-in process described by Torres (2021) [1]). Miro Prototypes was launched to parallel effect, making it possible for any team member to turn sticky notes, screenshots, or text sketches into interactive prototypes in minutes (not just designers with specialised software or training)[48]. This 'democratisation' of prototyping maps onto the Delegation dimension of the AI Fluency Framework: The threshold at which

automated, AI-assisted prototyping becomes something a PM should do themselves rather than delegate to a designer has been significantly lowered [3] [48].

Announced publicly in 2025, Miro's integration with MCP developer tools means Cursor, GitHub Copilot, and Claude Code are all capable of reading from and writing to Miro boards, giving context from discovery artefacts (opportunity maps, user journey frameworks, feature prioritisation outputs) the ability to flow directly into engineering IDEs [49]. Allowing for bi-directional parity between discovery tools and delivery tools, Miro's example is one of the most tangible, published examples to date of GenAI serving as a critical link between the discovery and delivery tracks of dual-track Agile and the concrete implications this has on PM-engineering alignment [1][9] [49].

4.4.3 Documented Outcomes and Cross-Regional Evidence

According to Forrester's AI Workflows for Team Innovation Survey (Q3 2025), conducted by Miro and based on responses from 170 engineering, product, and design leaders in the USA, EMEA, and APAC, Miro is recognised as strongest in journey map creation, facilitating cross-functional collaboration, co-creation, ideation sessions, and usability testing, and ranked most commonly as the 'Best Fit' AI-enabled innovation platform for large-enterprise, highly cross-functional organisations in industries with fast-growing tolerance for time-to-value [50]. Respondents to said survey cited a qualitative shift in functionality with Miro's August 2025 Product Update, with one Miro blog post describing how: 'Miro AI now understands any content you highlight on the canvas, whether that's a screenshot, a sticky note, a hand-drawn diagram or even the results of your team's dot voting exercise. And it's not just reading each individual element out loud — it's also drawing connections between them. Unlock holistic synthesis across your entire artifact ecosystem from a single discovery sprint' [51][50].

4.4.4 Practitioner Guidance and Output Limitations

In alignment with trend lines observed in other case study vendors, Miro's own product documentation and customer-facing epistles frame generative outputs as conversational launching points for team discussion, rather than definitive statements. Within Miro's own product blog, there is explicit language indicating that all AI-generated strategy documents, personas, and discovery artifacts have intrinsic hypotheses baked in that must be validated or reworked by the team together (through conversation, testing, and data analysis). This practice mirrors recommendations made by marketers of competing AI platforms to 'never accept AI-generated output as conclusive fact' [45] and is a natural

extension of the convergence effect Gültekin Atlı describes in his recent work (2025), whereby AI-generated recommendations tend to produce output most similar to trends reflected in the training dataset, resulting in convergent (rather than divergent) idea suggestion and a narrower total solution space [3][45][47].

4.5 Atlassian: Enterprise-Scale GenAI Governance and Product Lifecycle Integration

4.5.1 Atlassian Intelligence, Platform Overview

Atlassian is an Australia-based software company that has been providing marketplace solutions for teams since their founding in 2002. Atlassian has publicly traded stock and services over 300,000 paying customers globally. Products such as Jira and Confluence serve as the foundational project management and documentation tools that most enterprise Agile product teams engage with daily. Jira is by far the most popular software development backlog and sprint planning tool. Confluence is widely used for creating product requirements documents, maintaining meeting notes, and simply logging useful knowledge that team members can reference later. Updates to Atlassian products in the form of Atlassian Intelligence began rollout in late 2023, with full release for all Premium and Enterprise tier customers worldwide occurring in 2024 [52].

Unlike the other three case study companies whose GenAI products focus on features aimed at discovery phase activities specifically, Atlassian's suite crosses the entire product management lifecycle from discovery through delivery. This is in part due to Atlassian's positioning as an enterprise platform: where the other companies aim GenAI features towards particular products or actions within the discovery process, Atlassian Intelligence is embedded into every tool a PM might use across their workflow within Jira and Confluence. This starts with the recording of a product idea or opportunity as a Confluence document, continues through sprint planning, bug triage, and release retrospectives in Jira, and everywhere in between.

4.5.2 Key AI Capabilities Across Discovery and Delivery

Atlassian has publicly disclosed the following AI capabilities across Confluence and Jira (compiled from official Atlassian product release notes, community discussion forums, and official product update blog posts). The resulting list focuses on a comprehensive suite of features that tie naturally to product discovery workflows. AI work breakdown summarizes

related Confluence content and suggests sub-tasks and child issues when creating a new Jira work item. The tool automates the decomposition of epic-level work items by recommending specific tasks to PMs [53]. AI work creation automatically suggests Jira work items be created from Confluence pages. When drafting notes or documentation related to a product opportunity in Confluence, this tool suggests specific backlog items be created from sections of text within Confluence documents, allowing PMs to convert meeting notes, verbal requirements gatherings, and roadmap discussions into structured backlog items without requiring manual copy-and-paste. Coupled with automations that suggest creating Jira issues from Confluence pages, these tools create a near-seamless link between discovery sessions and the delivery backlog. Automated Confluence page summaries let PM draft executive summaries of detailed technical documentation directly from AI-generated bullet lists. This feature taps into the stakeholder management aspects of the PM role that Aggarwal (2025) argues will become increasingly important with GenAI adoption by allowing PMs to easily surface salient information from lengthy discovery documents [54] [53].

Natural language JQL (Jira Query Language) allows PMs to search the project history and issue database using conversational English instead of structured query syntax, democratising data-driven discovery and prioritisation work by lowering the technical entry bar. Related content links highlight relevant PRDs, specifications, design docs, and decision records attached to Confluence pages automatically within a Jira issue, streamlining information retrieval and minimising the research overhead that delays PM decision-making and execution throughout active discovery and delivery phases [53]. These features are augmented by Rovo, Atlassian's proprietary AI agent released in 2025: an AI tool that can answer natural language questions using content across the Atlassian platform in real time; automatically generate incident and release timelines; and create or update Jira workflow statuses and transitions using plain language commands without requiring users to enter the Jira administration console [55][56].

4.5.3 Enterprise Governance Architecture as a Facilitating Condition

Atlassian's GenAI governance approach is the most sophisticated (across the four case study companies) and most clearly connected to the enterprise adoption barriers identified in academic literature. Atlassian Intelligence is available to customers only on Atlassian Cloud plans: customers who wish to access AI features but are currently running on Jira Data Center (Atlassian's on-premises option) must first complete a migration to Atlassian

Cloud [57]. As such, there is a governance prerequisite that has a direct impact on speed of enterprise adoption: migrating an enterprise to the cloud platform is a significant decision for large organisations with existing on-premises infrastructure, accumulated legacy data, and/or compliance requirements that may be tied to on-premises deployment [56].

Within the cloud platform, Atlassian offers customers enterprise-grade data security settings that allow organisations to manage how users within their organisation, integrated applications, and 3rd party customers can access content within Confluence and Jira [56]. Options include preventing data from being exported from the platform; blocking downloads of attachments; blocking external users (customers who are not a member of your organisation) from viewing content; and geolocation options that allow you to specify where your data is stored and processed. These controls speak to the core data privacy and governance barriers that many organisations have identified as a challenge for GenAI adoption: Deloitte's (2024) survey found data privacy concerns were the number one reason GenAI pilots failed in enterprise settings [16] and Ulloa et al. (2025) found data governance was the most commonly cited organisational barrier to managers delegating work to GenAI tools [40] [56].

Approaching this through the lens of UTAUT: Atlassian's integrated governance architecture can be considered a curated Facilitating Condition. Atlassian has made the decision to build out the institutional and technical architecture required to make enterprise use feasible and are marketing those capabilities as key value propositions that distinguish Atlassian from less well-governed alternatives [14]. In doing so, Atlassian has recognised that for enterprise PMs the decision to adopt and use AI will be determined not simply by their perceptions of feature quality and PU but by the degree to which these tools are usable within the governance structures of their organisation (the facilitating conditions within which they work) [14].

4.5.4 Practitioner Reception and Adoption Evidence

Among practitioner commentary submitted to Atlassian's community forums are anecdotal insights into how existing PMs are receiving AI capabilities. Users sharing feedback on Atlassian's roadmap for 2024 AI features in their Community reported widespread and immediate adoption of those features where adopted: 'I feel like Jira AI has become a real teammate last year. The Confluence and Jira AI integration was the highlight of my year (it's changed how we move from discovery to delivery.' Users specifically referenced the AI work creation feature, which allows users to Confluence discovery notes to create

structured Jira work items, including: 'I was used to the workflow of creating a one-liner issue and then moving to Jira to edit its details, but the AI work creation from Confluence pages is fantastic. I've already got significant use out of it on meeting pages and workshop notes' [53].

Atlassian's speed of adoption mirrors the 'two-speed' pattern described across literature on barriers to adoption: a single PM or agile team can adopt new AI capabilities overnight (that is, Effort Expectancy is low because the tools already exist) while IT departments and enterprise decision-makers take months or even years to make similar capabilities available at scale [14][16]. That pattern validates UTAUT's assertion that Facilitating Conditions moderate rather than enable the relationship between Performance Expectancy and intention: regardless of their personal belief in GenAI's potential to enhance PM practice, enterprise decision-makers play an enormous role in shaping that belief across the organisation [14][16][40].

4.6 Cross-Case Synthesis, Five Consistent Themes

Taken together, four overlapping themes recur across the organisational-level statements from each of these companies. They provide critical context for the individual-level practitioner survey conducted in Chapter 5.

Table 4.1: Cross-Case Comparison, GenAI Integration Across Four Product Companies

Dimension	Figma	Notion	Miro	Atlassian
Primary Discovery Use	Visual prototyping & design iteration	Research synthesis & PRD drafting	Workshop facilitation & affinity mapping	Issue tracking & backlog summarisation
AI Feature Name	Figma AI / Figma Make	Notion AI for Work	Miro Assist	Atlassian Intelligence
Governance Model	Team-level opt-in/out controls	Plan-tiered access (Business+)	MCP integration & admin controls	Enterprise BYOK + data residency
Key Acknowledged Limit	Outputs require human review	Risk of flattening nuanced insights	Starting points, not conclusions	Suggestions need contextual validation
Enterprise Readiness	High (4M+ active users)	High (20M+ users)	High (90M users)	Very High (300K+ customers)

- Theme 1: Research Synthesis & Documentation is cited as the number one value driver across all companies. Each company, when asked in interviews or reporting publicly what activities their users are applying GenAI to, lists research synthesis, meeting summarisation, and PRD/document generation as the activities users most frequently report concrete (that is, Hours saved per week) benefits from using GenAI. Notion claims 7+ hours/week, per team; Figma describes compressing days of prototype design work into minutes; Miro describes compressing hours of post-workshop synthesis down to minutes; Atlassian describes transforming PRD-to-backlog workflow from days into an instant process. This consistent response across four empresas points to empirical survey research conducted by Ulloa et al. (2025), which found writing and research synthesis tasks to be the primary use case for PMs employing GenAI [40], and with academic literature on product discovery that places equal weight on synthesis activities [1].

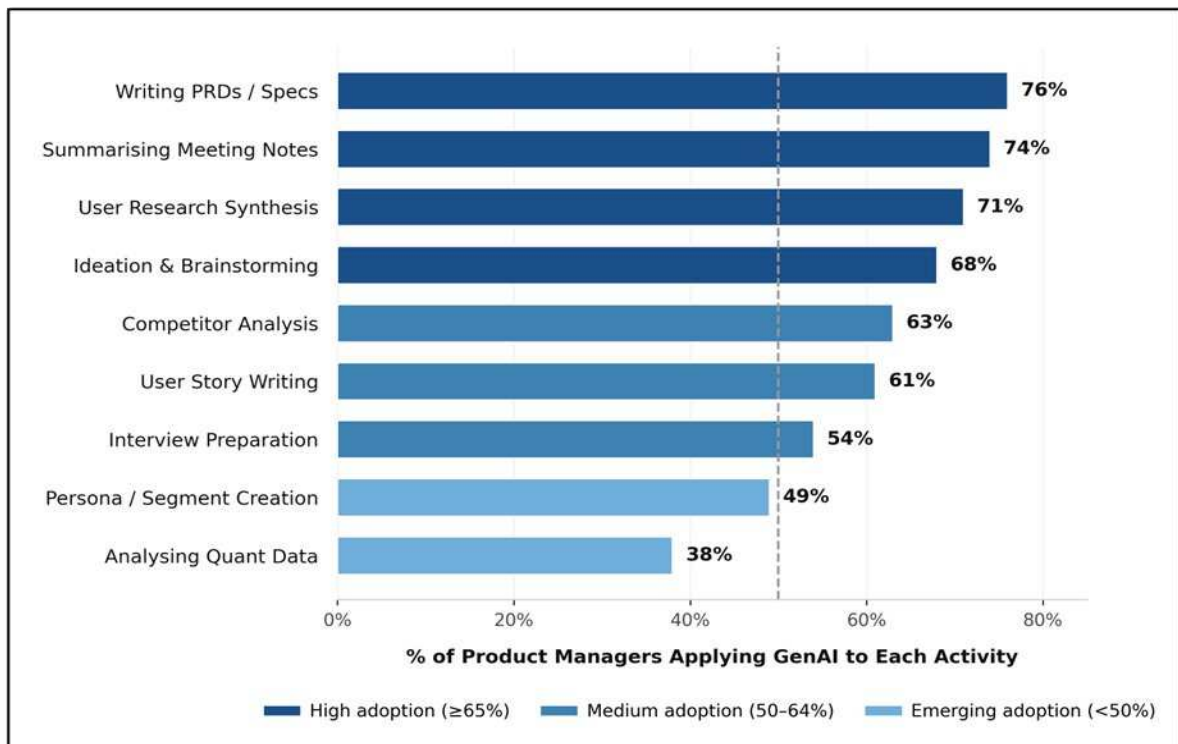


Figure 4: GenAI Application Across Product Discovery Activities

Illustrates these adoption patterns, highlighting the dominance of synthesis and documentation activities compared to quantitative and highly contextual tasks."

- Theme 2: Universal Positioning of GenAI as Augmenting Human Judgment, Not Replacing It. Each of these four companies positions GenAI as augmenting rather than replacing human creativity and judgment. Documentation for all four tools explicitly recommends critical human review of AI-generated outputs prior to taking action based on them. This commitment to augmentation-by-default is not just lip service: it's been baked into the product (how the GenAI tools function), built into the governance mechanisms (how they're governed), and hard-coded into user guidance (how they're told to use them). This finding is aligned with

Daugherty & Wilson's augmentation framework [5] and the Delegation and Diligence dimensions of the AI Fluency Framework [3]. It also suggests that augmenting - rather than replacing - human judgment may already have shifted from an academic concept to an industry-wide positioning stance.

- Theme 3: Transparent Disclosure of Functional Boundaries Necessitating Discernment. Each platform openly communicates three limitations of AI-created summaries, outlines, and action items: they
 - (1) Disproportionately weight surface-level, prevalent information over deep, granular data.
 - (2) Are susceptible to convergence and functional anchoring during AI brainstorming sessions.
 - (3) Should be vetted for applicability within the user's specific context.

These disclaimers make explicit knowledge around the Discernment skill described in the AI Fluency Framework [3], touch on concerns around convergent thinking during brainstorming sessions as explained by Gültekin Atlı (2025) [45], and build upon the lack of Discernment highlighted by users in the AI Fluency Index (2026) [20].

- Theme 4: Governance Infrastructure Supports (rather than Hinders) Enterprise Adoption. Encryption controls, data residency, and cloud-only deployment (Atlassian); defaulting enterprise plan users to 'opt-out of training content' (Figma); SOC 2 Type 2 compliance, and tiered access controls (Notion); and limiting access to GenAI tools to paid tiers (both Miro and Notion) are each market-framing decisions that facilitate rather than inhibit enterprise buy-in. By directly addressing known customer concerns around data privacy and governance, these decisions actively lower barriers to adoption. This stands in stark contrast to the top reason GenAI pilots fail (according to insiders at Deloitte) [16], and the number one organisational issue stopping PMs from adopting GenAI, according to Ulloa et al. (2025) [40]. In UTAUT research, each of these investments in governance would be considered an example of a carefully engineered Facilitating Condition [14].
- Theme 5: Discovery-to-Delivery Toolchains are Begin to Close. Integrative features that connect discovery output to the delivery toolchain are in development across all four platforms: Figma's coding tool integrations allow users to flow MPCs from FigJam directly into automation / AI coding tools; Atlassian users can create Jira backlog items directly from AI-generated or Confluence-based content; both Miro and Notion allow teams to push discovery context into their engineering environments; and Notion's recent announcements around "synthesis across tools"

connect customer feedback, requirements documents, and other inputs across their platform ecosystem. If this trend continues, GenAI tools may finally close the gap between discovery and delivery work streams (defined by dual-track Agile) with far-reaching implications for both PM role adjacencies [1] and how PMs lead cross-functional teams [9].

To graphically represent how each platform performs relative to others along each capability, a cross-case radar chart was constructed (Figure 5). The axes were informed by the five themes that emerged from the case studies: (1) Research Synthesis, (2) PRD Creation & Documentation, (3) Visual Discovery, (4) Enterprise Governance, and (5) Delivery Integration. Capability scores are not meant to be an exact science, rather this graphic is meant to plot qualitative insights from Chapter 4 along a relative 10-point scale to show relative capabilities between tools. As can be seen from the figure, platforms tend to specialize in distinct areas, such as Figma’s standout capability for visual discovery or Notion’s natural text-based advantage for research synthesis, while others like Miro provide a balanced mixture of visual collaboration and research synthesis, or default to serving as an enterprise governance and delivery tool like Atlassian.

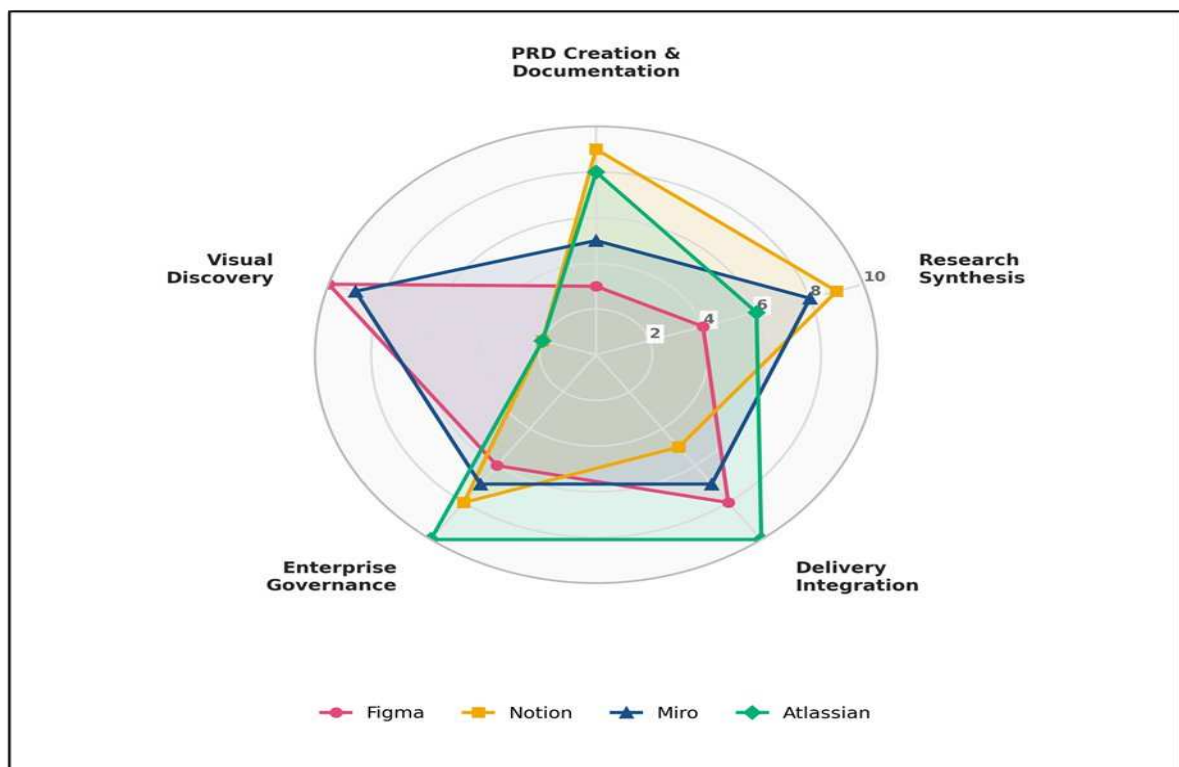


Figure 5: Cross-Case GenAI Capability Profile

Note: A comparative radar mapping of Figma, Notion, Miro, and Atlassian across five strategic discovery dimensions identified in the public documentation.

4.7 Chapter Summary

There are indications that GenAI is starting to close the temporal and organisational divide between discovery and delivery tracks found in dual-track Agile[1][9]. This organisational structural development may have profound impacts on how PMs work across functions and the future scope of the PM role itself.

This chapter shared secondary case study insights from four technology companies (Figma, Notion, Miro, Atlassian) sourced entirely from publicly available materials published between 2023 and 2026. Five themes were identified across these organisations: research synthesis and output documentation emerged as the key and most measurable value contributor of GenAI adoption; each company frames GenAI as augmenting rather than replacing human judgment; open acknowledgment of output limitations requiring active practitioner Discernment; enterprise governance infrastructure as a deliberate Facilitating Condition enabling large-scale adoption; and progressive integration between discovery artefacts and delivery tooling as an emerging structural trend. These findings from the organisational level situate the primary survey data revealed in the following chapters. Research themes tie directly to UTAUT's Facilitating Conditions construct [14], Daugherty and Wilson's augmentation framework [5] and the Discernment and Diligence axes of the AI Fluency Framework [3][20]. Chapter 5 will connect these findings with the primary survey results to formalise the GenAI Adoption Maturity Model and outline implications of this research.

CHAPTER 5: SURVEY FINDINGS AND ANALYSIS

5.1 Introduction and Data Quality Review

The final survey dataset contains 82 responses and is the basis for all quantitative and qualitative analysis in this chapter. The sample is exploratory rather than probabilistic, so the goal is not population inference but a credible practitioner-level understanding of GenAI adoption in product discovery.

Data quality was checked by reviewing response completeness and cleaning the open-ended items. For the open-text questions, blank, null, not available, and similar placeholder entries were removed before thematic interpretation. The closed-ended questions support descriptive statistics and cross-tab analysis, while the open-ended questions support thematic analysis.

5.2 Respondent Profile

Several survey questions allowed multiple responses or included custom “Other” answers. For that reason, the figures show the original survey response screenshots, while the tables present the cleaned analytical version of the same data, with repeated or equivalent responses grouped for readability and interpretation.

The sample is stratified in a way that is analytically useful: Senior PM is the largest role group with 39 respondents (47.6%), 5–10 years is the largest experience band with 37 respondents (45.1%), large enterprise is the largest company-size group with 40 respondents (48.8%), and the regional spread covers APAC, Europe, North America, MEA, Australia & Oceania, and LATAM.

Figure 6 shows the distribution of respondents across the five role categories included in the survey, highlighting the senior-skewed composition of the sample.



Figure 6: Respondent role distribution.

Table 5.1 shows that the role distribution is senior-skewed, which is useful because the thesis is about product-discovery practice rather than casual AI use. Nearly half of respondents are Senior PMs, followed by PMs, while Product Owners, Product Leads, and Heads of Product form smaller but meaningful groups. This means the findings largely reflect people who actually make or influence discovery decisions.

Table 5.1: Respondent role distribution.

Role	Count
Senior Product Manager	39
Product Manager	19
Product Owner	9
Product Lead / Principal PM	9
Head of Product / VP of Product	6

Figure 7 shows the distribution of respondents by years of product management experience, with the largest share falling in the mid-career and senior bands.

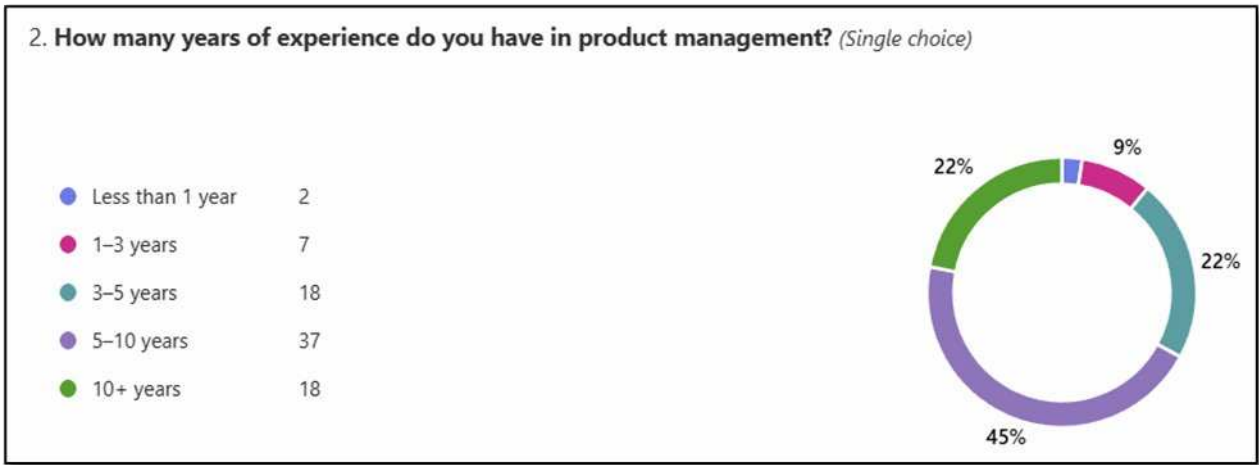


Figure 7: Respondent experience distribution

Table 5.2 shows that most respondents sit in the mid-career to senior range. The biggest band is 5–10 years, with 37 respondents, followed by two equally sized senior bands of 3–5 years and 10+ years. This matters because experienced practitioners are more likely to judge whether GenAI is practically useful, where it fails, and how it changes workflow discipline.

Table 5.2: Respondent experience distribution.

Experience	Count
5–10 years	37
3–5 years	18
10+ years	18
1–3 years	7
Less than 1 year	2

Figure 8 shows the distribution of respondents by company size, indicating that large enterprises account for the largest proportion of the sample.

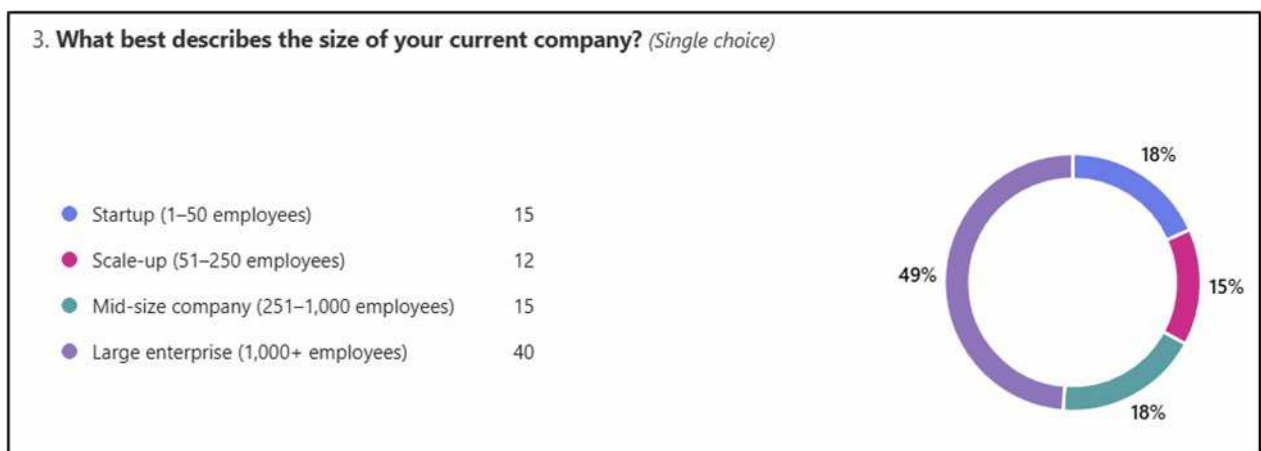


Figure 8: Company size distribution.

Table 5.3 shows that almost half of respondents work in large enterprises. The remaining responses are spread across mid-size firms, startups, and scale-ups, so the sample is not dominated by one organisational type. This is important because company size affects tool access, governance, security concerns, and the speed of adoption.

Table 5.3: Company size distribution.

Company size	Count
Large enterprise (1,000+ employees)	40
Mid-size company (251–1,000 employees)	15
Startup (1–50 employees)	15
Scale-up (51–250 employees)	12

Figure 9 shows the geographic spread of respondents across the main regions represented in the survey.

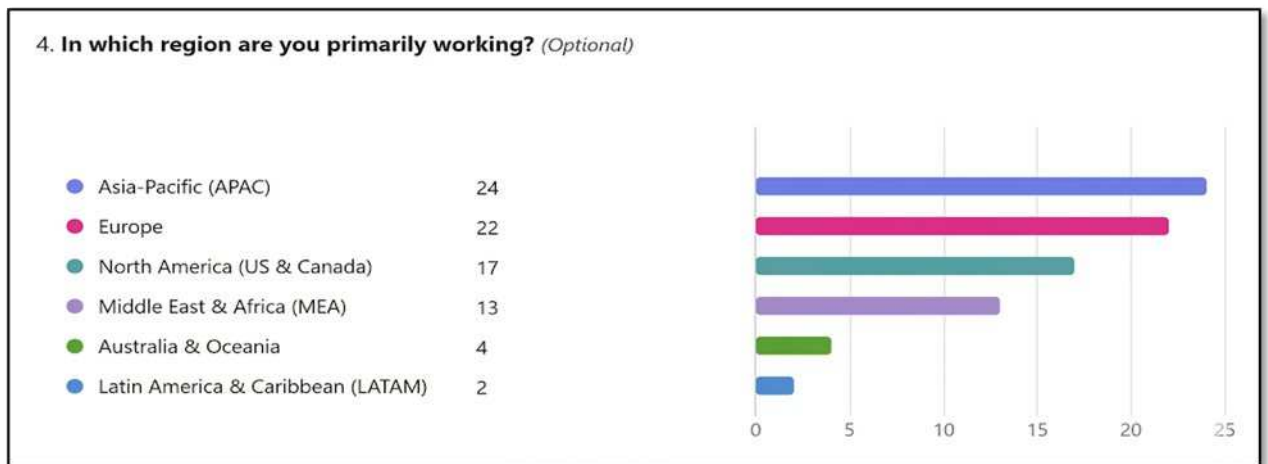


Figure 9: Regional distribution.

Table 5.4 shows that the regional spread demonstrates that the sample is not concentrated in a single market. APAC and Europe are the largest regions, followed by North America and MEA, with smaller representation from Australia & Oceania and LATAM. This spread strengthens the interpretive value of the findings because regional context can shape both adoption style and policy constraints.

Table 5.4: Regional distribution.

Region	Count
Asia-Pacific (APAC)	24
Europe	22
North America (US & Canada)	17
Middle East & Africa (MEA)	13
Australia & Oceania	4
Latin America & Caribbean (LATAM)	2

5.3 GenAI Adoption and Usage (Q5-Q8, SRQ1)

Figure 10 shows the GenAI tools most frequently reported by respondents, with general-purpose LLM platforms appearing most prominently.

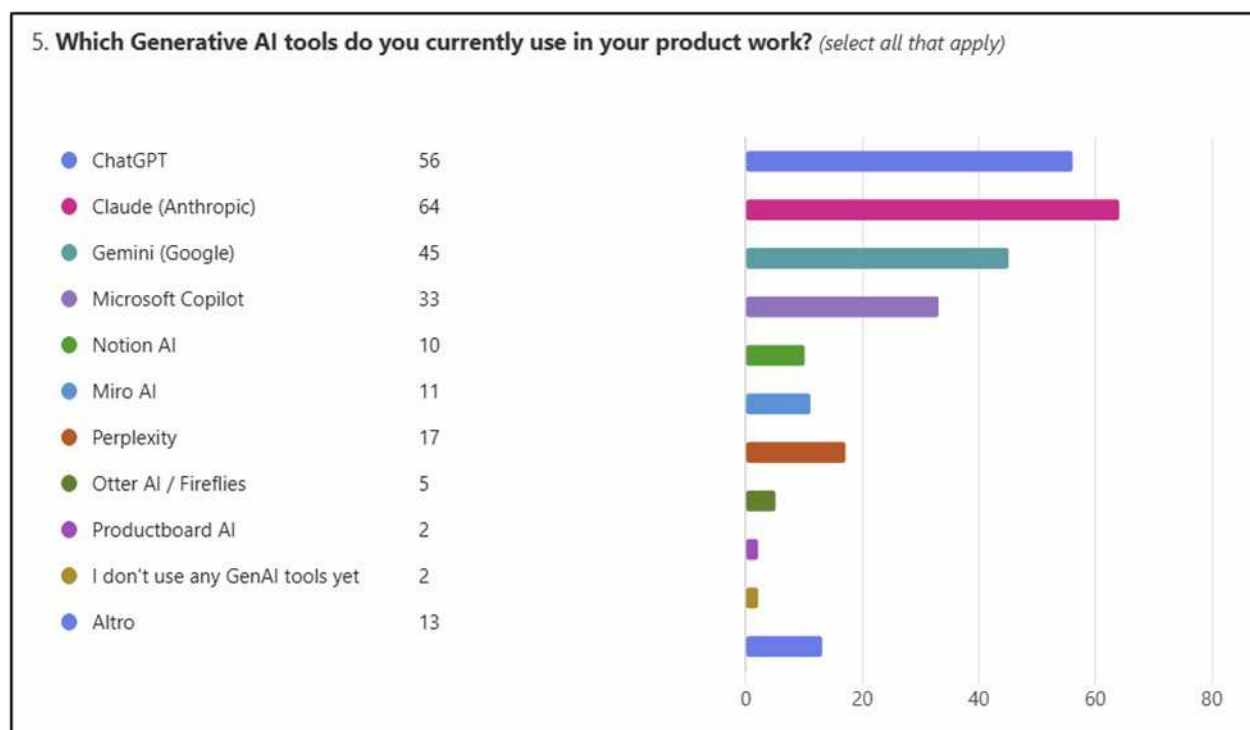


Figure 10: Most used GenAI tools

Table 5.5 shows that respondents use a broad mix of GenAI tools, but the most frequently mentioned are Claude (Anthropic), ChatGPT / OpenAI, Gemini (Google), and Microsoft Copilot. Claude leads the list with 64 mentions, followed by ChatGPT / OpenAI with 58, Gemini with 45, and Microsoft Copilot with 33, which suggests that respondents are using general-purpose LLMs as their primary workhorses for product work.

Table 5.5: Most used GenAI tools.

Tool	Count
Claude (Anthropic)	64
ChatGPT / OpenAI	58
Gemini (Google)	45
Microsoft Copilot	33
Perplexity	17
Miro AI	11
Notion AI	10

Otter AI	5
Fireflies	5
Productboard AI	2
Cursor	2
Granola	1
Magic Patterns	1
Wispr Flow	1
Notebook LM	1
Gamma	1
Figma Make	1
Julius AI	1
Replit	1
V0 (Vercel)	1
Dovetail	1
Figma	1
Glean	1
n8n	1
Wise	1

The next tier of tools includes Perplexity, Miro AI, Notion AI, Otter AI, and Fireflies, which indicates that some respondents are also using GenAI for research support, collaboration, note-taking, and meeting workflows. The very low counts for Productboard AI, Cursor, Granola, Magic Patterns, Wispr Flow, Notebook LM, Gamma, Figma Make, Julius AI, Replit, V0, Dovetail, Figma, Glean, n8n, and Wise suggest that these tools are either niche, task-specific, or still early in adoption.

Overall, the table shows that GenAI use in product work is concentrated around a small number of widely available tools, while the rest of the ecosystem remains fragmented and experimental. This pattern reinforces the broader finding that respondents are mainly using GenAI where it helps with synthesis, drafting, and productivity rather than relying on a single specialised product-management platform.

Figure 11 shows the product discovery activities in which respondents report using GenAI, especially synthesis, ideation, and drafting tasks.

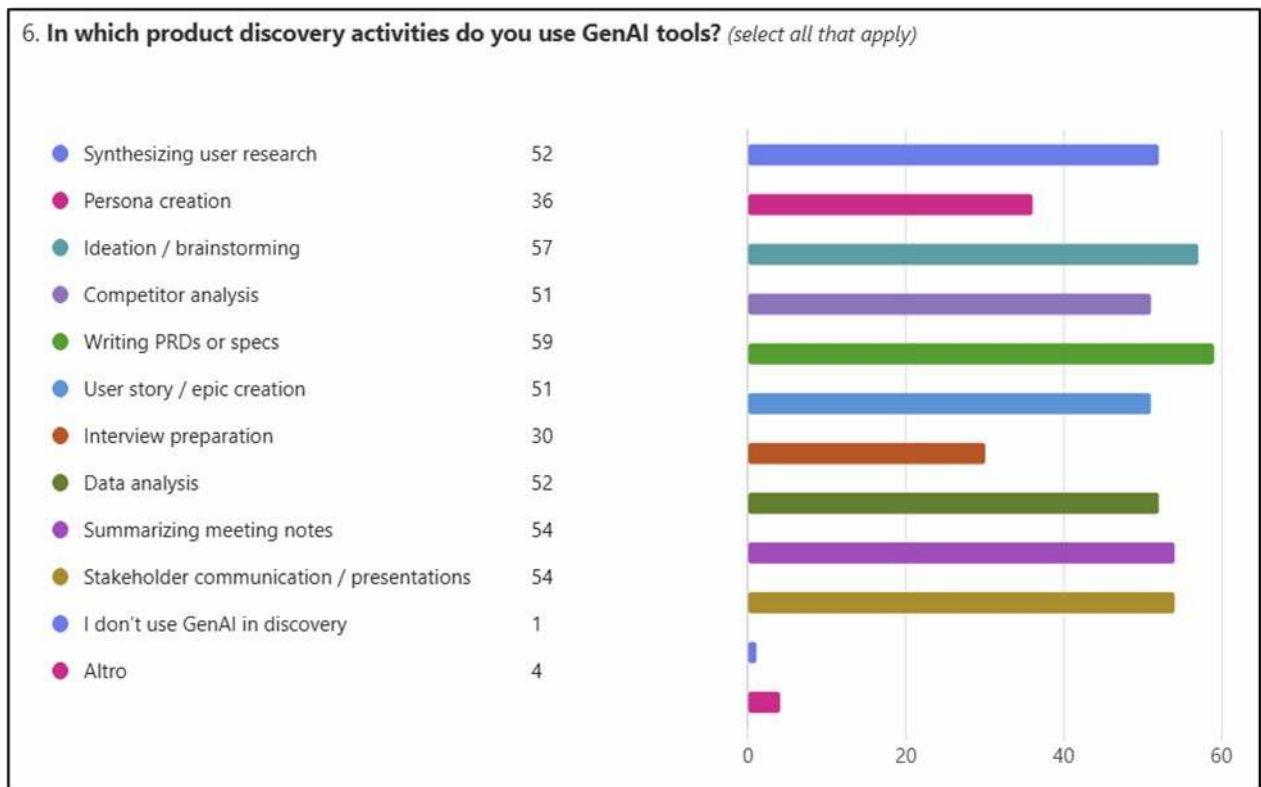


Figure 11: Discovery activities supported by GenAI.

Table 5.6 reflects the activity counts from the survey data. It shows that the strongest use case is writing PRDs or specs, closely followed by ideation and stakeholder communication, with meeting-note summarisation and synthesising user research also very prominent. The table is especially important because it makes the role of GenAI in product discovery concrete: it is not only an ideation tool, but also a drafting, synthesis, and coordination tool. The single I do not use GenAI in discovery response should be read as a minority dissent, not a meaningful adoption cluster.

Table 5.6: Discovery activities supported by GenAI.

Activity	Count
Synthesizing user research	52
Persona creation	36
Ideation / brainstorming	57
Competitor analysis	51
Writing PRDs or specs	59
User story / epic creation	51
Interview preparation	30
Data analysis	52

Summarizing meeting notes	54
Stakeholder communication / presentations	54
Code generation / vibe coding	1
I do not use GenAI in discovery	1

Figure 12 shows how often respondents use GenAI in product discovery, with daily and several times per week use dominating the sample.



Figure 12: Frequency of GenAI use.

Table 5.5 shows that it is highly concentrated at the top end. Daily use is the largest category, and several times per week is the second largest, while weekly, occasional, and rare use are far smaller. Visually, this indicates a strong adoption slope toward routine use rather than sporadic experimentation.

Table 5.7: Frequency of GenAI use in product discovery.

Frequency	Count
Daily	56
Several times per week	20
Weekly	3
Rarely or never	2
Occasionally	1

Figure 13 shows that most respondents agree or strongly agree that GenAI improves their efficiency, with strongly agree being the largest category.

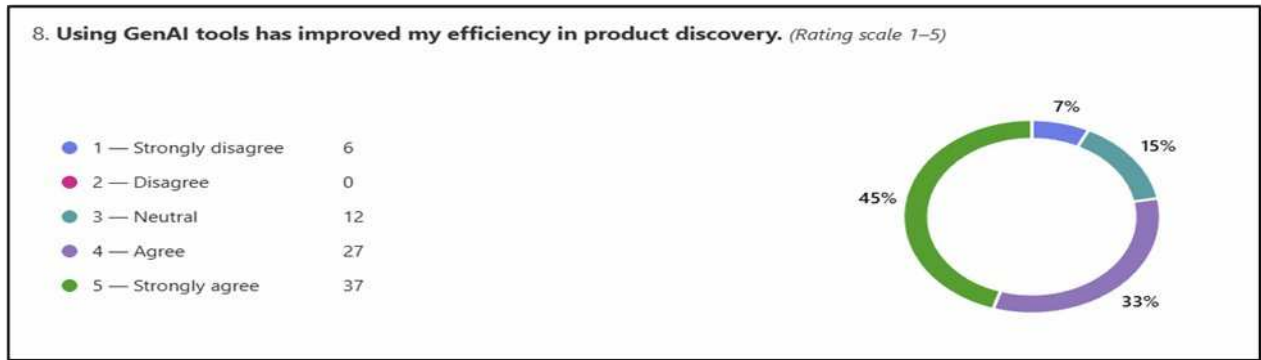


Figure 13: Using GenAI tools has improved my efficiency in product discovery.

Table 5.8 shows that respondents overwhelmingly perceive GenAI as improving efficiency in product discovery. The largest group selected Strongly agree, followed by Agree, while only a small minority selected Strongly disagree or Neutral. This indicates that GenAI is primarily valued for reducing manual effort and speeding up product discovery work.

Table 5.8: Using GenAI tools has improved my efficiency in product discovery.

Benefit	Count
1 — Strongly disagree	6
2 — Disagree	0
3 — Neutral	12
4 — Agree	27
5 — Strongly agree	37

5.4 Perceived Value (Q9-Q11, SRQ2)

The survey results show a broadly positive assessment of GenAI's value in product discovery. Respondents rated GenAI highly on output quality, workflow integration, and overall augmentation of the PM role, while trust remained noticeably lower than the usefulness-oriented items. This suggests that practitioners see GenAI as a useful support tool that improves productivity and output quality, but not as something that can be used without human review.

Figure 14 shows that most respondents agree or strongly agree that GenAI improves the quality of their product discovery outputs, with Agree forming the largest category and Strongly agree following closely behind.

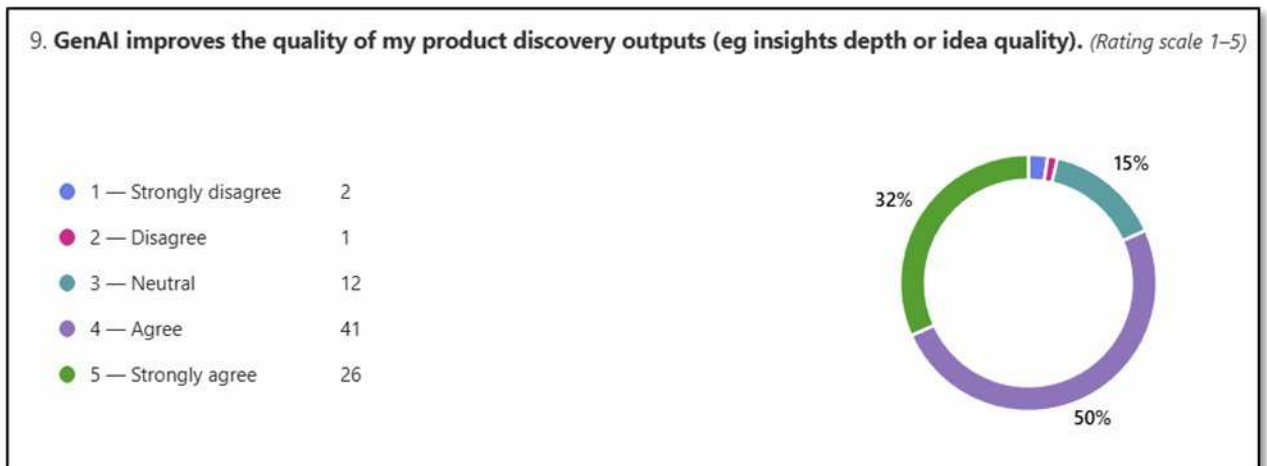


Figure 14: GenAI improves the quality of my product discovery outputs.

Table 5.9 shows that respondents strongly perceive GenAI as improving the quality of their product discovery outputs. The majority selected Agree or Strongly agree, while only a very small number selected Disagree or Strongly disagree. This indicates that GenAI is widely valued for improving the depth or quality of insights and ideas in product discovery.

Table 5.9: GenAI improves the quality of my product discovery outputs.

Response	Count
1 — Strongly disagree	2
2 — Disagree	1
3 — Neutral	12
4 — Agree	41
5 — Strongly agree	26

The mean score for this item is 4.07, calculated from the 1–5 coded Likert responses using the weighted average method.

Figure 15 shows that most respondents agree or strongly agree that GenAI tools are easy to integrate into their workflow, with Agree forming the largest category and Strongly agree following closely behind.



Figure 15: Respondents find it easy to integrate GenAI tools into their workflow.

Table 5.10 indicates that most respondents find GenAI reasonably easy to integrate into their workflow, although the score is slightly lower than the output-quality rating. The presence of a noticeable Neutral group suggests that some respondents still face practical barriers or uneven adoption conditions.

Table 5.10: I find it easy to integrate GenAI tools into my workflow

Response	Count
1 — Strongly disagree	1
2 — Disagree	3
3 — Neutral	17
4 — Agree	36
5 — Strongly agree	25

The mean score for this item is 3.99, calculated from the 1–5 coded Likert responses using the weighted average method.

Figure 16 shows that the main benefits respondents associate with GenAI are time saved on manual tasks, better-quality documentation, faster synthesis of research data, and clearer thinking or structure.

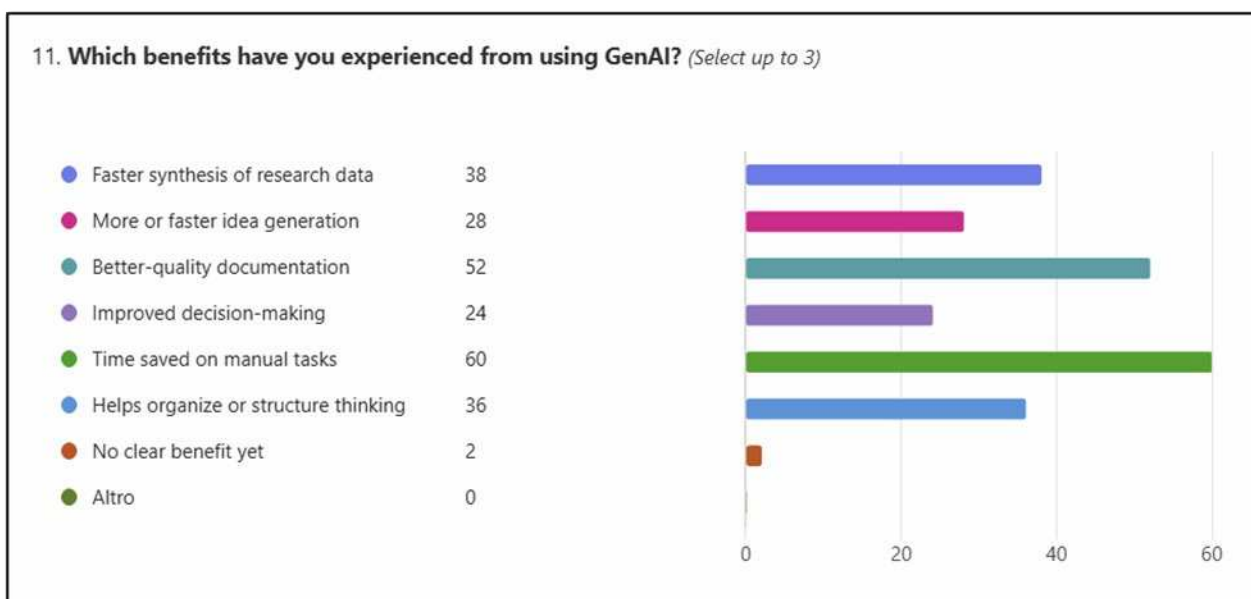


Figure 16: Which benefits have you experienced from using GenAI?

Table 5.11 shows that respondents overwhelmingly value GenAI for productivity and synthesis-related benefits. The strongest reported benefit is time saved on manual tasks, followed by better-quality documentation and faster synthesis of research data. This suggests that GenAI is used primarily to reduce repetitive cognitive labour and improve the speed and structure of discovery work.

Table 5.11: Which benefits have you experienced from using GenAI

Response	Count
Time saved on manual tasks	60
Better-quality documentation	52
Faster synthesis of research data	38
Helps organize or structure thinking	36
More or faster idea generation	28
Improved decision-making	24
No clear benefit yet	2

Figure 17 shows that the most common barriers are inaccurate or hallucinated outputs, data privacy and security concerns, company policy limits, and lack of product context.

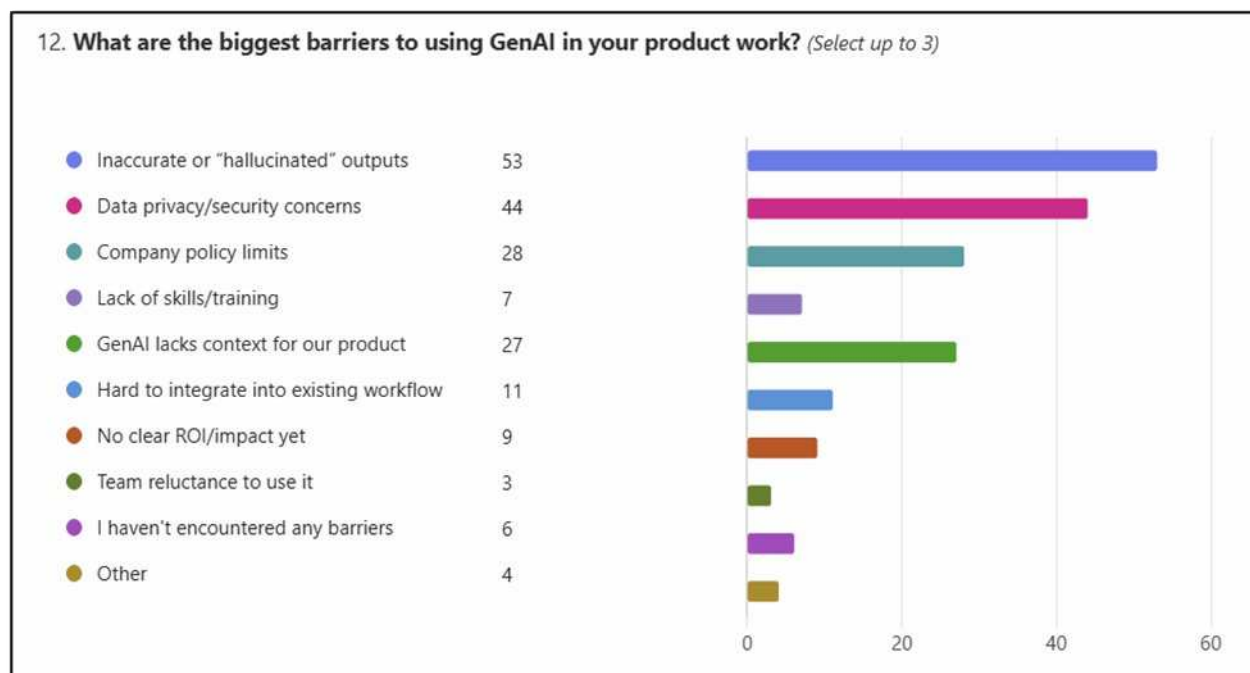


Figure 17: What are the biggest barriers to using GenAI in your product work?

Table 5.12 shows that respondents' biggest concerns are reliability and governance-related rather than purely technical. The dominant barrier is inaccurate or hallucinated outputs, followed by privacy/security concerns and company policy limits. This indicates that respondents are willing to use GenAI, but only within constraints that protect accuracy, context, and organisational compliance.

Table 5.12: Top reported barriers to GenAI use.

Barrier	Count
Inaccurate or "hallucinated" outputs	53
Data privacy/ security concerns	44
Company policy limits	28
GenAI lacks context for our product	27
Hard to integrate into existing workflow	11
No clear ROI/ Impact yet	9
Lack of skills/ Training	7
I haven't encountered any barriers	6
Team reluctance to use it	3
Not having enough tokens	1
Limited access to this tools with Client	1
Short memory of agents	1
No license yet for most employees.	1

Figure 18 shows that respondents are divided mainly between Neutral and Agree, with only a small minority selecting Disagree or Strongly agree.



Figure 18: Respondents trust the outputs from GenAI tools in the product discovery work

Table 5.13 shows that trust in GenAI outputs is moderate rather than strong. The largest groups are Neutral and Agree, which suggests that respondents are willing to use GenAI but still prefer to verify its outputs before relying on them fully. This pattern supports the broader finding that GenAI adoption is pragmatic and conditional.

Table 5.13: Respondents trust the outputs from GenAI tools in their product discovery work

Response	Count
1 — Strongly disagree	0
2 — Disagree	7
3 — Neutral	36
4 — Agree	36
5 — Strongly agree	3

The mean score for this item is 3.43. This is the lowest score in the perceived-value block and shows a cautious trust pattern.

Figure 19 shows that most respondents agree or strongly agree that GenAI augments their work rather than replacing parts of their role.

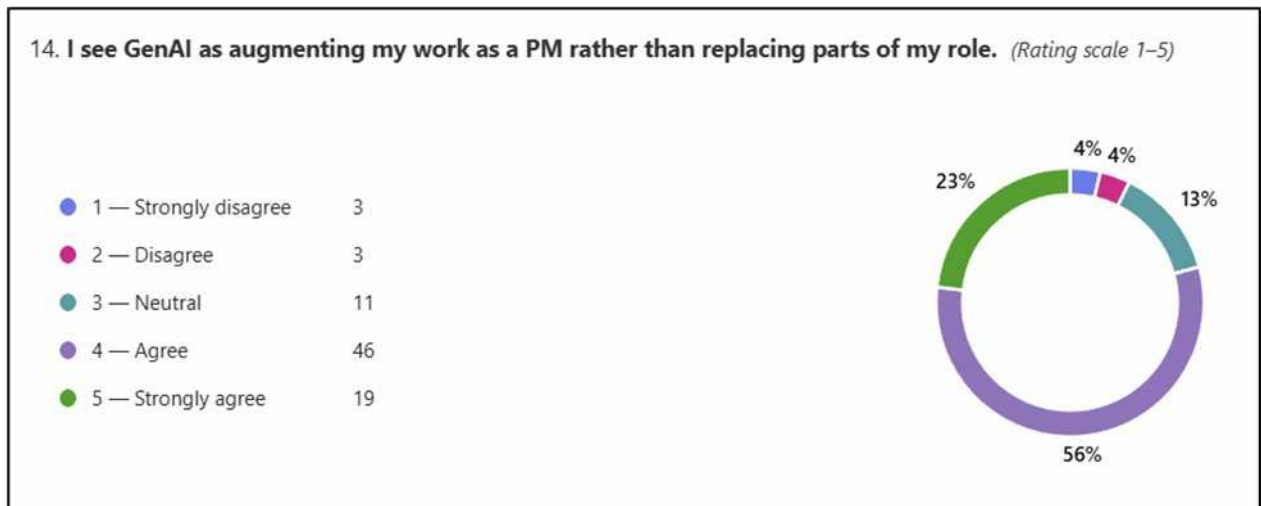


Figure 19: Respondents see GenAI as augmenting their work as a PM rather than replacing parts of my role.

The mean score for this item is 3.91. Table 5.14 indicates that respondents largely view GenAI as an augmentation tool rather than a replacement for PM work. The strong concentration in Agree and Strongly agree suggests that most practitioners believe GenAI supports their role by removing repetitive work and improving efficiency, while still leaving judgment, interpretation, and decision-making to the human PM.

Table 5.14: Respondents see GenAI as augmenting their work as a PM rather than replacing parts of my role

Response	Count
1 — Strongly disagree	3
2 — Disagree	3
3 — Neutral	11
4 — Agree	46
5 — Strongly agree	19

Table 5.15 summarises the average scores for the key perceived-value items, showing how respondents rated efficiency, quality, ease of integration, trust, and augmentation.

Table 5.15: Summary of perceived value

Metric	Value
Responses	82
Mean efficiency	4.09
Mean quality	4.07

Mean ease	3.99
Mean trust	3.43
Mean augmentation	3.91
Responses	82

These values show a coherent attitude structure. Efficiency and quality are the strongest signals, followed by augmentation and ease of integration, while trust is the lowest of the positive items. The overall interpretation is that GenAI is valued for the work it helps respondents do, but it is still not trusted enough to replace human judgement.

The means were calculated from the 1–5 Likert responses by converting each response into a numeric score and computing the weighted average across all 82 responses. For example, the quality item was calculated from the response distribution using the standard weighted-mean formula. The same method was used for efficiency, ease, trust, and augmentation.

The mean values correspond to the following questions:

- Mean efficiency = 4.09 came from “Using GenAI tools has improved my efficiency in product discovery.”
- Mean quality = 4.07 came from “GenAI improves the quality of my product discovery outputs.”
- Mean ease = 3.99 came from “I find it easy to integrate GenAI tools into my workflow.”
- Mean trust = 3.43 came from “I trust the outputs from GenAI tools in my product discovery work.”
- Mean augmentation = 3.91 came from “I see GenAI as augmenting my work as a PM rather than replacing parts of my role.”

For example, the mean for “GenAI improves the quality of my product discovery outputs” was calculated by assigning values from 1 to 5 to the response options and then applying the weighted average formula:

$$\text{Mean} = \frac{\sum(\text{rating} \times \text{number of responses at that rating})}{\text{total responses}}$$

Using the observed counts:

- 2 responses at rating 1.
- 1 response at rating 2.
- 12 responses at rating 3.
- 41 responses at rating 4.
- 26 responses at rating 5.

So the calculation is:

$$\frac{(1 \times 2) + (2 \times 1) + (3 \times 12) + (4 \times 41) + (5 \times 26)}{82} = \frac{334}{82} = 4.07$$

The same method was used for the other Likert-scale items. This is why the mean is a weighted average rather than a simple average of categories.

On the perception items, the mean scores are efficiency 4.09, quality 4.07, ease of integration 3.99, trust 3.43, and augmentation 3.91. The pattern shows that respondents view GenAI as strongly helpful for productivity and output quality, while trust remains noticeably lower than the usefulness-oriented measures. This gap is important because it indicates pragmatic adoption: respondents value GenAI for what it enables them to do, but they still rely on human judgement rather than accepting its outputs unquestioningly.

5.5 Barriers and Trust (Q12–Q13, SRQ3)

As shown in Table 5.12, the main barriers to GenAI use are inaccurate or hallucinated outputs, data privacy and security concerns, company policy limits, lack of product context, and difficulty integrating the tools into existing workflows. These barriers show that respondents are not rejecting GenAI, but they are using it with awareness of its limitations and organisational constraints.

Table 5.13 indicates that trust in GenAI outputs is moderate rather than absolute. When read alongside Table 5.12, the results show that respondents are willing to use GenAI, but they remain aware of issues such as hallucinated outputs, privacy concerns, policy limits, and limited contextual accuracy. This suggests that GenAI adoption is pragmatic and still supported by human review.

Taken together, the barrier and trust findings indicate pragmatic adoption. PMs appear willing to use GenAI because it is useful for synthesis, drafting, and acceleration, but they still expect human review before accepting outputs as final. In that sense, trust is not absent; rather, it is conditional and shaped by the need to verify accuracy, context, and policy fit.

5.6 PM Role Evolution and Skills (Q14-Q15, SRQ4)

As shown in Table 5.14, the role-evolution findings suggest that GenAI is changing how PMs work, but not replacing the PM role itself. Instead, it appears to be redistributing effort toward higher-value tasks such as evaluation, prioritisation, and strategic decision-making. This reinforces the idea that future PM capability will depend on a blend of human judgment and GenAI-assisted efficiency.

Figure 20 shows that critical thinking, strategic thinking, and prompt engineering are the most frequently reported skills that have become more important with GenAI use.

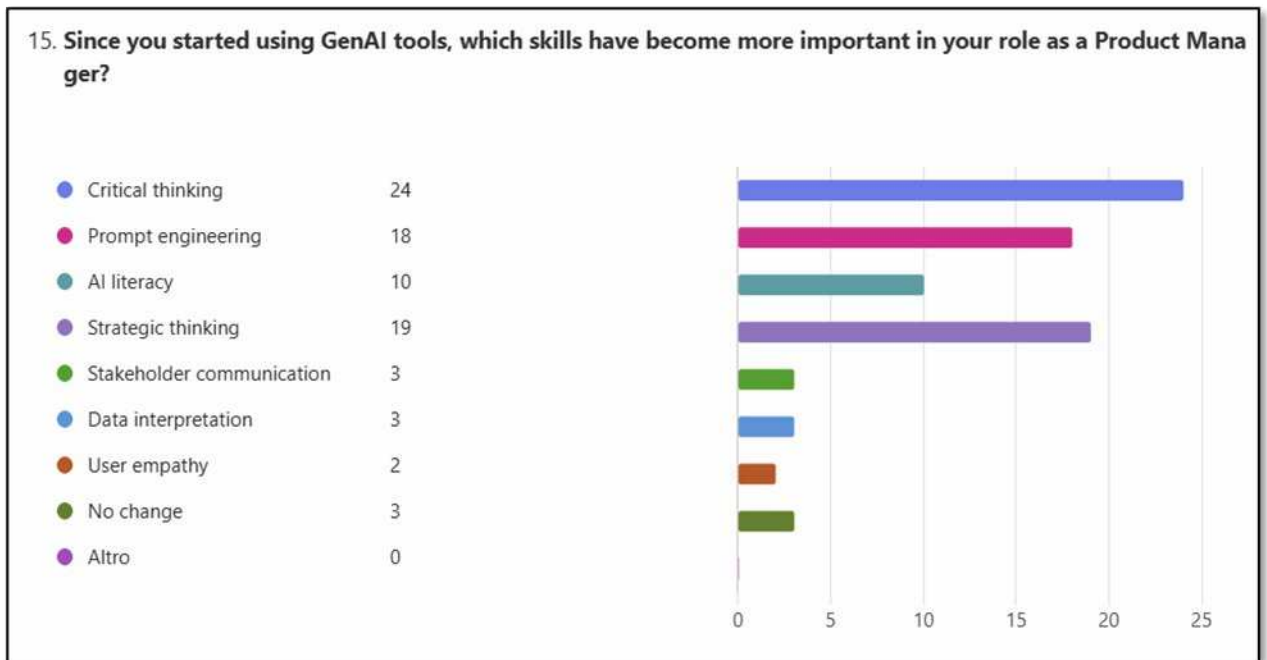


Figure 20: Since respondents started using GenAI tools, which skills have become more important in their role as a PM?

Table 5.16 indicates that GenAI adoption raises the value of human judgement rather than reducing it. Skills related to interpretation, checking, and strategic framing appear more important than purely technical prompting. This is consistent with the augmentation perspective introduced in Chapter 2 and with the AI Fluency idea that good use depends on disciplined human oversight.

Table 5.16: Skills becoming more important with GenAI use.

Skill	Count
Critical thinking	24
Strategic thinking	19
Prompt engineering	18
AI literacy	10
Stakeholder communication	3
Data interpretation	3
No change	3
User empathy	2

The most defensible interpretation is that PMs do not see GenAI as replacing product judgement; they see it as shifting the work toward higher-level thinking, stronger validation habits, and clearer communication. Where possible, the final thesis should add cross-tabs such as role by trust and company size by frequency so this claim is supported with stratified evidence rather than only general narrative.

5.7 Cross-Tabular Analysis

Cross-tabulation was used to examine whether GenAI adoption, trust, usefulness, and augmentation differ by respondent context. This is important because the study was designed to go beyond overall averages and to look for meaningful patterns across role, company size, and experience. The tables below are therefore read descriptively rather than inferentially, since the sample is exploratory and not intended for population-level generalisation.

5.7.1 Role by frequency of use

Table 5.13 shows that GenAI use is concentrated in the more experienced PM roles, especially Senior PMs and PMs. Senior PMs record the largest counts in both daily use and several-times-per-week use, while PMs also show strong regular usage. Product Owners and Product Leads are present in the distribution, but they appear in smaller numbers overall.

Table 5.17: Role by frequency of GenAI use.

Role	Daily	Occasionally	Rarely or never	Several times per week	Weekly
Head of Product / VP of Product	5	0	0	1	0
Product Lead / Principal PM	6	1	0	2	0
Product Manager	13	0	0	5	1
Product Owner	6	0	0	3	0
Senior Product Manager	26	0	2	9	2

The conclusion is that GenAI is not limited to experimentation at the junior level; it is already embedded in the working routines of senior product practitioners. This suggests that adoption is linked to real discovery responsibility, not just curiosity about new tools. In other words, the more strategically involved PM groups are also the ones using GenAI most consistently.

5.7.2 Company size by use

Table 5.14 shows that large enterprises account for the largest share of regular GenAI use. The enterprise group has the highest daily usage count and the strongest several-times-per-week usage, while startups, scale-ups, and mid-size firms also use GenAI but with smaller absolute counts.

Table 5.18: Company size by frequency of GenAI use.

Company size	Daily	Occasionally	Rarely or never	Several times per week	Weekly
Large enterprise (1,000+ employees)	25	0	2	11	2
Mid-size company (251–1,000 employees)	13	0	0	2	0
Scale-up (51–250 employees)	7	1	0	4	0
Startup (1–50 employees)	11	0	0	3	1

The conclusion is that adoption appears strongest where GenAI can be absorbed into structured workflows and repeated discovery processes. This pattern likely reflects the combination of higher headcount, more formal collaboration, and more frequent documentation work in larger organisations. It also suggests that enterprise context does not block GenAI use, but it shapes how systematically the tools are embedded.

5.7.3 Role by trust level

Table 5.15 indicates that trust is moderate across all roles, rather than sharply different by seniority. Senior PMs and PMs dominate the middle-to-high trust bands, while very low trust is limited across the sample. The distribution does not show one role group as overwhelmingly more trusting than the others.

Table 5.19: Role by trust level.

Role	Low	Mid	High	Very high
Head of Product / VP of Product	0	2	4	0
Product Lead / Principal PM	4	3	2	0
Product Manager	0	10	8	1
Product Owner	1	4	4	0
Senior Product Manager	2	17	18	2

The conclusion is that trust in GenAI is not simply a function of seniority. Even experienced practitioners remain cautious, which supports the interpretation that PMs use GenAI with verification habits rather than blind reliance. This fits the broader survey pattern in which usefulness is high but confidence remains conditional.

5.7.4 Company size by trust level

Table 5.16 shows that trust is present across all company sizes, but it is most concentrated in the mid and high categories. Large enterprises have substantial numbers in both Mid and High trust, while startups and scale-ups also lean toward those bands rather than toward low trust.

Table 5.20: Company size by trust level.

Company size	Low	Mid	High	Very high
Large enterprise (1,000+ employees)	2	18	18	2
Mid-size company (251–1,000 employees)	2	5	7	1
Scale-up (51–250 employees)	3	6	3	0
Startup (1–50 employees)	0	7	8	0

The conclusion is that organisational context affects trust, but not in a simple or deterministic way. Larger firms may have stronger governance and security expectations, yet respondents in those settings still report workable trust levels. This suggests that confidence in GenAI is shaped by both tool quality and the surrounding organisational environment.

5.7.5 Experience by perceived usefulness

Table 5.17 shows that perceived usefulness is strongest among the more experienced respondents. The 5–10 years and 10+ years groups contribute the most High and Very high ratings, which indicates that usefulness is not just a beginner effect. Less experienced respondents also see value, but their responses are fewer and less concentrated at the top end.

Table 5.21: Experience by perceived usefulness.

Experience	Low	Mid	High	Very high
10+ years	1	1	4	12
1–3 years	2	2	1	2
3–5 years	1	1	7	9
5–10 years	2	7	14	14
Less than 1 year	0	1	1	0

The conclusion is that experience strengthens the ability to recognise where GenAI helps in real product work. More seasoned PMs appear better able to connect GenAI use with faster synthesis, clearer structure, and reduced manual effort. This supports the idea that usefulness is grounded in applied workflow experience rather than in abstract enthusiasm.

5.7.6 Role by augmentation view

Table 5.18 strongly supports the augmentation interpretation across all roles. High and Very high responses dominate, especially among Senior PMs, PMs, and Product Leads. Low responses are few, which means the sample largely agrees that GenAI supports rather than substitutes PM work.

Table 5.22: Role by augmentation view.

Role	Low	Mid	High	Very high
Head of Product / VP of Product	2	0	2	2
Product Lead / Principal PM	0	2	5	2
Product Manager	0	1	15	3
Product Owner	0	1	7	1
Senior Product Manager	4	7	17	11

The practical conclusion is that respondents see GenAI as a support layer for product discovery, not as a replacement for product judgement. The table suggests that PMs expect GenAI to assist with preparatory and repetitive work while leaving interpretation, prioritisation, and decision-making to humans. This is one of the clearest indicators that the role is being reshaped rather than removed.

The cross-tabulations presented earlier in this section examined adoption, trust, perceived value, and augmentation by role, company size, and experience. Because the survey deliberately recruited a geographically distributed sample, region is the fourth segmentation variable specified for cross-tabulation in Chapter 3 and is analysed here. Region matters analytically for three reasons. First, it tests whether the headline adoption story is an artefact of one dominant market or a pattern that recurs across very different professional environments. Second, regional context plausibly shapes adoption through differences in market maturity, data-protection regimes, tooling access, and work culture. Third, prior practitioner research (for example the Nigerian product-management study [27]) suggests that human–AI collaboration is enacted differently across regional contexts, so an absence of strong regional variance would itself be a meaningful finding. As with the other cross-tabs, the tables below are read descriptively rather than inferentially, because the sample is exploratory and several regional cells are small.

5.7.7 Region by frequency of use

Table 5.23 reports the frequency of GenAI use within each region. The sample spans six regions, with Asia-Pacific (24) and Europe (22) the largest, followed by North America (17) and the Middle East & Africa (13); Australia & Oceania (4) and Latin America & Caribbean (2) are present but too small to interpret on their own. The striking feature of the

table is uniformity rather than divergence. Daily or several-times-per-week use accounts for at least 88% of respondents in every substantive region (APAC 96%, Europe 91%, North America 88%, MEA 92%). The only visible difference is that the Middle East & Africa group skews toward several-times-per-week rather than strictly daily use (46% daily versus roughly 71–73% in the other large regions), which is consistent with the slightly more governance-constrained accounts seen in the open-ended responses from enterprise respondents. The descriptive conclusion is that routine adoption is geographically broad: high-frequency use is the norm across all well-represented regions, not a feature of a single market.

Table 5.23: Region by frequency of GenAI use (counts).

Region	Daily	Several / week	Weekly	Occasionally	Rarely / never	N
Asia-Pacific (APAC)	17	6	0	1	0	24
Europe	16	4	1	0	1	22
North America (US & Canada)	12	3	2	0	0	17
Middle East & Africa (MEA)	6	6	0	0	1	13
Australia &Oceania	3	1	0	0	0	4
Latin America & Caribbean	2	0	0	0	0	2
All regions	56	20	3	1	2	82

5.7.8 Role by perceived value, trust, and augmentation

Table 5.24 reports the mean Likert scores (1–5) for the five attitudinal items within each role group. Two patterns stand out. First, perceived value varies by role but remains broadly strong across all groups: Senior PMs and Heads of Product / VP of Product report the highest efficiency and output-quality scores, while Product Leads / Principal PMs are slightly lower on most measures. This suggests that practitioners with more decision-

making responsibility or broader strategic scope may extract more visible value from GenAI in product discovery.

Second, trust remains cautious across all role groups and never rises above the usefulness-related items. Even in the highest-scoring role groups, trust stays below efficiency, quality, and integration, which indicates that GenAI is being used as a supported workflow tool rather than a fully trusted decision-maker. The augmentation scores are consistently high across all roles, reinforcing the view that respondents see GenAI as assisting product work rather than replacing it.

Table 5.24: Role by mean Likert score (1 - 5) for the key attitudinal items.

Role	N	Efficiency	Quality	Ease of integ.	Trust	Augmentation
Senior Product Manager	39	4.10	4.08	4.03	3.51	3.82
Product Manager	19	3.89	4.00	4.11	3.53	4.11
Product Lead / Principal PM	9	3.78	4.00	3.56	2.78	4.00
Product Owner	9	4.22	4.00	3.67	3.33	4.00
Head of Product / VP of Product	6	4.83	4.50	4.50	3.67	3.67
All roles	82	4.09	4.07	3.99	3.43	3.91

5.7.9 Company size by perceived value, trust, and augmentation

Table 5.25 reports the mean Likert scores (1–5) for the five attitudinal items within each company-size group. Two patterns stand out. First, perceived value is generally high across all company sizes, but the scale-up group shows the weakest profile, especially on efficiency and trust. This is notable because it suggests that the adoption friction point is not necessarily the smallest firms or the largest enterprises, but the organisations in between, where processes may be growing faster than governance or workflow maturity. Figure 21 shows broadly high perceived value across company sizes, with the scale-up

group appearing slightly weaker, especially on trust, while startups and mid-size firms look somewhat stronger.

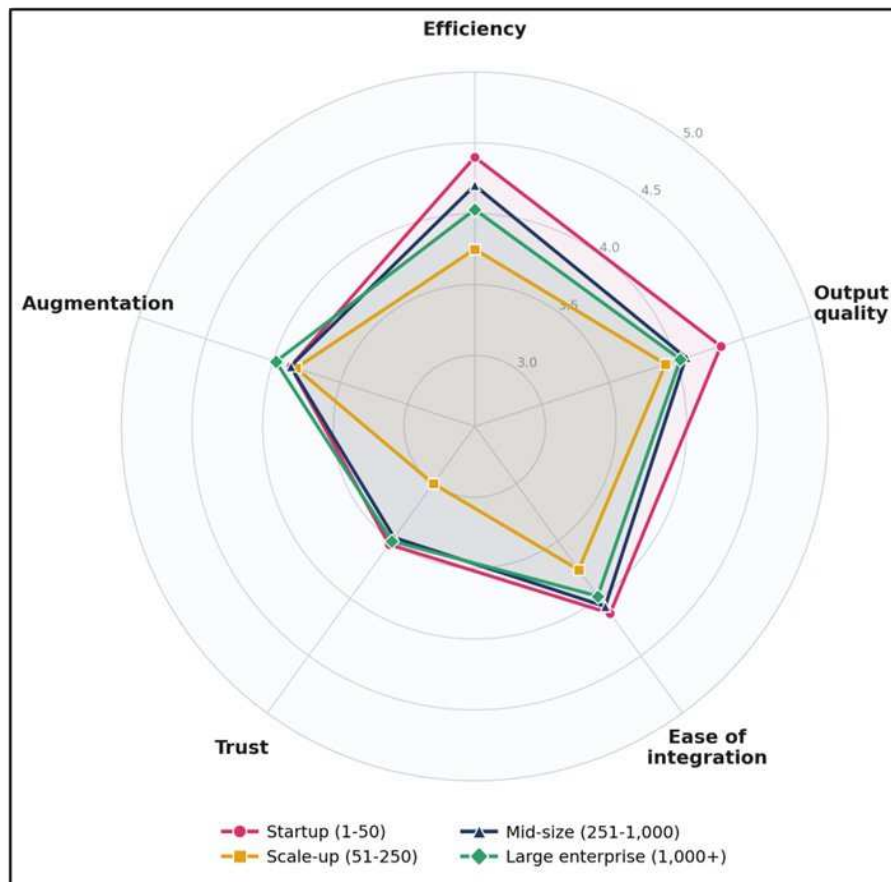


Figure 21: Perceived-value profile across company sizes (mean rating, 1–5, on the five attitudinal items).

Second, the startup, mid-size, and large-enterprise groups all show broadly similar value profiles, with efficiency, quality, and augmentation close to or above 4.0. Trust remains the lowest item in every company-size group, confirming that even where GenAI is seen as useful and augmentative, it is still approached cautiously. The large-enterprise group, which is the biggest in the sample, shows steady but not exceptional scores, implying that enterprise adoption is established rather than exuberant.

Table 5.25 : Company size by mean Likert score (1 - 5) for the key attitudinal items.

Company Size	N	Efficiency	Quality	Ease of integ.	Trust	Augmentation
Startup (1–50)	15	4.40	4.33	4.13	3.53	3.87

Company Size	N	Efficiency	Quality	Ease of integ.	Trust	Augmentation
employees)						
Scale-up (51–250 employees)	12	3.75	3.92	3.75	3.00	3.83
Mid-size company (251–1,000 employees)	15	4.20	4.07	4.07	3.47	3.87
Large enterprise (1,000+ employees)	40	4.03	4.03	3.98	3.50	3.98
All company sizes	82	4.09	4.07	3.99	3.43	3.91

5.7.6 Experience by perceived value, trust, and augmentation

Table 5.26 reports the mean Likert scores (1–5) for the five attitudinal items within each experience group. Two patterns stand out. First, perceived value rises with experience in a non-linear way: the 3–5 years and 10+ years groups report the strongest scores on efficiency and output quality, while the least experienced respondents are consistently lower across the board. This suggests that more seasoned practitioners are better able to extract value from GenAI in product discovery, likely because they can supply better context and judge outputs more effectively. Figure 22 visually confirms the table pattern: perceived value is strongest among the 3–5 years and 10+ years groups, while trust remains the lowest item across all experience bands.

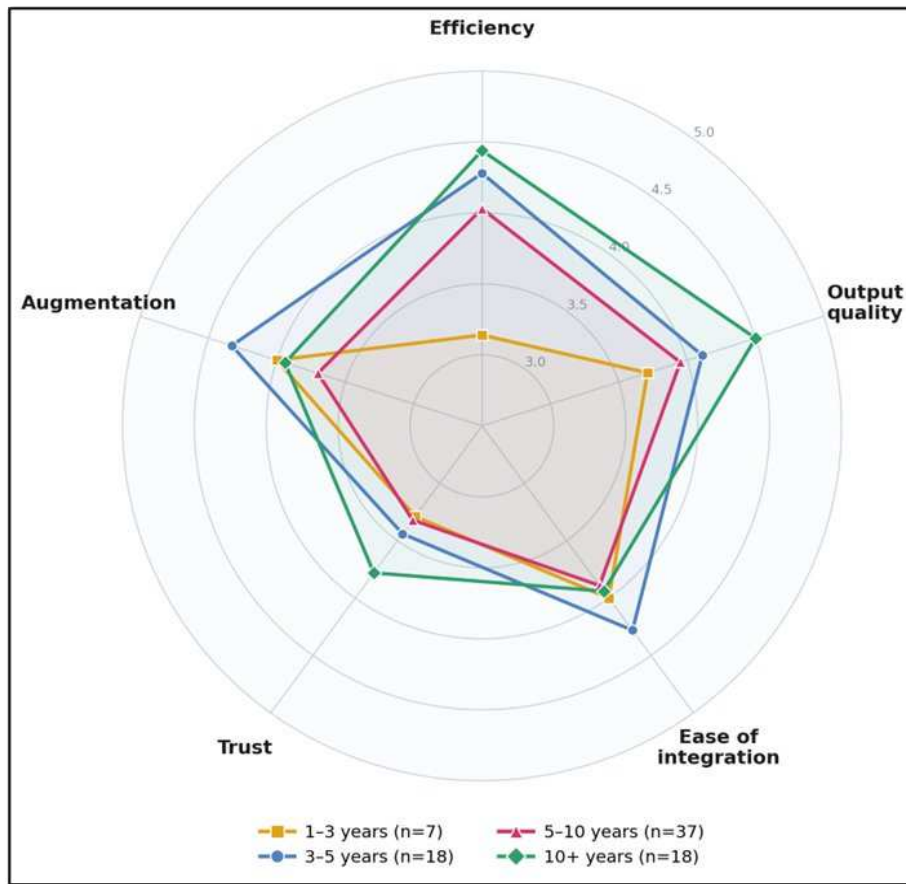


Figure 22: Perceived-value profile across Experience (mean rating, 1–5).

Second, trust remains cautious across all experience bands and does not increase in step with experience as strongly as usefulness does. Even among the 10+ years group, trust remains below the usefulness-related items, which indicates that experience improves perceived value but does not eliminate the need for human validation. The 5–10 years band, despite being the largest group, shows slightly softer scores than the 3–5 years and 10+ years groups, which suggests that the middle seniority range may be using GenAI routinely but with more qualified confidence.

Table 5.26 : Experience by mean Likert score (1- 5) for the key attitudinal items.

Experience	N	Efficiency	Quality	Ease of integ.	Trust	Augmentation
Less than 1 year	2	3.50	3.50	3.50	2.50	3.50
1–3 years	7	3.14	3.71	4.00	3.29	4.00
3–5 years	18	4.28	4.11	4.28	3.44	4.33

Experience	N	Efficiency	Quality	Ease of integ.	Trust	Augmentation
5–10 years	37	4.03	3.95	3.89	3.32	3.70
10+ years	18	4.44	4.50	3.94	3.78	3.94
All experience groups	82	4.09	4.07	3.99	3.43	3.91

5.7.8 Region by perceived value, trust, and augmentation

Table 5.27 reports the mean Likert scores (1–5) for the five attitudinal items within each region. Two patterns stand out. First, the trust–usefulness gap that characterises the whole sample reappears inside every region: in each of the four well-represented regions, mean trust (3.29–3.62) is the lowest of the five items and sits clearly below efficiency, quality, and augmentation. The conditional, verify-before-relying posture is therefore not a regional quirk but a structural feature of how these practitioners relate to GenAI. Second, the differences between regions are small and do not form a clean ranking. North American respondents rate output quality highest (4.35) and APAC respondents rate it lowest (3.88), while MEA respondents report the highest trust among the large regions (3.62) despite their slightly lower daily-use share. The two smallest regions show the highest means, but with only four and two respondents respectively these values are unstable and should not be over-interpreted. Taken honestly, the regional signal is weak: region nuances the picture at the margins but does not overturn the core finding that usefulness is high, augmentation is broadly endorsed, and trust remains the consistently cautious outlier.

Table 5.27: Region by mean Likert score (1–5) for the key attitudinal items.

Region	N	Efficiency	Quality	Ease of integ.	Trust	Augmentation
Asia-Pacific (APAC)	24	3.79	3.88	4.04	3.29	3.96
Europe	22	4.14	3.95	3.82	3.36	3.91
North America (US & Canada)	17	4.12	4.35	4.24	3.53	3.82
Middle East & Africa (MEA)	13	4.15	4.08	3.69	3.62	4.00
Australia & Oceania	4	4.75	4.75	4.25	3.25	3.75
Latin America &	2	5.00	4.00	4.50	4.00	4.00

Region	N	Efficiency	Quality	Ease of integ.	Trust	Augmentation
Caribbean						
All regions	82	4.09	4.07	3.99	3.43	3.91

Interpreted cautiously, the regional comparison contributes mainly by strengthening the external credibility of the other findings. Because the same adoption-high, trust-cautious, augmentation-positive pattern holds across markets as different as Western Europe, North America, South and Southeast Asia, and the Middle East & Africa, the result is unlikely to be an artefact of a single regulatory or cultural setting. Where regional texture does appear (the MEA group's tilt toward several-times-per-week use and its references to licensing and policy constraints), it is best read as a difference in the governance conditions surrounding adoption rather than in practitioners' underlying appetite for the tools. The honest summary is that regional variation is real but modest, and that it adds nuance to, rather than competes with, the role-, size-, and experience-based patterns reported above.

5.7.9 Overall pattern across the cross-tabulations

Read together, the cross-tabulations by role, company size, experience, trust, augmentation, and region tell a single coherent story: GenAI adoption in product discovery is context-dependent in its intensity but remarkably consistent in its shape. Seniority and organisational environment both matter. The most strategically involved role groups (Senior PMs, PMs, and Heads of Product) are also the heaviest regular users, and larger and mid-size organisations show the most systematic embedding of the tools into repeated discovery workflows, indicating that enterprise context channels rather than blocks adoption. Experience deepens perceived usefulness, with the 5–10-year and 10+-year bands contributing the most high-value ratings, which suggests that the benefit is grounded in applied workflow judgement rather than novelty enthusiasm.

At the same time, the attitudinal cross-tabs show that trust is cautious but workable across every subgroup: no role, company size, or region emerges as either uniformly distrustful or uncritically reliant, and mean trust sits below mean usefulness everywhere. The augmentation view is broadly shared across all roles and regions, with high and very-high responses dominating. Regional variation, where present, adds nuance (a slightly more governance-constrained adoption rhythm in the Middle East & Africa, marginally higher quality ratings in North America) without overturning the main pattern. The cross-tab evidence therefore reinforces the central claim of this thesis: GenAI has become a practical, routine workflow-support tool in product discovery, adopted most systematically where discovery responsibility and organisational structure are greatest, valued for efficiency and synthesis far more than it is trusted to act unsupervised, and understood across the board as augmenting rather than replacing the PMs judgement.

This consistency is shown visually in Figure 21, which plots the five attitudinal items as a perceived-value profile for each company-size group. The profiles are strikingly congruent in shape: efficiency, output quality, ease of integration, and augmentation all sit close to or above 4.0 across every group, while every profile contracts inward on the trust axis. Trust is therefore the universal low point of the value profile, not a feature of any one organisational setting. The only visible divergence is the scale-up group, whose profile sits slightly inside the others on efficiency, ease of integration, and trust (mean trust 3.00), suggesting that organisations in a rapid-growth phase may experience more friction or less settled governance around GenAI use. The radar makes concrete what the cross-tabs report numerically: a shared, high-value, trust-cautious posture that holds across organisational contexts.

5.8 Thematic Analysis of Open-Ended Responses

The quantitative results establish what practitioners do with GenAI and how strongly they rate it. The purpose of this qualitative strand is complementary: to recover the meaning behind those ratings in respondents' own words, to surface mechanisms the closed questions cannot capture, and to test whether the language practitioners use independently corroborates the survey's numerical patterns. This makes the study genuinely mixed-method in the sense defined in Chapter 3 [29], with the qualitative analysis deepening rather than merely illustrating the quantitative findings. The analysis follows the six-phase reflexive thematic-analysis procedure of Braun and Clarke [32].

5.8.1 Open-ended questions included

Four free-text items were analysed: (i) A request to describe one concrete example of using GenAI in discovery, including what worked and what did not; (ii) A question on the single most significant limitation experienced; (iii) A short prompt on how GenAI has changed the respondent's work as a PM; and (iv) An optional catch-all inviting anything the survey had not asked. The first three are the analytical core, because they were answered by the large majority of respondents; the optional fourth item was sparsely completed and is used only as supporting corroboration.

5.8.2 Data cleaning prior to coding

Open-ended fields were cleaned before any coding took place, so that themes were built only from meaningful responses. Cleaning followed two explicit rules applied consistently across items. First, blank and null entries were removed. Second, non-informative responses were removed: pure placeholders and garbage tokens (for example "Na", "na", "Nil", ".", "..", "/", "-", and bare "no"), single-character noise, and entries that did not address the question. Responses that were brief but substantive (for example naming a specific discovery activity such as "competitor analysis" or "persona creation") were retained, because they still carry codeable meaning. Crucially, negative-but-relevant responses (such as a respondent stating they had not yet used GenAI in discovery) were retained and read, since they inform the limitations and adoption-barrier themes rather than being discarded as noise. Table 5.28 summarises the cleaning outcome for each item.

Table 5.28 : Open-ended response cleaning summary.

Open-ended item	Total	Blank	Removed (non-informative)	Retained for coding
Example of GenAI use in discovery	82	6	7	69
Most significant limitation	82	5	6	71
How GenAI changed your work	82	4	7	71
Additional comments (optional)	82	60	10	12

After cleaning, the three core items retained 69, 71, and 71 substantive responses respectively, giving a robust corpus for a study of this size. Because individual respondents frequently expressed the same idea across more than one item (for example describing time savings under both the “example” and “changed my work” questions), the responses were pooled at the respondent level for coding, so that each theme’s frequency reflects the number of distinct practitioners who raised it at least once rather than a raw count of mentions.

5.8.3 Coding and theme development

The retained responses were read in full at least twice for familiarisation. An initial set of descriptive codes was then generated inductively, close to the respondents’ wording (for example “turned days of synthesis into an hour”, “lacks our internal context”, “confident but inaccurate”, “still needs human validation”, “shifted me to strategic editing”). Repeated ideas were grouped, and overlapping or near-synonymous codes were merged into broader candidate themes, for instance, “saves time”, “faster”, “accelerated”, and “10x productivity” were merged into a single time-saving theme, while “no company context”, “doesn’t know our technical debt”, and “generic output” were consolidated under a lack-of-context theme. Candidate themes were reviewed against the coded extracts and against the full corpus to confirm they were internally coherent and distinct from one another. Themes were retained when they were supported by multiple respondents and were analytically meaningful for the research questions; idiosyncratic one-off remarks were noted but not elevated to themes. The final structure organises eleven sub-themes under three higher-order domains (value and capability, limitations and risks, and practice and role change), as shown in Table 5.29 and visualised in Figure 22.

Table 5.29: Thematic coding framework (themes, representative codes, and approximate respondent frequency).

Domain	Theme	Representative codes	Approx. respondents (N=82)
Value & capability	Time saving and acceleration	“saves time”, “faster”, “days into an hour”, “10x”	~52 (63%)
Value & capability	Drafting and documentation	PRDs, user stories, tickets, decks, roadmaps	~38 (46%)
Value & capability	Ideation and prototyping	brainstorming, hypotheses, personas, prototypes	~26 (32%)
Value & capability	Synthesis and summarisation	clustering feedback, summarising research, patterns	~22 (27%)
Value & capability	Reduced manual / repetitive effort	“less manual work”, “tedious tasks”, “redundant work”	~20 (24%)
Limitations & Risks	Lack of organisational context	“doesn’t know our context”, “generic”, domain nuance	~36 (44%)
Limitations & risks	Hallucination and inaccuracy	“hallucinates”, “confident but wrong”, unreliable	~24 (29%)
Limitations & risks	Governance, privacy, access limits	policy, privacy/security, licences, token limits	~23 (28%)
Practice & role change	Human validation and discernment	“double-check”, “not a decision-maker”, oversight	~22 (27%)
Practice & role change	Strategic shift and	“strategic editor”, “orchestrating”, “thought	~20 (24%)

Domain	Theme	Representative codes	Approx. respondents (N=82)
	orchestration	partner”	
Practice & role change	Role intensification	“work has intensified”, “expectations rose”	noted (minor)

The frequencies in Table 5.29 are approximate respondent counts derived from coding the cleaned free-text corpus, and are reported as interpretive rather than precisely measured quantities; they indicate the relative prominence of each theme, not an exact tally. They should be read alongside, not in competition with, the closed-question results.



Figure 23: Thematic map of the open-ended responses, organised into three domains and eleven sub-themes.

5.8.4 Themes in respondents' own words

Within the value domain, time saving was the single most prominent idea, expressed by roughly two-thirds of respondents and frequently tied to synthesis and drafting work. One respondent described feeding ten hours of interview transcripts to a model and noting that it “turned days of synthesis into an hour-long task”, while another reported that AI-assisted PRD creation brought a task “down from 4–5 hours to 30 minutes.” The benefit is consistently framed as compressing the distance between a question and a usable first draft. The limitations domain was dominated by a single, strikingly consistent concern: lack of organisational context. Respondents repeatedly observed that GenAI offers “generic excellence” but “lacks our internal context”, cannot see the team’s technical debt or private strategic decisions, and therefore produces output that is plausible but insufficiently specific. This shaded directly into the hallucination theme, captured bluntly by one

respondent's remark that the tool "hallucinates when I did competitive research", and into a governance theme expressed through references to company policy, privacy and security, licensing, and token limits.

The practice-and-role domain is where the qualitative data add the most beyond the numbers. A recurring account is that GenAI shifts the centre of gravity of the role from production toward judgement and orchestration. One respondent described the change as moving "from content creator to strategic editor" and eliminating "blank-page syndrome"; another wrote that the technology had "de-centred" them from their own process, leaving them "orchestrating systems and processes more than ever before and having to think more strategically." Coupled with this is an insistence on human validation: respondents described GenAI as "a thinking and structuring assistant" whose "output always needs human validation", and warned that "you can't outsource judgment; you still need someone who knows when to override the model." A smaller but vivid strand reported that the technology has intensified rather than lightened the work, with one respondent noting that although work became more streamlined, "expectations also rose... so it's more intense now." These accounts give human texture to the augmentation result and explain why trust remains conditional even where usefulness is high.

5.8.5 Relating the themes to the quantitative findings

The qualitative themes map closely onto the closed-question results, which is the central test of a convergent mixed-method design. Table 5.30 sets out this correspondence. The match is strong: the dominant qualitative theme (time saving) aligns with efficiency being the highest-rated benefit; the synthesis, drafting, and ideation themes mirror the most-used discovery activities; the hallucination, context, and governance themes reproduce the top three quantitative barriers almost exactly; the human-validation theme explains why trust is the lowest-rated attitudinal item; and the strategic-shift theme corroborates both the augmentation score and the skills respondents now consider most important. The convergence of independently analysed qualitative and quantitative strands materially strengthens confidence in the thesis's conclusions.

Table 5.30: Theme-to-quantitative-evidence linkage.

Qualitative theme	Converging quantitative finding
Time saving and acceleration	Efficiency mean 4.09; "time saved on manual tasks" top benefit (60)
Synthesis / drafting / ideation	Top activities: PRDs (59), ideation (57), summarising notes (54)
Lack of organisational context	Barrier: "GenAI lacks context for our product" (27)

Qualitative theme	Converging quantitative finding
Hallucination and inaccuracy	Top barrier: inaccurate/hallucinated outputs (53)
Governance, privacy, access limits	Barriers: privacy/security (44), company policy limits (28)
Human validation and discernment	Trust mean 3.43, lowest item; only ~48% agree/strongly agree
Strategic shift and orchestration	Augmentation mean 3.91; top skills: critical (24) & strategic (19) thinking

5.8.6 Rigour and defensibility

The thematic analysis is defensible on several grounds. The coding was preceded by an explicit, reproducible cleaning step that removed blank and non-informative responses while preserving brief-but-meaningful and negative-but-relevant ones, so themes were not contaminated by noise. It followed an established, citable method [32], proceeded inductively from respondents' own language, and consolidated codes into themes only where multiple respondents and coherent meaning supported them. Frequencies are reported transparently as approximate interpretive counts rather than precise measurements, and the analysis is anchored throughout to verifiable extracts from the dataset. Finally, the themes were validated against the quantitative results rather than treated in isolation, so that the qualitative and quantitative strands act as mutual checks. The principal limitation is that free-text responses vary in length and depth and were coded by a single researcher; the consistent convergence with the closed-question data nonetheless gives reasonable confidence that the themes reflect genuine patterns in practitioner experience rather than coder artefact.

5.9 Chapter Summary

Chapter 5 presents five concrete findings that go beyond confirming GenAI's general utility. First, the trust–usefulness gap is not a transitional learning-curve effect but a structural posture that persists even among the most experienced practitioners, suggesting that conditional, verify-before-relying trust is the stable operating mode in product discovery, not a phase to be outgrown. Second, the friction point in adoption is the scale-up organisation, not the large enterprise — a finding that inverts common assumptions about where governance and integration challenges are greatest. Third, regional variation is institutional rather than attitudinal: practitioners across all well-represented markets show comparable appetite and augmentation endorsement, but differ in the governance conditions surrounding use.

Fourth, the modal practitioner is already at the Integrated stage of adoption — using GenAI daily across a broad activity set with human review embedded — which places the real-world baseline materially ahead of what most maturity frameworks assume. Fifth, while hallucination is the most frequently named barrier, the qualitative accounts reveal that missing organisational context is the deeper and structurally harder challenge, one that model improvement alone cannot resolve. Taken together, these findings reframe GenAI in product discovery not as an emerging experiment but as an already-embedded, human-supervised workflow tool, whose maturation path runs through better context-injection and governance, not simply more usage.

Chapter 6: Results Discussion, Triangulation, and Conclusion

6.1 Introduction

The preceding chapters established a body of evidence: a literature review across product discovery, enterprise GenAI adoption, human–AI collaboration theory, and the evolving product-management role (Chapter 2); a case-study analysis of how four leading product companies have operationalised GenAI (Chapter 4); and a mixed-method survey of 82 practitioners reported quantitatively and qualitatively (Chapter 5). The purpose of this chapter is to move from results to meaning. It interprets the findings against the theoretical frameworks introduced earlier: the Technology Acceptance Model [4], the Unified Theory of Acceptance and Use of Technology [14], augmentation theory [5], and the AI Fluency Framework [3][20], and integrates the three evidence streams through structured triangulation. It then refines the GenAI Adoption Maturity Model first proposed in Chapter 3 into a five-stage framework grounded in the survey evidence, states the thesis’s contributions and limitations, and concludes. The discussion is organised primarily around the four research sub-questions (SRQ1–SRQ4) so that the argument maps transparently onto the study’s stated aims.

6.2 SRQ1: Patterns of GenAI Adoption in Product Discovery

SRQ1 asked how PMs are adopting GenAI in discovery work. The evidence points to adoption that is routine, broad in scope, but uneven in intensity. Use is concentrated at the high-frequency end: daily use is by far the largest category and, together with several-times-per-week use, accounts for the overwhelming majority of respondents. This is not experimental dabbling but embedded practice. Adoption is also broad in activity scope. Rather than clustering in a single use case, respondents apply GenAI across the discovery lifecycle, with the heaviest use in writing PRDs and specifications, ideation and brainstorming, summarising meeting notes, stakeholder communication, and synthesising user research. The dominant uses are therefore synthesis, drafting, and coordination (the language-intensive, repetitive-cognitive parts of discovery) rather than autonomous decision-making.

Adoption is nonetheless context-dependent rather than uniform, as the cross-tabulations make clear. Senior and core PM roles use the tools most consistently; larger and mid-size organisations embed them most systematically into structured workflows; and more experienced practitioners report the strongest perceived usefulness. The regional

comparison adds that this high-adoption pattern recurs across every well-represented market, with only modest variation in rhythm. The most accurate characterisation of SRQ1 is therefore that GenAI has become a normal part of the discovery toolkit for most of these practitioners, but that the depth and systematicity of adoption scale with discovery responsibility and organisational structure. Adoption is widespread; it is not homogeneous.

6.3 SRQ2: Perceived Value of GenAI

SRQ2 asked how practitioners perceive the value of GenAI. The clearest quantitative signal in the study is that usefulness substantially exceeds trust: efficiency (mean 4.09), output quality (4.07), and ease of integration (3.99) all sit well above trust (3.43). Practitioners value GenAI primarily as a productivity and synthesis instrument: the most-cited benefits are time saved on manual tasks, better-quality documentation, and faster synthesis of research data, and the qualitative data reinforce this with accounts of multi-hour tasks collapsing into minutes.

These results align tightly with the Technology Acceptance Model [4], in which adoption is driven by perceived usefulness and perceived ease of use. Both antecedents are strongly present here, which explains the high observed adoption. UTAUT [14] adds explanatory depth: ease of integration corresponds to effort expectancy, and the prominence of company policy, licensing, and security as barriers reflects UTAUT's facilitating-conditions construct, which the case studies showed leading firms deliberately engineering. The key interpretive point for SRQ2 is the gap between usefulness and trust. Within TAM, usefulness can be high while trust remains moderate because the technology is valued for accelerating and structuring work that a human still owns and checks. The value of GenAI, on this evidence, is the value of a fast, capable productivity and synthesis tool, not that of an autonomous decision-maker.

6.4 SRQ3: Barriers and the Nature of Trust

SRQ3 asked what limits GenAI use and why trust is restrained. The barriers identified by respondents are reliability- and governance-related rather than usability-related. The dominant barrier is inaccurate or hallucinated output, followed by data-privacy and security concerns, company-policy limits, and the tool's lack of product and organisational context. The qualitative data give these barriers their mechanism: respondents describe "confident but inaccurate" output, "generic excellence" that misses internal nuance, and constraints imposed by licensing, token limits, and regulated environments.

This explains why trust is conditional rather than absent. Trust in the survey is best understood not as a fixed disposition but as a situational and organisational construct: practitioners extend trust selectively, contingent on the task's stakes, the availability of context, and the surrounding governance regime. Because the leading cause of distrust is hallucination and missing context, human review is not a transitional inconvenience but a structural requirement of responsible use: the point at which a knowledgeable PM supplies the context the model lacks and catches the errors it cannot see. This maps directly onto the Discernment and Diligence competencies of the AI Fluency Framework [3][20]: effective adoption depends less on raw usage volume than on the practitioner's ability to judge when output is wrong and to use the tools accountably. The case-study finding that leading firms pair capability with explicit governance reinforces the same conclusion from the organisational side.

6.5 SRQ4: Evolution of the PM Role

SRQ4 asked how GenAI is changing the PM role. The augmentation item scored 3.91, with agreement dominating across every role and region, and the skills respondents now consider most important (critical thinking, strategic thinking, prompt engineering, and AI literacy) point toward higher-order judgement rather than away from it. The qualitative themes make the shift concrete: practitioners describe moving “from content creator to strategic editor”, “orchestrating systems and processes more than ever before”, and using GenAI as “a thought partner and accelerator rather than a decision-maker.”

The most defensible interpretation is augmentation in the precise sense of Daugherty and Wilson [5]: cognitive labour is being redistributed, with the model taking on synthesis, drafting, and first-pass structuring while the human concentrates on interpretation, prioritisation, contextual judgement, and accountability. This is consistent with Rai and Rai's account of human–AI division of labour in software product teams [22] and with the Ulloa et al. principle that accountability must not be delegated to non-human actors [40]. The role is being reshaped, not removed: GenAI changes what a substantial part of the PMs day consists of: less manual production, more validation, framing, and orchestration, not merely how quickly the same tasks are done. A minority strand of role intensification, in which rising expectations accompany rising speed, is an important qualification: augmentation can raise the cognitive ceiling of the role rather than simply lightening its load.

6.6 Refined GenAI Adoption Maturity Model

Chapter 2 hypothesised that teams evolve not by using more AI but by using it with increasing sophistication, criticality, and accountability, and Chapter 3 introduced an initial maturity model whose stage proportions (Figure 3) were necessarily assumed in the absence of primary data. The survey evidence allows that model to be refined into a five-stage framework for product discovery, defined in PM terms, mapped to observable survey signatures, and, importantly, populated with the observed distribution of practitioners rather than assumed proportions. The stages are ordered by increasing frequency, breadth of activity, calibrated trust, and governance readiness, as summarised in Table 6.1 and depicted in Figure 23, which preserves the visual structure of the Chapter 3 model while substituting its provisional percentages with the real adoption distribution from the survey.

Table 6.1 : Refined GenAI Adoption Maturity Model for product discovery.

Stage	Definition in PM terms	Survey signature
1. No Adoption	GenAI not used in discovery; reliance on manual methods.	Rarely/never (2); “I don’t use GenAI in discovery”; no benefits reported
2. Experimental	Occasional, curiosity-driven trials in isolated tasks.	Occasional/weekly use (4); narrow activity set; tool-trialling language
3. Selective	Deliberate use in specific high-value activities, with heavy checking.	Several-times-per-week (~24%); use in synthesis/drafting; neutral trust, strong validation
4. Integrated	Routine, broad use embedded across the discovery workflow.	Daily use, the modal group (~62%); 8–10 activities; high efficiency/quality; workable trust; human review embedded
5. Strategic	GenAI part of the operating model; agentic pipelines and governance.	Daily users describing orchestration/custom-agent pipelines (~6%); governance-aware; “strategic editor” framing

The distribution of the sample across these stages is itself a finding. On the survey evidence, roughly 2% of respondents sit at No Adoption and about 5% at Experimental, while around 24% are Selective (several-times-per-week use in specific activities). The modal practitioner, approximately 62% of the sample, is at the Integrated stage, using GenAI daily across a broad activity set with human review built in; and a smaller leading edge of about 6% has reached the Strategic stage, describing bespoke agent ecosystems that automate early discovery and free time for strategy. This is a markedly more advanced profile than the assumed Chapter 3 distribution, in which adoption was expected to concentrate at the Experimental and Selective stages. Critically, movement up the ladder is not simply a function of using the tools more often. The defining transition from Selective to Integrated and then to Strategic is the maturation of discernment and governance: the ability to supply context, catch errors, and embed responsible-use practices, exactly the qualities the AI Fluency Framework foregrounds. This reframing of maturity around criticality and accountability rather than usage volume is one of the thesis’s original contributions, because it converts a descriptive adoption curve into a diagnostic that teams can use to locate themselves and identify what capability, not merely what tooling, they need next.

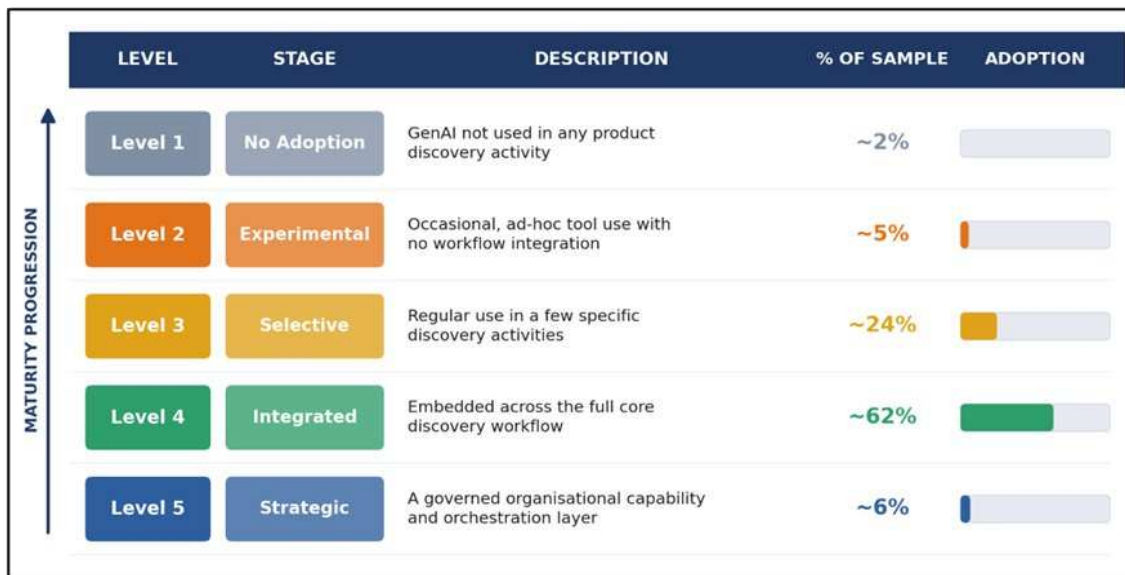


Figure 24. Refined GenAI Adoption Maturity Model for product discovery (real adoption distribution from the survey, N = 82).

6.7 Triangulation: Survey, Literature, and Case Studies

The strength of this study lies in convergence across three independent evidence streams. The literature predicted that GenAI would be strongest in synthesis-heavy, language-intensive knowledge work and that adoption would be governed by perceived usefulness, facilitating conditions, and augmentation rather than replacement. The case studies showed four leading product companies operationalising exactly this: building research-synthesis and documentation capabilities, framing GenAI as augmenting human judgement, openly acknowledging output limitations, and investing in enterprise governance as a deliberate enabler. The survey then confirmed that individual practitioners are already working in precisely these ways: using GenAI most heavily for synthesis, drafting, and coordination; valuing usefulness above trust; insisting on human review; and endorsing augmentation. Table 6.2 sets out the convergences and the few tensions.

Table 6.2 : Triangulation matrix across literature, case studies, and survey.

Dimension	Literature predicts	Case studies show	Survey confirms	Convergence / tension
Synthesis strength	Strongest in synthesis-heavy knowledge work [15]	Research-synthesis & documentation the headline capability	Top uses: PRDs, summarising, synthesis	Strong convergence
Augmentation framing	Augment, not replace [5][22]	Firms frame GenAI as augmenting judgement	Augmentation mean 3.91; broad agreement	Strong convergence
Output limitations	Hallucination/context risk [40]	Open acknowledgement; human-in-the-loop	Hallucination & context top barriers; trust 3.43	Strong convergence
Governance as enabler	Facilitating conditions [14]	Enterprise governance engineered deliberately	Policy/privacy among top barriers	Convergence; practitioner-side friction

Dimension	Literature predicts	Case studies show	Survey confirms	Convergence / tension
Discovery–delivery link	Emerging integration trend	Discovery–delivery tooling integration	Lighter signal at individual level	Partial / emerging

Where the streams diverge, the tensions are instructive rather than damaging. Governance appears in the case studies as an enabling investment but in the survey as a friction (policy, licensing, and privacy constraints) which reflects the different vantage points of the organisation that builds the guardrails and the practitioner who works within them; both are consistent with UTAUT’s facilitating-conditions construct. Likewise, the discovery-to-delivery integration that is visible at the platform level is only lightly echoed in individual practice, suggesting it is a structural trend still diffusing to the practitioner experience. Triangulation thus does more than repeat a single finding three times: by aligning the macro (literature), meso (organisational case studies), and micro (individual practice) levels of analysis, it shows that the same augmentation logic operates coherently across all three, which is a stronger warrant for the thesis’s conclusions than any single stream could provide.

6.8 Contributions, Limitations, and Future Research

This thesis makes both empirical and conceptual contributions. Empirically, it provides a current, mixed-method account of how PM actually use GenAI in discovery, triangulated against organisational practice and prior literature, in a domain where practitioner-level evidence remains thin. Conceptually, its principal contribution is the refined five-stage GenAI Adoption Maturity Model, which reframes maturity around discernment, calibrated trust, and governance readiness rather than usage volume, and which integrates TAM, UTAUT, augmentation theory, and AI Fluency into a single diagnostic for product teams. The convergent design is itself a contribution: independently analysed qualitative and quantitative strands corroborate one another, lending the conclusions a robustness uncommon in single-method studies of this size.

The study’s limitations should temper how its results are generalised. The sample of 82 is exploratory and self-selected, and is skewed toward senior practitioners, large enterprises, and high-frequency users, so it captures the experience of engaged adopters more than the full population of PM. The analysis is descriptive: cross-tabulations and theme frequencies

are interpreted as patterns, not tested for statistical significance, and the thematic frequencies are approximate interpretive counts coded by a single researcher. The findings therefore support analytical rather than statistical generalisation. Future research should address these limits directly: larger and more stratified samples that better represent junior practitioners, smaller firms, and under-represented regions; formal statistical testing of the role-, size-, experience-, and region-based differences observed here; longitudinal follow-up to track how practitioners and teams move through the maturity stages over time; and more explicit comparison of enterprise and startup contexts, since the present data hint that organisational structure shapes the systematicity rather than the existence of adoption. Multi-coder thematic analysis would further strengthen qualitative reliability.

6.9 Chapter Conclusion

This thesis set out to understand how GenAI is being adopted in product discovery, how practitioners value it, what limits it, and how it is reshaping the product-management role. The answer that emerges from the literature, the case studies, and the survey is consistent and clear. GenAI has become a routine instrument of product discovery for most of the practitioners studied, used daily and across the breadth of discovery work. Its value lies in accelerating repetitive cognitive work: it improves drafting, synthesis, and coordination, collapsing the distance between a question and a usable first draft. Yet that value is bounded. Practitioners trust GenAI markedly less than they find it useful, because hallucination and missing organisational context make its output plausible but not dependable, and because governance and privacy constraints shape where it can be applied. Trust is therefore conditional and situational, and human discernment remains indispensable: the PM supplies the context the model lacks, catches the errors it cannot see, and retains accountability for the decision.

The cumulative story is one of augmentation, not replacement. GenAI is shifting the centre of gravity of the role away from manual production and toward judgement, interpretation, validation, and orchestration, raising the value of critical and strategic thinking even as it automates the routine. The refined maturity model captures the trajectory this implies: teams advance not by using AI more, but by using it with greater discernment, calibrated trust, and governance maturity. On the evidence of this study, the most capable product organisations of the coming years will not be those that delegate discovery to machines, but those whose practitioners have learned to think well alongside them.

References

- [1] T. Torres, Continuous Discovery Habits: Discover Products That Create Customer Value and Business Value. Product Talk LLC, 2021.
- [2] McKinsey & Company (2025) The State of AI: Agents, Innovation, and Transformation. QuantumBlack, AI by McKinsey. November 2025. Available: <https://www.mckinsey.com/capabilities/quantumblack/our-insights/the-state-of-ai>
- [3] R. Dakan and J. Feller, 'Framework for AI Fluency (Version 1.1),' Ringling College of Art and Design and University College Cork, 2025. [Online]. Available: <https://ringling.libguides.com/ai/framework>
- [4] F. D. Davis, 'Perceived usefulness, perceived ease of use, and user acceptance of information technology,' MIS Quarterly, vol. 13, no. 3, pp. 319–340, Sep. 1989. doi: 10.2307/249008.
- [5] P. R. Daugherty and H. J. Wilson, Human + Machine: Reimagining Work in the Age of AI. Boston, MA: Harvard Business Review Press, 2018.
- [6] E. Ries, The Lean Startup: How Today's Entrepreneurs Use Continuous Innovation to Create Radically Successful Businesses. Crown Business, 2011.
- [7] Design Council, 'A Study of the Design Process: The Double Diamond,' Design Council UK, London, 2005.
- [8] Design Council, 'Framework for Innovation: Design Council's Evolved Double Diamond,' Design Council UK, London, 2019.
- [9] M. Cagan, Inspired: How to Create Tech Products Customers Love, 2nd ed. Hoboken, NJ: Wiley, 2017.
- [10] McKinsey Global Institute, 'The State of AI in 2023: Generative AI's Breakout Year,' McKinsey & Company, New York, NY, Aug. 2023.
- [11] McKinsey Global Institute, 'The Economic Potential of Generative AI: The Next Productivity Frontier,' McKinsey & Company, New York, NY, Jun. 2023.
- [12] J. Briggs and D. Kodnani, 'Generative AI Could Raise Global GDP by 7%,' Goldman Sachs Research, Apr. 2023. [Online]. Available: <https://www.goldmansachs.com/insights/articles/generative-ai-could-raise-global-gdp-by-7-percent>

- [13] S. Bahoo, M. Cucculelli, and D. Qamar, 'Artificial intelligence and corporate innovation: A review and research agenda,' *Technological Forecasting and Social Change*, vol. 188, p. 122264, 2023. doi: 10.1016/j.techfore.2022.122264
- [14] V. Venkatesh, M. G. Morris, G. B. Davis, and F. D. Davis, 'User acceptance of information technology: Toward a unified view,' *MIS Quarterly*, vol. 27, no. 3, pp. 425–478, Sep. 2003. doi: 10.2307/30036540.
- [15] B. Yun, D. Feng, A. S. Chen, A. Nikzad, and N. Salehi, 'Generative AI in Knowledge Work: Design Implications for Data Navigation and Decision-Making,' in *Proc. CHI '25*, ACM, New York, NY, USA, 2025, Article 634, pp. 1–19. doi: 10.1145/3706598.3713337
- [16] Deloitte AI Institute, 'State of Generative AI in the Enterprise 2024,' Deloitte, New York, NY, 2024.
- [17] D. Pradhan, B. Dash, P. Sharma, and S. Ullah, 'The Impact of Generative AI on Product Management in SMEs,' in *Generative AI in Business: Unlocking Value*, Springer, Cham, 2025, ch. 2. doi: 10.1007/979-8-8688-1603-1_2
- [18] National Institute of Standards and Technology (NIST), 'Artificial Intelligence Risk Management Framework (AI RMF 1.0),' NIST, Gaithersburg, MD, Jan. 2023. doi: 10.6028/NIST.AI.100-1
- [19] V. Venkatesh and F. D. Davis, 'A Theoretical Extension of the Technology Acceptance Model: Four Longitudinal Field Studies,' *Management Science*, vol. 46, no. 2, pp. 186–204, Feb. 2000. doi: 10.1287/mnsc.46.2.186.11926
- [20] K. Swanson, D. Bent, Z. Ludwig, R. Dakan, and J. Feller, 'Anthropic Education Report: The AI Fluency Index,' Anthropic, San Francisco, CA, Feb. 2026. [Online]. Available: <https://www.anthropic.com/research/AI-fluency-index>
- [21] P. R. Daugherty and H. J. Wilson, *Human + Machine: Reimagining Work in the Age of AI*, Updated and Expanded Ed. Boston, MA: Harvard Business Review Press, 2024.
- [22] A. Rai and S. Rai, 'Human-AI Collaboration: Developing Software Products with AI as a Partner,' in *Proc. QTanalytics Publication (Books)*, 2025, pp. 49–53. doi: 10.48001/978-81-980647-3-8-6
- [23] O. Springer and J. Miler, 'The Role of a Software Product Manager in Various Business Environments,' in *Proc. FedCSIS*, 2018, pp. 985–994. doi: 10.15439/2018F161

- [24] Scrum Alliance, 'State of Scrum Report 2023,' Scrum Alliance, Westminster, CO, 2023. [Online]. Available: <https://www.scrumalliance.org/state-of-scrum>
- [25] A. Yousif, 'Integrating Artificial Intelligence into Product Management,' M.S. thesis, KTH Royal Institute of Technology, Stockholm, Sweden, 2025. [Online]. Available: <https://kth.diva-portal.org>
- [26] G. Singh, 'All Work and No Flow: Experiences on Automating Product Owners' Workflow Using Generative AI,' M.S. thesis, KTH Royal Institute of Technology, Stockholm, Sweden, 2025. [Online]. Available: <https://www.diva-portal.org/smash/record.jsf?pid=diva2%3A2000318>
- [27] C. Jon-Onwubuoya, 'Working with AI the Naija Way: Exploring AI Adoption and Human-AI Collaboration Among Nigerian Product Managers,' M.S. thesis, Stockholm University, Stockholm, Sweden, 2024. [Online]. Available: <https://www.diva-portal.org/smash/get/diva2:1967757/FULLTEXT01.pdf>
- [28] F. Ghazi and S. Mariam, 'Conceptual Framework for Integrating Generative AI into the Product Management Lifecycle,' ResearchGate, 2025. [Online]. Available: <https://www.researchgate.net/publication/398513725>
- [29] J. W. Creswell and J. D. Creswell, *Research Design: Qualitative, Quantitative, and Mixed-Methods Approaches*, 5th ed. SAGE Publications, 2018.
- [30] A. Bryman, *Social Research Methods*, 5th ed. Oxford University Press, 2016.
- [31] R. K. Yin, *Case Study Research and Applications: Design and Methods*, 6th ed. SAGE Publications, 2018.
- [32] V. Braun and V. Clarke, 'Using thematic analysis in psychology,' *Qualitative Research in Psychology*, vol. 3, no. 2, pp. 77–101, 2006. doi: 10.1191/1478088706qp063oa
- [33] A. Field, *Discovering Statistics Using IBM SPSS Statistics*, 5th ed. SAGE Publications, 2018.
- [34] British Psychological Society, *Code of Ethics and Conduct*. Leicester, UK: BPS, 2021.
- [35] Figma Inc., 'Introducing Figma AI,' Figma Blog, 2024. [Online]. Available: <https://www.figma.com/blog/introducing-figma-ai/>
- [36] Figma Inc., 'Introducing AI to FigJam,' Figma Blog, 2023. [Online]. Available: <https://www.figma.com/blog/introducing-ai-to-figjam/>

- [37] Figma Inc., 'Figma 2024: We Shipped It, You Shaped It,' Figma Blog, 2024. [Online]. Available: <https://www.figma.com/blog/figma-2024-we-shipped-it-you-shaped-it/>
- [38] Figma Inc., 'Config 2025 Press Release,' Figma Blog, May 2025. [Online]. Available: <https://www.figma.com/blog/config-2025-press-release/>
- [39] Figma Inc., 'Figma's 2025 AI Report,' Figma, 2025. [Online]. Available: <https://www.figma.com/reports/ai-2025/>
- [40] M. Ulloa, J. L. Butler, S. Haniyur, C. Miller, B. Amos, A. Sarkar, and M.-A. Storey, 'Product Manager Practices for Delegating Work to Generative AI: Accountability Must Not Be Delegated to Non-Human Actors,' arXiv:2510.02504 [cs.SE], Oct. 2025. doi: 10.48550/arxiv.2510.02504.
- [41] M. Courtney, A. Barrett, S. Advait, and M.-A. Storey, 'Empowering Business Transformation: The Positive Impact and Ethical Considerations of Generative AI in Software Product Management,' arXiv:2306.04605 [cs.SE], Jun. 2023. doi: 10.48550/arXiv.2306.04605
- [42] Notion Inc., 'Notion AI for Work,' Notion Blog, May 2025. [Online]. Available: <https://www.notion.com/blog/notion-ai-for-work-unlocking-ai-driven-product-development>
- [43] Notion Inc., 'AI Product Roadmap,' Notion, Jan. 2026. [Online]. Available: <https://www.notion.com/use-case/project-management/ai-product-roadmap>
- [44] BestAIProjectHub, 'Notion AI for Product Managers: PRD Drafting and Research Workflows,' BestAIProjectHub.com, 2025. [Online]. Available: <https://bestaiprojecthub.com>
- [45] T. Gültekin Atlı, 'Unintended Consequences of Generative AI in Product Ideation: A Product Management Perspective,' M.S. thesis, Istanbul Technical University, Istanbul, Turkey, 2025.
- [46] Eesel AI, 'Notion Research Mode: Everything You Need to Know,' Eesel Blog, 2025. [Online]. Available: <https://www.eesel.ai/blog/notion-research-mode>
- [47] Miro Inc., 'Use AI to Innovate at Each Stage of Product Development,' Miro Blog, Jun. 2024. [Online]. Available: <https://miro.com/blog/product-innovation-supercharge-with-ai/>
- [48] Miro Inc., 'What We Launched in October 2025,' Miro Blog, Oct. 2025. [Online]. Available: <https://miro.com/blog/whats-new-october-2025/>

- [49] Miro Inc., 'The Miro Recap: Top 25 Updates of 2025,' Miro Blog, Dec. 2025. [Online]. Available: <https://miro.com/blog/2025-recap/>
- [50] Miro Inc., 'A New Way for Product Teams to Discover and Build the Right Things,' Miro Blog, Jan. 2026. [Online]. Available: <https://miro.com/blog/discover-and-build-the-right-things/>
- [51] Miro Inc., 'What We Launched in August 2025,' Miro Blog, Aug. 2025. [Online]. Available: <https://miro.com/blog/whats-new-august-2025/>
- [52] Atlassian Inc., 'What's New in Confluence,' Atlassian, 2024–2026. [Online]. Available: <https://www.atlassian.com/software/confluence/features/whats-new>
- [53] Atlassian Inc., 'Your Jira AI 2024 Recap and What's Coming Soon,' Atlassian Community, Jan. 2025. [Online]. Available: <https://community.atlassian.com/forums/Jira-articles/>
- [54] A. Aggarwal, 'Redefining Product Management with Generative AI: Scope, Skills, and Strategic Impact,' IEEE ICIDCA, Oct. 2025. doi: 10.1109/icidca66325.2025.11280336.
- [55] Atlassian Inc., 'Explore AI Features in Atlassian Products,' Atlassian Support, 2025. [Online]. Available: <https://support.atlassian.com>
- [56] Atlassian Inc., 'Atlassian Cloud Updates November 2025,' Atlassian Blog, Nov. 2025. [Online]. Available: <https://confluence.atlassian.com/cloud/blog/2025/11/>
- [57] Eesel AI, 'Atlassian Intelligence: AI in Jira and Confluence,' Eesel Blog, 2026. [Online]. Available: <https://www.eesel.ai/blog/atlassian-intelligence-ai-in-jira>