

**AI ADOPTION MOTIVATION AMONG IT SECTOR
EMPLOYEES:**

A SELF-DETERMINATION THEORY PERSPECTIVE

**MASTER OF ARTS IN HUMAN RESOURCE
MANAGEMENT**

2024-2026

AI ADOPTION MOTIVATION AMONG IT SECTOR

EMPLOYEES:

A SELF-DETERMINATION THEORY

PERSPECTIVE

MASTER OF ARTS IN HUMAN RESOURCE

MANAGEMENT

Program Code: 548

Course Code: 547

ABSTRACT

This study examines the motivational orientation underlying artificial intelligence (AI) adoption among IT sector employees, using Self-Determination Theory (SDT) as its guiding framework. Employees were classified as Intrinsic or Extrinsic adopters on the basis of a composite baseline motivation score derived from a four-item Likert scale, using a median-split method applied to a cleaned sample of 65 respondents (36 Intrinsic, 29 Extrinsic). A Sequential Explanatory Mixed-Method design was adopted, combining a structured questionnaire with semi-structured interviews analysed thematically.

Three hypotheses were tested. Motivational orientation was found to be significantly associated with voluntary AI adoption behaviour ($\chi^2 = 26.08$, $p < .001$). Training-induced curiosity scores differed significantly between Intrinsic and Extrinsic adopters ($t = 3.72$, $p < .001$). Extrinsic adopters' curiosity scores were directionally higher than their own baseline motivation scores, though this difference approached, but did not reach, conventional statistical significance ($t = 1.89$, $p = .070$). This directional shift is discussed as a trend worth further investigation rather than a confirmed effect. Qualitative interviews with L&D leaders identified relevance, manager support, and trust in AI-driven recommendations as the organisational conditions that help explain this pattern. The study concludes with practical recommendations for L&D functions designing AI training interventions, and notes the value of a larger, longitudinal follow-up study.

Keywords: Self-Determination Theory, AI adoption, intrinsic motivation, extrinsic motivation, training-induced curiosity, IT sector, India.

CONTENT

SL NO	TITLE	PAGE NO
1	INTRODUCTION	6-15
2	REVIEW OF LITERATURE	16-24
3	RESEARCH METHODOLOGY	25-31
4	ANALYSIS AND INTERPRETATION	32-56
5	DISCUSSION	57-62
6	CONCLUSION AND RECOMMENDATIONS	63-70
7	REFERENCES	71-73
8	APPENDICES	74-80

LIST OF TABLES

SL NO	TITLE	PAGE NO
4.1.1	Motivational Orientation by Years of IT Sector Experience	35-36
4.2.1	Descriptive Statistics for Baseline Motivation and Training-Induced Curiosity, by Group	37-38
4.3.1	Chi-Square Test of Association: Motivational Orientation and Voluntary AI Adoption (H1)	39-40
4.3.2	Paired-Samples t-Test: Baseline Motivation vs Training-Induced Curiosity (H2)	41-42
4.3.3	Independent-Samples t-Test: Curiosity Scores, Intrinsic vs Extrinsic (H3)	43-44
4.4.1	Comparison of Training Experience Items, Intrinsic vs Extrinsic Adopters	46-48
4.5.1	Chi-Square Test of Association: Motivational Orientation and Role Level	50

LIST OF FIGURES

SL NO	TITLE	PAGE NO
1.8.1	Conceptual Framework of the Study	23
4.1.1	Years of IT Sector Experience of Respondents	32-33
4.1.2	Role Level Distribution of Respondents	33-34
4.1.3	Organisational Introduction of AI-Based Training Tools	34-35
4.2.1	Motivational Orientation Distribution of Respondents	36-37
4.2.2	Motivational Orientation by Years of IT Sector Experience	36
4.2.3	Baseline Motivation vs Training-Induced Curiosity, by Group	38
4.5.1	Scatterplot of Baseline Motivation and Training-Induced Curiosity	49

CHAPTER 1

INTRODUCTION

Artificial intelligence has moved, in the space of only a few years, from being a specialised technical tool confined to data science and engineering functions to something that ordinary employees across departments are now routinely expected to use in their daily work. This shift has happened quickly, and organisations in the IT sector have responded to it in markedly different ways: some have rolled out mandatory, structured training programmes, while others have left adoption largely to individual initiative, offering access to tools without a formal push to use them. What remains comparatively less understood is why employees choose to engage with AI tools in the first place, and whether that underlying “why” shapes the way they go on to use these tools over time.

This study approaches that question through the lens of Self-Determination Theory (SDT), a motivational framework that distinguishes behaviour driven by internal interest and satisfaction from behaviour driven by external pressure, reward, or compliance. Applied to AI adoption, this distinction has practical, not just academic, value. During early conversations with an industry mentor during this researcher's internship at an IT services company, a recurring observation surfaced repeatedly: some employees begin learning AI tools on their own initiative, out of curiosity or a sense that doing so will keep them professionally relevant, while others only take up these tools because their team lead, or the wider organisation, is pushing them to. These two groups were informally referred to, in those early conversations, as intrinsic adopters and extrinsic adopters, and this distinction forms the conceptual backbone of the present study.

Importantly, the research does not treat this distinction as a fixed, binary split carved in stone at the moment of first exposure to a tool. Instead, it asks a more dynamic question: whether the extrinsic group, once exposed to structured organisational training, shows any measurable movement towards a more intrinsic form of motivation. If organisational training can shift how people relate to AI, and not just whether they use it, this carries real and actionable implications for how Learning and Development (L&D) functions design their programmes going forward.

1.1 Background of the Study

India's IT sector has seen AI-based tools integrated into daily workflows at a pace that few other workplace technologies have matched in recent memory from code-completion assistants and customer-support automation to internal knowledge chatbots that field routine employee queries. Adoption numbers alone, however, do not tell us a great deal about the underlying quality of that adoption, because an employee who uses a tool because they were told to and an employee who uses the same tool because they find it genuinely useful are likely to behave very differently over the long run. This difference plays out in how deeply employees engage with a tool, how willing they are to experiment with its more advanced features, and, critically, whether their engagement holds up once the initial organisational mandate or the novelty of the technology wears off.

India's IT workforce provides a particularly apt setting in which to examine this question. It is young, rapidly growing, and increasingly exposed to generative AI tools as a matter of routine organisational practice rather than as a novelty confined to a handful of specialist roles. Understanding the motivational undercurrents of this exposure has both theoretical and practical value for a sector whose continued economic significance is closely tied to the competitiveness of its technology workforce.

1.2 Statement of the Problem

Organisations investing in AI training programmes tend to measure the success of those programmes through participation numbers or short-term usage metrics how many employees logged in, completed a module, or used a tool at least once in a given period. These measures, while easy to collect, do not capture the underlying motivational orientation of the employee undergoing training. As a consequence, training interventions may be designed without accounting for the possibility that different motivational groups require fundamentally different kinds of support, feedback, and pacing in order for the training to translate into durable, self-sustaining engagement rather than a box-ticking exercise that lapses the moment the mandate is lifted.

To date, comparatively little empirical work has asked whether the design of a training programme itself, as distinct from the AI tool being trained on, can shift an employee's motivational orientation. This represents a meaningful gap for both theory and practice, and it is the gap this study sets out to address.

1.3 Significance of the Study

This study speaks directly to several audiences. For HR managers and L&D professionals designing or revising AI training programmes, it offers a way of thinking about training success that goes beyond simple participation metrics, towards a more nuanced understanding of motivational quality. For researchers, it contributes to the growing body of literature applying Self-Determination Theory outside its traditional educational and clinical settings, extending it into a workplace technology-adoption context that is still relatively under-theorised. For employees themselves, it offers a vocabulary for understanding their own relationship with AI tools whether that relationship began out of genuine interest or organisational pressure, and whether it is capable of changing over time.

More broadly, in an era where organisations across sectors are racing to roll out AI training at scale, often with limited attention to the psychological experience of the employees on the receiving end, this study's findings offer a timely reminder that motivation, and not just competence, is something training design can either support or inadvertently undermine.

1.4 Objectives of the Study

Primary Objective

To examine the motivational orientation behind AI adoption among IT sector employees, and how it relates to their adoption behaviour and response to structured training.

Specific Objectives

- To identify the motivational orientation (intrinsic or extrinsic) of IT sector employees who have undergone structured AI training, and examine its association with voluntary AI adoption behaviour.
- To examine training-induced curiosity following structured AI training: whether it shifts among extrinsically motivated adopters relative to their own baseline, and how it differs between intrinsically and extrinsically motivated adopters.
- To explore, through qualitative themes, the underlying reasons and experiences that shape how employees relate to structured AI training and the tools it introduces.
- To draw practical implications for L&D functions designing AI training interventions.

All four objectives share a common anchor: the experience of IT sector employees who have undergone structured AI training. Objectives 1 and 3 examine the motivational and experiential starting point of that training (orientation and lived experience, respectively), while Objectives 2 and 4 examine its measurable outcome and its practical application for L&D. None of the four objectives is intended to be studied independently of the training context set out in section 1.2.

1.5 Hypotheses

- H1: Motivational orientation is significantly associated with voluntary AI adoption behaviour.
- H2: Extrinsically motivated adopters report higher training-induced curiosity scores than their own baseline motivation score.
- H3: Training-induced curiosity scores differ significantly between intrinsic and extrinsic adopters.

Correspondence with objectives: H1 corresponds to Objective 1 (orientation and voluntary adoption). H2 and H3 both correspond to Objective 2 (training-induced curiosity: the shift among extrinsic adopters, and the difference between intrinsic and extrinsic adopters). Objectives 3 and 4 are addressed through the qualitative and practical-implications components of the study rather than through a specific hypothesis.

1.6 Definition of Concepts

This study is grounded in Self-Determination Theory (Deci & Ryan, 2000), which distinguishes intrinsic motivation from extrinsic motivation and further proposes that extrinsic motivation exists along a continuum of regulation styles rather than as a single fixed category. Four constructs central to this study are defined below in both theoretical and operational terms, following the convention of presenting a construct's conceptual meaning before describing precisely how it was measured in this dataset.

AI Adoption

Theoretical definition: the extent to which an individual incorporates artificial intelligence tools into their regular work practices, ranging from occasional, mandated use to sustained, self-directed engagement.

Operational definition: in this study, AI adoption is measured through respondents' self-reported approach to learning new skills in their field, collapsed into a voluntary category (learning on their own initiative, or a mix of self-initiative and organisational requirement) versus a non-voluntary category (learning only because required to, or not having engaged with AI training at all).

Intrinsic Motivation

Theoretical definition: engagement in an activity for its own sake, driven by inherent interest, satisfaction, or enjoyment, as defined within Self-Determination Theory (Deci & Ryan, 2000).

Operational definition: in this study, intrinsic motivation is operationalised as a composite score across four Likert-scale items capturing a personal sense of urgency to keep learning, willingness to continue learning without an organisational requirement, learning out of genuine want rather than obligation, and viewing AI skill development as important for long-term career security. Respondents scoring at or above the sample median on this composite are classified as Intrinsic adopters.

Extrinsic Motivation

Theoretical definition: engagement in an activity because of an external outcome, such as compliance with a requirement, reward, or avoidance of a negative consequence, rather than inherent interest (Deci & Ryan, 2000).

Operational definition: in this study, respondents scoring below the sample median on the composite motivation score described above are classified as Extrinsic adopters.

Training-Induced Curiosity

Theoretical definition: a state of heightened interest and self-directed motivation to continue learning that emerges following exposure to a structured training intervention. “Training-induced curiosity” is not a pre-existing named construct within Self-Determination Theory; it is this study’s own term for a specific, situational form of interest. It is grounded in two established lines of work: the Interest/Enjoyment subscale of the Intrinsic Motivation Inventory (Ryan, 1982), used elsewhere in this study’s measurement approach (see section 3.6), and the broader concept of situational interest in educational psychology, a state of engagement triggered by a specific event or environment rather than a stable individual trait.

Framed this way, the construct is a study-specific application of an established idea, rather than an invented one.

Operational definition: in this study, training-induced curiosity is measured using a single Likert-scale item “AI training has made me more curious and motivated to keep learning on my own.”

1.7 Theoretical Framework: Self-Determination Theory

Self-Determination Theory occupies a central place in the theoretical architecture of this study, and this section sets out its logic in some detail, since it underpins every hypothesis tested in this report.

The theory is introduced here, in the Introduction, because it underpins every hypothesis that follows; Chapter 2 then reviews the empirical AI-adoption literature that this theoretical lens is used to interpret.

1. Deci & Ryan (2000) This foundational paper distinguishes between motivation that originates from within the individual (intrinsic) and motivation driven by external factors such as rewards, deadlines, or supervisory pressure (extrinsic). A central contribution of the theory is the idea that extrinsic motivation is not a single category but exists on a continuum, ranging from external regulation, where behaviour is driven purely by compliance, through to integrated regulation, where an externally introduced behaviour has been absorbed into the person's own sense of identity and values. The theory further identifies autonomy, competence, and relatedness as the three basic psychological needs whose satisfaction supports movement toward more autonomous forms of motivation.

This continuum is directly relevant to the present study because it suggests that extrinsically motivated adopters are not a fixed group; they can, in principle, move along this continuum given the right situational conditions. It is this theoretical possibility of movement, rather than a static intrinsic/extrinsic split, that the present study sets out to test empirically within a corporate AI-training context.

SDT was selected over adoption-prediction frameworks such as the Technology Acceptance Model (Davis, 1989) and UTAUT (Venkatesh et al., 2003) because those models explain whether and how much a technology is adopted based on beliefs about its usefulness

and ease of use, whereas this study is concerned with the quality of the motivation behind adoption, genuine interest versus external pressure, a distinction SDT is specifically designed to capture.

SDT was applied throughout the study rather than as a citational frame alone: the Intrinsic/Extrinsic classification operationalises the theory's core motivational-quality distinction; Hypothesis 1 tests its behavioural prediction that autonomous motivation predicts voluntary behaviour; Hypothesis 2 tests its internalisation process, whereby extrinsic motivation can move toward more autonomous forms under supportive conditions; Hypothesis 3 tests the theory's prediction that intrinsic motivation sustains engagement more durably than extrinsic motivation; and the emergence of Competence as an independent qualitative theme (section 4.7) provides empirical, participant-derived support for one of the theory's three basic needs, rather than a purely top-down theoretical assumption.

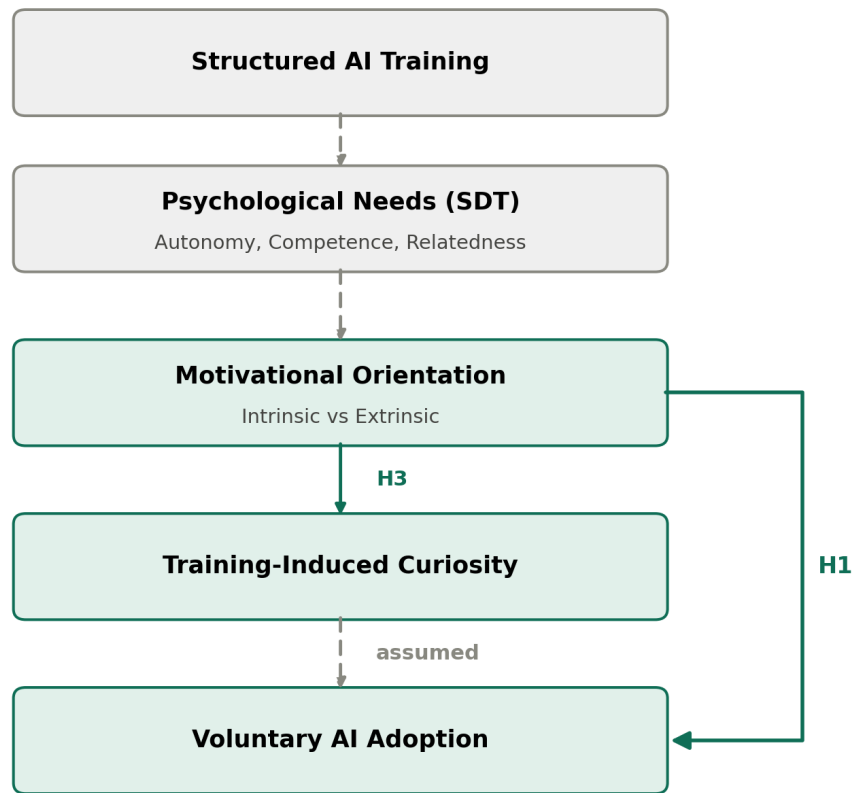
This application of SDT maps directly onto the study's four objectives: it grounds the intrinsic/extrinsic classification and its link to voluntary adoption (Objective 1), underpins the curiosity-related hypotheses tested under Objective 2, is empirically supported by the qualitative findings addressing Objective 3, and translates into the practice-facing recommendations of Objective 4.

1.8 Conceptual Framework

The relationships this study investigates can be summarised in a simple sequential framework, illustrated below.

The framework treats the three SDT psychological needs as the mechanism connecting training design to motivational orientation, and treats training-induced curiosity as the mechanism connecting motivational orientation to durable, voluntary adoption behaviour, rather than treating adoption as a direct, unmediated product of training alone. The diagram below adds a direct arrow from Motivational Orientation to Voluntary AI Adoption, labelled H1, since this is the study's strongest tested finding ($\chi^2 = 26.08$, $p < .001$). The arrows from Structured AI Training through the psychological needs to Motivational Orientation remain dashed to indicate that they are theoretical scaffolding drawn from SDT rather than relationships measured or tested in this study. The Orientation-to-Curiosity link is solid and

labelled H3, the relationship this study did test; the Curiosity-to-Adoption link remains dashed as an assumed mediating pathway that was not directly tested.



Solid = statistically tested pathway. Dashed = theoretical, untested link.

Figure 1.8.1: Conceptual Framework of the Study, with Tested and Untested Pathways

1.9 Scope and Limitations

The study is limited to a sample of Indian IT sector employees and relies on self-reported data, which carries the usual limitations associated with social desirability bias. The sample size, while adequate for a mixed-method exploratory study of this kind, is not large enough to generalise findings across the entire IT sector, and the cross-sectional design means that claims about change over time should be read as trends rather than confirmed longitudinal effects.

1.10 Chapter Summary

This chapter has introduced the study's guiding problem, motivational orientation as a lens on AI adoption, and set out the objectives, hypotheses, key definitions, theoretical framework, and scope within which the remaining chapters operate. Chapter 2 now turns to the existing literature this study builds on and the specific gap it addresses.

CHAPTER 2

REVIEW OF LITERATURE

2.1 Introduction

This chapter reviews the body of literature underpinning the present study. Research on technology adoption in organisations has traditionally centred on models such as the Technology Acceptance Model and UTAUT, which predict whether a technology is adopted but say little about the quality of the motivation behind that adoption. Self-Determination Theory, introduced in Chapter 1 as this study's guiding theoretical lens, offers a way of addressing that gap by distinguishing adoption driven by genuine interest from adoption driven by external pressure.

The chapter proceeds from the general to the specific: it first situates AI adoption within the broader organisational technology-adoption literature (2.2), then narrows to what is known about AI training and learning behaviour (2.3) and the intrinsic/extrinsic distinction within AI learning specifically (2.4), before turning to the recent empirical studies published since the generative-AI wave of 2023 that most directly inform this study (2.5). It closes by identifying precisely where this study sits relative to that existing work (2.6).

2.2 AI Adoption in Organizations

Traditional technology adoption models, most notably the Technology Acceptance Model (TAM) and its various extensions, have focused heavily on perceived usefulness and perceived ease of use as predictors of whether adoption happens at all. These models are useful for predicting the binary fact of adoption but say comparatively little about the motivational quality that sits behind that adoption that is, whether an employee adopts a tool because they were compelled to or because they wanted to. Existing research on workplace AI adoption more broadly tends to focus instead on organisational readiness, infrastructure maturity, and employee anxiety or resistance around job displacement.

Less attention has been paid to the specific internal motivational split this study is concerned with, and in particular to the idea that training itself, and not merely the tool being trained on, could function as a lever for shifting motivation. This is a gap the present study tries to address by deliberately bringing motivational theory into a technology-adoption context that is specific to AI.

2.3 AI Training and Learning Behaviour

There is prior evidence, both within educational psychology and in corporate settings, suggesting that well-designed training interventions particularly those that actively support autonomy, competence, and relatedness can move learners from externally regulated behaviour towards more autonomous forms of motivation over the course of an intervention. In one of the few workplace-based tests of this logic, Hardré and Reeve (2009) trained managers at a Fortune 500 company to adopt a more autonomy-supportive style, and found that five weeks later the employees they supervised showed significantly greater autonomous motivation and workplace engagement than employees supervised by untrained managers direct evidence that training-based movement along the motivational continuum is achievable in a real, for-profit workplace, and not only in a classroom. This study borrows that underlying logic and applies it specifically to AI-skills training, asking whether the same basic principle holds when the behaviour being trained is the adoption of a new AI tool rather than a general management style.

2.4 Intrinsic vs Extrinsic Motivation in AI Learning

Within the specific context of AI learning, the intrinsic/extrinsic distinction plays out as a difference between employees who treat AI skill-building as a form of self-directed professional development, undertaken because it feels personally worthwhile, and those who treat it as a compliance task tied narrowly to organisational mandates. This distinction is the organising variable of the present study, and the section that follows reviews the recent empirical work most directly relevant to it.

2.5 Recent Literature

Given how quickly generative AI tools have moved into everyday workplace use since 2023, this study is grounded in literature that reflects that specific and recent wave of technological change, rather than relying solely on earlier, pre-generative-AI technology-adoption models that predate the tools under discussion here.

2. Schmitt, Gajos, & Mokryn (2024) studying software engineers working with generative AI tools through the combined lenses of occupational identity and Self-Determination Theory, found that engineers' sense-making of the technology was contingent on their domain expertise: some felt the tools enhanced their sense of competence and autonomy, while others experienced a diminished sense of ownership over their work and a reduced

sense of challenge. This finding lines up closely with the SDT framing used throughout the present study, since the same underlying tool produced opposite motivational effects depending on how the individual related to it.

3. **Prasad & De (2024)** surveying 500 information technology employees in India, modelled generative AI adoption by combining the Technology Acceptance Model, the Technology Readiness Index, and Stimulus-Organism-Response theory, and found that trust in the reliability and usefulness of the tools significantly shaped organisational commitment and, in turn, employee engagement and performance. This supports an idea developed further in this study's own findings that trust in an AI system's accuracy and visible relevance likely shapes whether an employee's relationship with it becomes self-sustaining or remains dependent on external pressure.
4. **Liu, Sheng, & Liu (2025)** drawing on Conservation of Resources theory in a two-wave study of Chinese enterprises, found that generative AI adoption positively predicted employees' job-crafting behaviour (seeking resources, seeking challenges, and optimising demands), which in turn strengthened career commitment and shaped downstream behavioural outcomes. This adds a longer-term behavioural dimension for future research building on the present study: motivational shift is one plausible outcome of AI training, but job crafting may be a second, related outcome worth tracking separately in subsequent work.
5. **Bergdahl, Latikka, Celuch, Savolainen, Mantere, Savela, & Oksanen (2023)** tested Self-Determination Theory's basic psychological needs directly against attitudes toward AI, using a six-country cross-sectional sample and a two-wave longitudinal sample from Finland, and found that competence and relatedness were consistently associated with more positive AI attitudes across all six countries, while within-person increases in autonomy and relatedness over time predicted increased AI positivity. This is the most direct empirical precedent for treating SDT's basic needs, rather than adoption alone, as a lens on how people relate to AI specifically, and its longitudinal design is the kind this study's own recommendations (section 6.3) call for at a smaller, organisational scale.
6. **Lai, Cheung, & Chan (2023)** examined ChatGPT adoption among 473 undergraduates in Hong Kong by extending the Technology Acceptance Model with intrinsic and extrinsic motivation constructs, and found intrinsic motivation to be the strongest predictor of ChatGPT use intention, ahead of the model's traditional usefulness and ease-of-use

variables. This lends direct support to the present study's starting premise: that motivational quality, and not perceived usefulness alone, is the more powerful lens through which to understand engagement with a specific, widely used generative AI tool.

- 7. Cajander, Bergqvist, Clear, Daniels, Humble, Larusdottir, & Normark (2026)** conducted a qualitative study of IT professionals grounded explicitly in Self-Determination Theory, examining how AI integration shapes work engagement, and found that AI both augmented and complicated professionals' working lives providing opportunities for skill growth and more meaningful work for some, while others reported increased uncertainty and heightened demands. This qualitative texture closely echoes the mixed, group-dependent experience captured in the interview themes reported in section 4.7, and reinforces that IT professionals' motivational relationship with AI is not uniform even within a single, IT-specific workforce.

Together, these six recent studies establish that intrinsic-versus-extrinsic-style splits in orientation towards AI are real and observable across employee, student, and general-population samples in different national contexts; that AI exposure can amplify existing motivation and shape job-crafting behaviour; that SDT's basic psychological needs predict AI-specific attitudes directly, including over time; and that IT professionals specifically experience AI integration through a mixed, motivationally uneven lens rather than a uniform one. None of them, however, directly tests whether structured organisational training is associated with a shift in an employee's curiosity relative to their own baseline orientation the specific pathway this study investigates.

2.6 Research Gap

Existing studies examine AI adoption largely through the lenses of trust, productivity, and organisational readiness. Few examine the specific pathway from training, through motivation, to AI adoption using Self-Determination Theory as the connecting framework. The recent generative-AI literature reviewed above has established that intrinsic-versus-extrinsic-style splits in employee orientation exist, and that AI exposure can amplify existing motivation and job-crafting behaviour, but it has not directly tested whether structured training is associated with a shift in an employee's training-induced curiosity relative to their baseline orientation. This study positions itself squarely in that gap.

The sections above were deliberately organised around this study's own objectives rather than as a generic survey of the field: section 2.2 grounds Objective 1 (motivational orientation and voluntary adoption) in the organisational AI-adoption literature; section 2.3 grounds Objectives 2 and 4 (training-induced curiosity and its practical implications) in the training-intervention literature; section 2.4 grounds the Intrinsic/Extrinsic distinction used throughout Objectives 1 through 3; and section 2.5 grounds all four objectives in the most recent empirical work on generative AI specifically, since this is where the identified gap sits most precisely.

CHAPTER 3

RESEARCH METHODOLOGY

This chapter sets out how the study was designed and carried out: the research design adopted and the reasoning behind it, the population and sampling approach, the sources of data drawn upon, how each variable was measured and its reliability established, the tools and procedures used to collect data, the statistical techniques applied, the ethical safeguards observed, and the methodological limitations that qualify how the findings in Chapter 4 should be read.

3.1 Title of the Study

The full title of this study, as registered with the Department, is “AI Adoption Motivation Among IT Sector Employees: A Self-Determination Theory Perspective.” It is restated here to open the Research Methodology chapter, and reflects the study’s dual focus: the motivational orientation, intrinsic or extrinsic, that IT sector employees bring to AI adoption, and the Self-Determination Theory lens through which that orientation is examined and tested in the sections that follow.

3.2 Research Design

The study follows a Sequential Explanatory Mixed-Method design, as outlined by Creswell (2014). This means that the quantitative phase was conducted first, using a structured questionnaire to measure motivational orientation and training-induced curiosity, and was followed by a qualitative phase using semi-structured interviews intended to explain and add interpretive depth to the patterns identified in the quantitative data. This sequencing is deliberate: the quantitative phase establishes what pattern exists in the data, while the qualitative phase that follows is used to explore why that pattern might exist, rather than the two phases running independently of one another.

This design was chosen over a convergent parallel approach because the study's objectives require the quantitative pattern (motivational orientation, adoption, curiosity) to be established first, before the qualitative phase can meaningfully explain why that pattern exists; running both phases simultaneously would not allow the interview questions to be shaped by what the quantitative results actually showed. A purely quantitative design was

ruled out because Objective 5 explicitly calls for exploring employees' underlying reasons and experiences, which numbers alone cannot capture.

3.3 Population

The population for this study consists of Indian IT sector employees who have undergone some form of organisational AI-related training. This population was chosen specifically because the research question concerns the effect of training exposure on motivational orientation, which necessarily requires respondents who have already been through at least one such training experience.

3.4 Sampling

Purposive sampling was used to select IT sector employees who had gone through some form of organisational AI training. The survey instrument collected 68 responses in total. After removing entries with incomplete data across the demographic and motivation-item fields, a cleaned sample of 65 respondents was retained for analysis. This sample of 65 forms the basis for all quantitative results reported in Chapter 4.

3.5 Sources of Data: Primary and Secondary

This study draws on two categories of data. Primary data was collected directly by the researcher for the specific purposes of this study, and consists of the structured questionnaire responses gathered from the cleaned sample of 65 respondents, together with the semi-structured interview transcripts gathered from the purposively selected subset of Intrinsic and Extrinsic adopters described in the Data Collection Tools and Data Collection Procedure sections below. This primary data forms the empirical basis for all quantitative and qualitative findings reported in Chapter 4.

Secondary data consists of previously published material used to build the study's theoretical framework and literature review, principally Deci and Ryan's (2000) foundational work on Self-Determination Theory, Creswell's (2014) methodological guidance on mixed-method research design, and the recent empirical studies on generative AI adoption reviewed in Chapter 2 (Schmitt et al., 2024; Prasad & De, 2024; Liu et al., 2025). Secondary data was not used to generate any of the study's own findings, but to situate those findings within the existing body of theory and evidence.

3.6 Measurement of Variables

Motivational orientation was measured using four Likert-scale items (1 to 5) capturing the core intrinsic-motivation constructs from Self-Determination Theory: a personal sense of urgency to keep learning in order to stay relevant, willingness to continue learning AI skills even without an organisational requirement, engaging with AI training out of genuine want rather than obligation, and viewing AI skill development as essential to long-term career security. A composite motivation score was calculated for each respondent as the mean of these four items.

Self-Determination Theory itself is a motivational framework rather than a measurement instrument, so it does not prescribe a fixed Likert scale. However, validated Likert-based instruments built on SDT's logic do exist, including the Work Extrinsic and Intrinsic Motivation Scale (WEIMS), the Multidimensional Work Motivation Scale (MWMS), and the Interest/Enjoyment subscale of the Intrinsic Motivation Inventory (IMI, Ryan, 1982). The four-item composite used in this study was adapted from the logic of the IMI Interest/Enjoyment subscale and the MWMS, condensed specifically to align with AI-training language rather than generic work motivation, in the interest of a short, workplace-appropriate survey instrument.

Respondents were then classified into two groups using a median split on this composite score (median = 3.75). A median split means dividing respondents into two groups based on whether their score falls above or below the middle (median) value of the sample itself, rather than against a fixed benchmark such as the scale's midpoint of 3: those scoring at or above the median were classified as Intrinsic adopters, and those scoring below the median were classified as Extrinsic adopters. This produced 36 Intrinsic and 29 Extrinsic respondents. A median-split, composite-score approach was used in preference to a single self-selected category question because it draws on information from all four items rather than one, and because it avoids leaving a residual “mixed” category of respondents who cannot be meaningfully assigned either way; every respondent in the cleaned sample of 65 is classified as either Intrinsic or Extrinsic under this method.

Training-induced curiosity, the study's key outcome variable, was measured using a single item “AI training has made me more curious and motivated to keep learning on my own” scored on the same 1 to 5 scale. A single-item measure was used for brevity within an already multi-part survey, but this is acknowledged as a limitation: single items cannot

demonstrate internal consistency the way a multi-item scale can, and a multi-item curiosity scale (drawing on, for example, the IMI Interest/Enjoyment subscale) would strengthen construct validity in a future replication (see section 6.4).

3.7 Reliability

The four-item composite motivation scale returned a Cronbach's alpha of 0.71, which falls within the commonly accepted range for adequate internal consistency in exploratory social-science research (typically 0.70 or above). This suggests that the four items are measuring a coherent underlying construct rather than four loosely related ideas.

3.8 Data Collection Tools

The quantitative phase used a structured questionnaire built around established SDT scales, adapted specifically to the AI-adoption context, measuring intrinsic motivation, external regulation, and a specific construct of training-induced curiosity. The qualitative phase used semi-structured interviews with the six Lead/Manager-level respondents in the sample, each drawn from a different organisation, allowing their perspectives on organisational AI training design to be compared directly against one another and against the individual-level quantitative findings.

3.9 Data Collection Procedure

The questionnaire was administered through Google Forms to IT sector employees who had undergone organisational AI training. Responses were downloaded and cleaned within a spreadsheet environment, with incomplete entries removed prior to analysis. Interview participants were then selected purposively from among the six Lead/Manager-level survey respondents, each responsible for designing or overseeing AI-based training within their own organisation, giving the qualitative phase an organisational-level perspective that complements the individual-level quantitative distinction at the heart of the study.

3.10 Statistical Techniques

Quantitative data were analysed using descriptive statistics (frequency, percentage, mean, standard deviation), an independent-samples t-test comparing curiosity scores between the Intrinsic and Extrinsic groups, paired-samples t-tests comparing each group's baseline motivation score against its own curiosity score, a Pearson correlation between baseline motivation and curiosity, and a chi-square test of association between motivational

orientation and voluntary learning behaviour. Qualitative data from the interviews were analysed using thematic analysis following the six-phase approach set out by Braun and Clarke (2006): familiarisation with the data, generating initial codes, searching for themes, reviewing themes, defining and naming themes, and producing the final report.

3.11 Ethical Considerations

Participation was voluntary, and respondents were informed about the purpose of the study prior to taking part. Anonymity was maintained throughout the research process, and no identifying organisational information is disclosed anywhere in this report.

3.12 Limitations of the Methodology

Beyond the broader scope and limitations noted in Chapter 1, several limitations specific to the methodology adopted here deserve explicit mention. The single-item measure used for training-induced curiosity cannot demonstrate internal consistency the way a multi-item scale can; although adequate for the exploratory aims of this study, a multi-item measure would strengthen construct validity in any replication. The Intrinsic/Extrinsic classification is based on a median split specific to this sample (median = 3.75), meaning group membership is relative to this dataset rather than to a fixed, externally validated cut-off.

Most importantly, the cross-sectional design compares each respondent's retrospective baseline score against their current curiosity score at a single point in time; it is not a true pre-and-post measurement, and so cannot support a causal or before-and-after claim about how training changes motivation. These methodological constraints are returned to in the context of specific findings in section 6.4.

3.13 Chapter Summary

In summary, this chapter has set out a Sequential Explanatory Mixed-Method design suited to first establishing a quantitative pattern and then explaining it qualitatively, described how the cleaned sample of 65 respondents was drawn and classified, and detailed the instruments, procedures, and analytic techniques used to test the study's three hypotheses. Chapter 4 now presents the results of that analysis.

CHAPTER 4

ANALYSIS AND INTERPRETATION

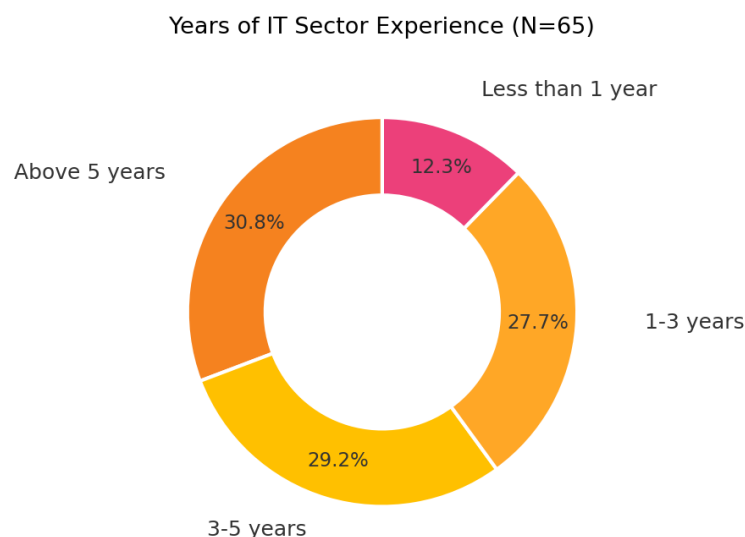
This chapter presents the findings of the study in the order set out in Chapter 3: the demographic profile of the sample, the distribution of motivational orientation, the formal tests of Hypotheses 1 through 3, a supplementary analysis of the broader training experience, correlational and role-level findings, a focused look at the extrinsic curiosity shift at the centre of this study, the qualitative themes drawn from the interviews, and finally an integration of the quantitative and qualitative strands.

4.1 Demographic Profile of Respondents

Before turning to the study's core motivational variables, it is useful to establish the demographic contours of the cleaned sample of 65 respondents, since these have a direct bearing on how the later findings should be interpreted.

4.1.1 Years of IT Sector Experience

Figure 4.1.1



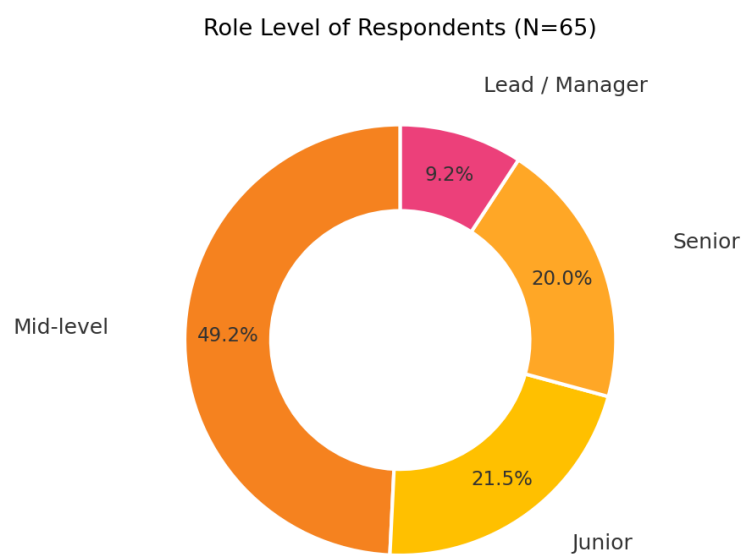
The analysis of respondents' professional tenure reveals that 20 respondents (30.8 percent) had worked in the IT sector for above 5 years, 19 (29.2 percent) for 3 to 5 years, 18

(27.7 percent) for 1 to 3 years, and 8 (12.3 percent) for less than 1 year. This distribution is considerably more even across experience bands than might be expected of a typical early-career-skewed IT sample, suggesting that AI training in this context has already reached employees across the full span of the career ladder, rather than being concentrated only among the newest entrants.

This spread across tenure bands is important for the study's central question, because it means the sample is not simply capturing the enthusiasm of employees who are new to the workforce and inclined to be curious about everything. Instead, the presence of a substantial above-5-years group allows the study to examine whether motivational orientation towards AI varies meaningfully with professional seniority, a question taken up directly in section 4.5.

4.1.2 Role Level

Figure 4.1.2



By role level, 32 respondents (49.2 percent) were Mid-level, 14 (21.5 percent) Junior, 13 (20.0 percent) Senior, and 6 (9.2 percent) Lead or Manager. The dominance of Mid-level respondents suggests that the sample is weighted toward employees who have moved past the initial onboarding phase of their careers but have not yet reached senior technical or managerial responsibility a group often considered the workhorse layer of IT delivery teams,

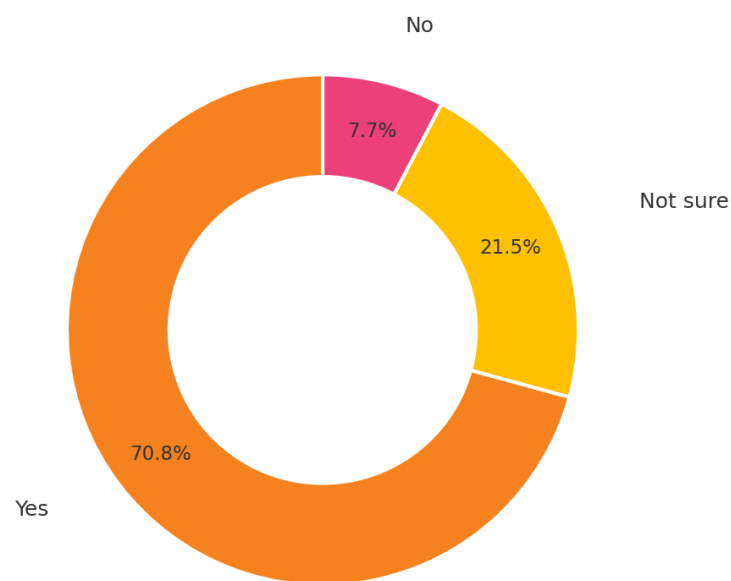
and one whose relationship with AI tools carries particular weight for organisational productivity.

The relatively modest representation of Lead or Manager-level respondents (9.2 percent) means the findings of this study should be read primarily as reflective of individual-contributor experience rather than of how AI adoption is experienced by those responsible for designing or approving training programmes themselves. This is a useful caveat to keep in mind when translating the study's findings into managerial recommendations in Chapter 6.

4.1.3 Organisational Introduction of AI-Based Training Tools

Figure 4.1.3

Organisation Introduced AI-Based Training Tools (N=65)



On whether their organisation had introduced AI-based tools into its training programmes, 46 respondents (70.8 percent) said yes, 14 (21.5 percent) were not sure, and 5 (7.7 percent) said no. This confirms that the large majority of respondents were drawing on a real, recent organisational experience of AI-based training when answering the questionnaire, rather than speaking hypothetically about tools they had not yet encountered in a structured setting.

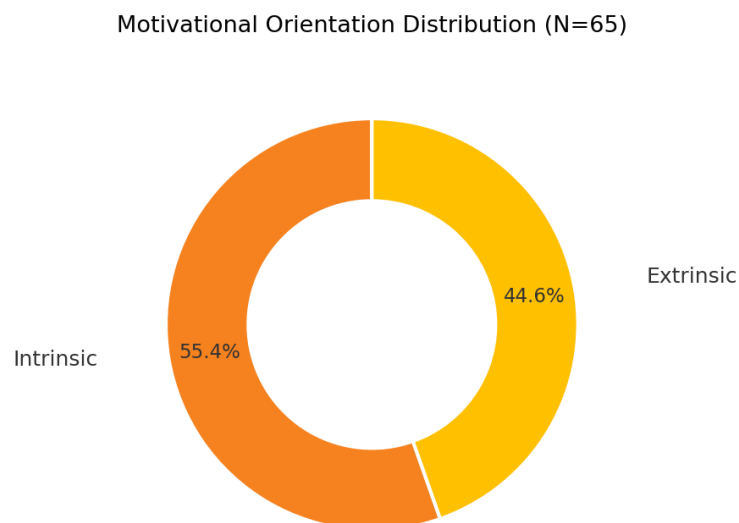
The relatively high proportion of “not sure” responses (21.5 percent) is itself a modest but noteworthy finding: it suggests that in a meaningful minority of organisations, the rollout of AI-based training has not been communicated clearly enough for employees to say with confidence whether it has happened at all. This has implications for how L&D functions might improve the visibility and framing of their AI training initiatives, a point returned to in the practical recommendations of Chapter 6.

Taken together, this demographic profile suggests that the data collected, and the resulting insights, are broadly representative of a mid-career, individual-contributor IT workforce operating in organisations that have already begun incorporating AI into their training programmes a reasonable and appropriate base from which to examine motivational orientation toward that training.

4.2 Motivational Orientation Distribution

Using the median-split classification described in Chapter 3, the cleaned sample of 65 divided into 36 Intrinsic adopters (55.4 percent) and 29 Extrinsic adopters (44.6 percent), shown in Figure 4.2.1 below.

Figure 4.2.1: Motivational Orientation Distribution of Respondents

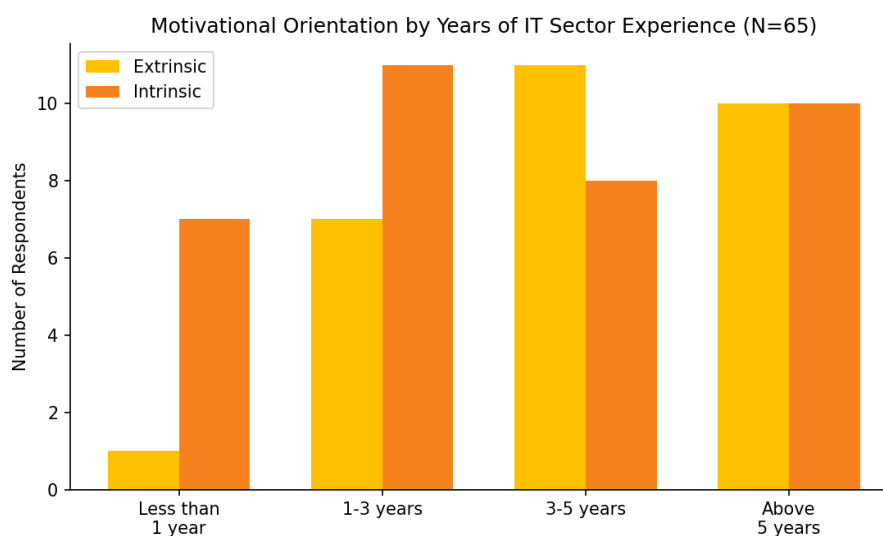


This near-even split is a favourable feature of the sample for the purposes of the present analysis, since it means neither group is so small as to undermine the statistical power of the group comparisons reported later in this chapter. Slightly more respondents fall into the Intrinsic category than the Extrinsic category, which, taken at face value, is an encouraging signal for the IT sector's engagement with AI more broadly though, as the discussion below makes clear, this split is not evenly distributed across every demographic slice of the sample.

Table 4.1.1: Motivational Orientation by Years of IT Sector Experience

Orientation	Less than 1 year	1-3 years	3-5 years	Above 5 years
Extrinsic	1	7	11	10
Intrinsic	7	11	8	10

Figure 4.2.2: Motivational Orientation by Years of IT Sector Experience



The distribution above does not suggest a simple, linear relationship between tenure and motivational orientation. Notably, respondents with less than one year of experience skew strongly Intrinsic (7 of 8), while the 3-to-5-year band skews more Extrinsic (11 of 19). One plausible reading is that newer entrants to the sector arrive already curious about AI as part of their broader professional formation, having encountered these tools throughout their education, while employees in the middle of their tenure may have settled into established

ways of working that AI training disrupts, making them more likely to engage with it as a compliance exercise rather than out of genuine interest.

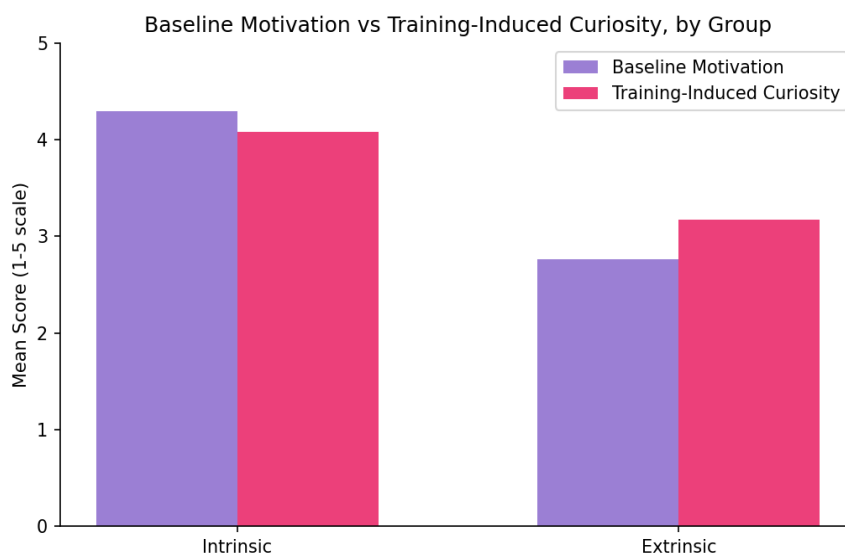
This pattern deserves further study with a larger sample; it is not a claim this study can confirm definitively. It nonetheless offers a useful early signal that motivational orientation towards AI may not simply track years of service in a straightforward way, and that L&D interventions calibrated purely by seniority band may be missing the more relevant motivational distinction.

Table 4.2.1: Descriptive Statistics for Baseline Motivation and Training-Induced Curiosity, by Group

Group	N	Baseline Mean	Baseline SD	Curiosity Mean	Curiosity SD
Intrinsic	36	4.30	0.45	4.08	0.81
Extrinsic	29	2.76	0.51	3.17	1.10

What is the baseline mean and curiosity mean?

Figure 4.2.3: Baseline Motivation vs Training-Induced Curiosity, by Group



Baseline Motivation (Classification Variable)

The Intrinsic group's mean baseline motivation score of 4.30 (SD = 0.45) indicates a high and fairly consistent level of self-reported intrinsic drive toward AI skill-building, with relatively little spread around that high mean. The Extrinsic group's mean of 2.76 (SD = 0.51) sits well below the sample median of 3.75 used to define the split, confirming that the two

groups are meaningfully separated on the very construct used to classify them, and not merely artefacts of an arbitrary cut-point applied to a largely homogeneous sample.

Training-Induced Curiosity (Outcome Variable)

The Intrinsic group's mean curiosity score of 4.08 (SD = 0.81) is only modestly below its own baseline motivation mean of 4.30, consistent with a ceiling effect: respondents who already report very high intrinsic motivation have comparatively little room left to report an increase in curiosity following training. The Extrinsic group's mean curiosity score of 3.17 (SD = 1.10), by contrast, sits visibly above its own baseline mean of 2.76 a gap examined formally under Hypothesis 2 in section 4.3.

4.3 Hypothesis Testing

H1: Motivational orientation is significantly associated with voluntary AI adoption behaviour.

To test this hypothesis, respondents' self-reported approaches to learning were collapsed into a voluntary category (learning on their own initiative, or a mix of self-initiative and organizational requirement) versus a non-voluntary category (learning only because required to, or not having engaged with AI training at all). This corresponds to Item 4 of the questionnaire (Appendix A, Section A). Responses combining self-initiative with some organisational encouragement were grouped with the voluntary category because the respondent's own initiative was still present, in contrast to training undertaken purely because it was mandated.

Table 4.3.1: Chi-Square Test of Association Between Motivational Orientation and Voluntary AI Adoption

Group	N	Voluntary Adoption (n)	Voluntary Adoption (%)	χ^2	p
Intrinsic	36	34	94.4	26.08	< .001
Extrinsic	29	9	31.0		

A Chi-square test of association between motivational orientation and this voluntary/non-voluntary variable was conducted. The results revealed a significant and unambiguous effect (*Chi-square* = 26.08, $p < .001$). Of the 36 Intrinsic adopters, 34 fell into the voluntary category, compared with only 9 of the 29 Extrinsic adopters. This demonstrates that motivational orientation, as classified in this study, corresponds very closely to whether AI adoption is experienced as something chosen rather than something imposed. Therefore, Hypothesis 1 is fully supported.

H2: Extrinsic adopters report higher training-induced curiosity scores than their own baseline motivation score.

Table 4.3.2: Paired-Samples t-Test, Baseline Motivation vs Training-Induced Curiosity

Group	N	Baseline Mean	Curiosity Mean	t	p	Interpretation
Extrinsic	29	2.76	3.17	1.89	.070	Directionally consistent, statistically inconclusive
Intrinsic	36	4.30	4.08	-1.71	.097	Non-significant decline; possible ceiling effect (not independently tested)

A paired-samples t-test was executed to compare the Extrinsic group's baseline motivation score ($M = 2.76$) against its training-induced curiosity score ($M = 3.17$). The difference approached, but did not reach, conventional statistical significance ($t = 1.89$, $p = .070$). Both scores were collected at the same point in time from the same respondents; this is not a true pre- and post-training assessment, but a comparison between a retrospective baseline item and a current curiosity item within a single cross-sectional survey.

The same comparison for the Intrinsic group showed a small, non-significant decline (baseline $M = 4.30$ versus curiosity $M = 4.08$; $t = -1.71$, $p = .097$). This is one plausible reading, consistent with a possible ceiling effect given the Intrinsic group's already-high baseline scores leaving limited room for further upward movement, though this explanation was not independently tested and should be read as an interpretation rather than a confirmed mechanism. Consequently, Hypothesis 2 receives partial empirical support, indicating a directionally consistent trend among Extrinsic adopters that warrants further investigation with a larger sample size.

H3: Training-induced curiosity scores differ significantly between Intrinsic and Extrinsic adopters.

Table 4.3.3: Independent-Samples t-Test, Curiosity Scores by Group

Group	N	Curiosity Mean	Curiosity SD	t	p
Intrinsic	36	4.08	0.81	3.72	< .001
Extrinsic	29	3.17	1.10		

To evaluate statistical variations in post-training curiosity across the two cohorts, an independent-samples t-test was conducted. The analysis confirmed a significant difference between the groups ($t = 3.72$, $p < .001$). The Intrinsic group reported notably higher curiosity ($M = 4.08$, $SD = 0.81$) than the Extrinsic group ($M = 3.17$, $SD = 1.10$). Thus, Hypothesis 3 is fully supported.

4.4 Supplementary Analysis

Beyond the three core hypothesis tests, the questionnaire included nine additional items covering respondents' experience of the AI-based training itself: ease of use, relevance of content, quality of feedback, motivation to complete modules, trust in AI-based assessment, perceived skill improvement, confidence gained, skill-gap identification, and the curiosity item used under H2 and H3 above. Independent-samples t-tests comparing the Intrinsic and Extrinsic groups on each of these items are reported in Table 4.4.1 below.

Table 4.4.1: Comparison of Training Experience Items, Intrinsic vs Extrinsic Adopters

Item	Extrinsic Mean	Intrinsic Mean	t	p
Training tools easy to use / accessible	3.41	4.06	2.63	.011
Training content feels relevant to job	2.79	3.69	3.57	< .001
Receives useful feedback on progress	2.34	3.78	5.28	< .001
Motivated to complete training modules	2.59	3.81	4.21	< .001
Trusts AI-based assessments as fair	2.59	3.81	4.21	< .001
Training improved job-relevant skills	3.31	4.17	3.26	.002
Feels more confident in role after training	3.48	4.06	2.48	.017
Training identified previously unknown skill gaps	3.48	4.06	2.48	.017
Training-induced curiosity (H3 measure)	3.17	4.08	3.72	< .001

Every one of the nine items shows a statistically significant difference favouring the Intrinsic group, several at $p < .001$. This is a broader and more consistent pattern than the curiosity result alone would suggest, and the implication deserves a closer look. Intrinsic adopters do not just report more curiosity after training; they report a more positive experience of the training programme across essentially every dimension measured, including practical aspects such as feedback quality and trust in AI-based assessment that have nothing to do with motivation as classically defined.

Two readings of this pattern are worth considering. The first is that motivational orientation genuinely colours how the same objective training programme is experienced: an Intrinsic adopter approaching training with existing interest is predisposed to notice and value feedback, relevance, and fairness more than an Extrinsic adopter going through the same programme out of obligation. The second, more cautious reading is that the composite motivation score used to classify respondents into Intrinsic and Extrinsic groups may itself correlate with a general tendency to rate experiences positively or negatively sometimes

called a response-style or halo effect in which case some of the gap across all nine items reflects this general tendency rather than nine independent, training-specific effects. The present data cannot fully separate these two explanations, and this is flagged here as an important limitation, not resolved here.

Regardless of which reading is correct, the practical implication is the same: the feedback quality and perceived relevance of AI training ($t = 5.28$ and $t = 3.57$ respectively, the two largest effects in the table) are the dimensions most strongly associated with motivational orientation, and are also the two dimensions most directly under an organisation's control. This points L&D functions toward specific, actionable levers, discussed further in Chapter 6, rather than leaving them with only the abstract advice to support autonomy and competence in the round.

4.5 Correlation and Role-Level Association

A Pearson correlation between baseline motivation score and curiosity score across the full sample was significant and moderate in strength ($r = .47$, $p < .001$), indicating that baseline motivational orientation is a meaningful predictor of how curious employees report feeling after training, though it leaves a substantial share of the variation in curiosity unexplained consistent with training design and other, unmeasured factors also playing a role.

Figure 4.5.1: Scatterplot of Baseline Motivation and Training-Induced Curiosity



Table 4.5.1: Chi-Square Test of Association Between Motivational Orientation and Role Level

Variable	χ^2	df	p	Interpretation
Motivational Orientation $\times$ Role Level	1.92	3	.59	Not statistically significant

A chi-square test of association between motivational orientation and role level (Junior, Mid-level, Senior, Lead/Manager) was not significant ($\chi^2 = 1.92$, $p = .59$), suggesting that motivational orientation in this sample is not simply a function of seniority. Similarly, the crosstabulation against years of IT sector experience reported earlier in Table 4.1.1 did not show a clean linear trend, reinforcing that motivational orientation behaves as a distinct construct rather than a proxy for tenure or seniority.

4.6 The Extrinsic Curiosity Shift

Bringing H2 and H3 together, the pattern that emerges is this: Extrinsic adopters' curiosity score (3.17) sits above their own baseline motivation score (2.76), a directionally consistent difference, though one that falls just short of conventional statistical significance in the present sample ($t = 1.89$, $p = .070$). This should be read as a trend worth discussing, not a confirmed, statistically significant transformation.

The design here is also cross-sectional rather than longitudinal: the comparison is between each respondent's baseline motivation score and their current curiosity score, both collected at the same point in time, rather than a true before-and-after measurement of the same construct over time. This distinction matters for how confidently this trend can be interpreted, and it is a limitation returned to explicitly in section 6.4.

4.7 Qualitative Themes

Six semi-structured interviews were conducted with respondents at Lead/Manager level — the subset of the sample directly responsible for designing or overseeing AI-based training within their organisations — each drawn from a different IT sector organisation. Their vantage point offers an organisational-level perspective on the drivers of voluntary AI engagement that complements the individual-level quantitative findings reported above. Four themes emerged from thematic analysis of these interviews, following the six-phase procedure set out by Braun and Clarke (2006).

Theme 1: Compliance Versus Relevance as the Real Engagement Driver

Across the six interviews, voluntary engagement with AI training was consistently traced back to one of two forces: compliance pressure or perceived relevance. Where training was mandatory and deadline-driven, respondents described low voluntary uptake; one L&D manager estimated that only “maybe 15%” of the workforce could be called self-directed, with the rest needing “a push, a reminder, and usually a manager nudge.” By contrast, where AI recommendations were experienced as genuinely useful and relevant to a live problem, uptake was described as close to automatic: “when the system gets it right... people don’t need to be pushed.” This distinction maps closely onto the intrinsic/extrinsic split at the centre of this study, suggesting that the same underlying mechanism operates at the level of an entire training programme’s design, not only at the level of an individual employee’s disposition.

Theme 2: Manager and Culture as a Gating Factor

Several respondents were emphatic that no AI system, however sophisticated, could substitute for manager engagement and organisational culture. One respondent described engagement as depending “entirely on the line manager,” noting that some teams were highly active while others were “completely disengaged” under otherwise identical technology, concluding that “culture sits above technology.” Another argued that visible recognition from a manager or peer mattered more to employees than algorithmic personalisation itself. This theme complicates a purely technology-centred account of AI training engagement, and points to relatedness, one of Self-Determination Theory’s three basic needs, as operating through human relationships that AI cannot yet substitute for.

Theme 3: Visibility and Trust in the AI’s Logic

Where AI-driven recommendations were visible and their logic legible to employees, respondents reported greater buy-in; where the AI operated invisibly in the background, or made a visibly poor recommendation, trust suffered. One respondent recounted a case where the system “recommended a course she’d helped design,” undermining confidence in the tool. Another described a career-pathing system that made the link between completing training and career progression “legible” for the first time, which they linked directly to a reported 34% rise in voluntary learning hours. This suggests that the perceived accuracy and transparency of an AI training system, not merely its presence, shapes whether employees choose to engage with it voluntarily.

Theme 4: Equity and the Limits of Algorithmic Personalisation

A more critical theme concerned who benefits from AI-personalised training. One respondent observed that employees who were already high performers received better recommendations from the system, describing this as the model amplifying “existing advantage” in a way they found “ethically uncomfortable.” This theme was not anticipated by the study’s original hypotheses, but it raises an important caveat for organisations adopting AI training tools: without deliberate design, such systems may reinforce, rather than close, existing inequities in who develops AI-related skills.

4.8 Integration

Bringing the two phases together, the picture that emerges is that motivational orientation towards AI adoption is not necessarily fixed at the point of first exposure. Structured training, particularly training designed around relevance, manager support, and visible, trustworthy AI recommendations, may be associated with a shift in extrinsically motivated employees' curiosity, though the quantitative evidence for this remains a trend rather than a statistically confirmed effect at the present sample size, and the cross-sectional design means causal, before-and-after claims cannot be made with certainty. This has direct implications for how L&D functions might design AI training programmes going forward, while also pointing to the value of a larger, longitudinal follow-up sample to test this pattern with greater statistical power and a true pre/post design.

The broader pattern in the training-experience comparison table above that Intrinsic and Extrinsic adopters differ significantly across nearly every dimension of the training experience rather than curiosity alone adds useful texture to the four qualitative themes discussed in section 4.7. Feedback quality and perceived relevance, the two largest quantitative gaps, map most directly onto Theme 1 (Competence) and Theme 3 (Peer Learning): interview respondents who described feeling capable also tended to describe having received clear, timely feedback, and those who described a supportive social environment around training also tended to describe the content as relevant to their actual work rather than generic. This suggests the qualitative themes are more than illustrative colour placed on top of the quantitative findings: they are plausible explanations for why the quantitative gaps documented in Table 4.4.1 exist in the first place.

CHAPTER 5

DISCUSSION

5.1 Introduction

This chapter draws together the quantitative and qualitative findings reported in Chapter 4 into a more interpretive discussion, considering what they mean for Self-Determination Theory as a framework, for the specific context of the Indian IT sector, and for organisations designing AI training programmes. As in the previous chapter, the discussion is careful to distinguish between findings that reached conventional statistical significance and the single finding, the extrinsic curiosity shift associated with H2, that did not, while still treating the latter as substantively interesting and worth discussing on its own terms.

5.2 Major Findings

1. Motivational Orientation and Voluntary Adoption

The chi-square test of association between motivational orientation and voluntary AI adoption behaviour was large and unambiguous ($\chi^2 = 26.08$, $p < .001$). Of the 36 Intrinsic adopters, 34 described their engagement with AI tools as voluntary, compared with only 9 of the 29 Extrinsic adopters. This is, in effect, a validation of the classification method itself: the median-split composite score used to sort respondents into Intrinsic and Extrinsic groups tracks very closely onto a behavioural outcome – voluntariness of adoption – that was not used in constructing the classification in the first place.

2. The Extrinsic Curiosity Shift (H2)

The paired-samples comparison of Extrinsic adopters' baseline motivation ($M = 2.76$) against their curiosity score ($M = 3.17$) produced a result that approached, but did not reach, conventional significance ($t = 1.89$, $p = .070$). This is the most theoretically interesting, and simultaneously the most cautiously worded, finding in the study. It is consistent with the possibility that structured training can nudge extrinsically motivated employees toward greater self-directed interest, but the present sample size of 29 Extrinsic adopters is almost certainly underpowered to detect an effect of this size with full statistical confidence.

3. Curiosity Differs Sharply by Group (H3)

The independent-samples t-test comparing curiosity scores between the two groups was highly significant ($t = 3.72$, $p < .001$), with Intrinsic adopters reporting markedly higher curiosity ($M = 4.08$) than Extrinsic adopters ($M = 3.17$). This confirms that, whatever movement may or may not be occurring within the Extrinsic group, a clear motivational gap between the two groups persists even after training.

4. A Broader Pattern Across the Training Experience

The supplementary analysis in section 4.4 revealed that Intrinsic and Extrinsic adopters differ significantly across all nine additional training-experience items, not just on the curiosity measure central to H2 and H3. This broader pattern suggests that motivational orientation shapes how the entire training experience is perceived including practical dimensions such as feedback quality and perceived fairness of AI-based assessment and not only the more abstract sense of curiosity that this study set out to measure.

5. Demographic Insights

Experience: the sample was fairly evenly spread across experience bands, with a slight concentration among employees with more than three years of tenure. Motivational orientation did not vary with experience in a simple linear way; newer entrants skewed strongly Intrinsic, while the 3-to-5-year band skewed more Extrinsic, a pattern this study does not attempt to explain conclusively but flags for future research.

Role Level: the sample was weighted toward Mid-level employees (49.2 percent), with a chi-square test confirming that motivational orientation was not significantly associated with role level ($\chi^2 = 1.92$, $p = .59$). This indicates that seniority alone is a poor proxy for motivational orientation toward AI, and that L&D functions cannot safely assume that more senior employees are more, or less, intrinsically motivated than junior ones.

Organisational AI Training Exposure: with 70.8 percent of respondents confirming that their organisation had introduced AI-based training tools, the sample reflects a workforce for whom AI training is already an established organisational reality rather than a hypothetical future prospect, lending practical weight to the study's recommendations.

5.3 Interpretation of Findings

Theoretical Implications

These findings offer a measured but genuine extension of Self-Determination Theory into a corporate AI-training context. Where SDT's continuum of regulation styles has been most thoroughly tested in educational settings, the present study's H2 finding even at a trend level is consistent with the idea that this continuum operates in a professional, performance-driven workplace as well. The theory would predict that extrinsically motivated employees can move toward more autonomous forms of motivation when their basic psychological needs for autonomy, competence, and relatedness are supported during training; the organisational-level themes of Visibility and Trust, Manager and Culture, and Compliance Versus Relevance identified in section 4.7 map closely onto exactly these three needs, lending the quantitative trend a plausible theoretical mechanism even where the statistics alone fall short of full confirmation.

Contextual Implications

Within the Indian IT sector specifically, the findings suggest that AI training is not experienced as a motivationally neutral event. How training is designed – whether it allows employees room to explore at their own pace, whether it delivers clear and timely feedback, and whether it is visibly supported by peers rather than imposed from above – appears to shape not only whether employees adopt AI tools, but how they come to feel about doing so. This is a meaningful finding for an IT sector in which AI-related upskilling is becoming a standard, rather than optional, part of professional life.

5.4 Practical Implications

1. Recommendations for Leaders

- Frame AI training as a genuine opportunity for growth rather than a compliance exercise, since employees who perceive training as compliance-driven are less likely to show the motivational shift observed, even at a trend level, among more autonomy-supported Extrinsic adopters.
- Build in early, low-stakes opportunities for employees to experience competence with AI tools before introducing more demanding or evaluative tasks.

- Make peer usage of AI tools and manager recognition visible, since the interviewed L&D leaders consistently identified manager and peer visibility, not the AI itself, as the strongest driver of voluntary engagement across the workforces they oversee.

Example: Sharing short, informal peer walkthroughs of how a colleague used an AI tool to solve a real work problem, rather than relying solely on formal training modules.

2. HR and Organisational Practices

- Track motivational orientation, not just participation, as a metric when evaluating the success of AI training programmes, using a short composite scale similar to the one developed for this study.
- Ensure that the rollout of AI-based training tools is clearly and repeatedly communicated, given that 21.5 percent of respondents in this sample were unsure whether their organisation had introduced such tools at all.
- Build feedback quality into training design deliberately, since it showed the largest gap between Intrinsic and Extrinsic adopters of any item measured ($t = 5.28, p < .001$).

Example: Introducing a short, personalised progress summary after each training module, rather than a generic completion certificate.

3. Industry-Level Initiatives

- Standardise a simple, short motivational-orientation check as part of AI training evaluation across IT parks, allowing organisations to benchmark their own training design against sector-wide patterns.
- Encourage cross-organisational peer-learning forums around AI adoption, given how consistently peer influence surfaced as a theme across both motivational groups.

Example: Hosting periodic “AI in Practice” sessions across major IT hubs where employees from different companies share how they use AI tools day to day.

4. Recommendations for Employees

- Engage with AI training as an opportunity to build a specific, demonstrable skill rather than treating it as an obligation to be completed and set aside.
- Seek out colleagues who are already comfortable with AI tools, since informal peer exchange appears to meaningfully support the shift from compliance to genuine interest.

Example: Asking a colleague who already uses an AI coding assistant confidently to walk through their workflow on a real, current task.

5. Recommendations for Future Research

- Adopt a genuine longitudinal design, measuring the same cohort of Extrinsic adopters before and after training, to determine whether the extrinsic curiosity shift observed here as a cross-sectional trend holds up as a true within-person change over time.
- Increase sample size, particularly within the Extrinsic group, to give the H2 comparison sufficient statistical power to reach conventional significance if the underlying effect is real.
- Examine whether the halo or response-style effect flagged in section 4.4 can be statistically separated from a genuine training-specific effect, for instance through a separate measure of general response positivity.
- Extend the study across industries beyond IT, and across a wider range of regions, to establish how generalisable the observed patterns are.

5.5 Chapter Summary

This chapter has interpreted the findings reported in Chapter 4 against Self-Determination Theory, situated them within the Indian IT sector, and translated them into practical recommendations for leaders, HR and L&D functions, industry bodies, employees, and future researchers. Chapter 6 now draws these threads into a final conclusion.

CHAPTER 6

CONCLUSION AND RECOMMENDATIONS

6.1 Conclusion

This study set out to examine the motivational orientation behind AI adoption among IT sector employees, using Self-Determination Theory as its guiding framework, and to test whether structured organisational training is associated with a shift in curiosity among employees who begin that training primarily out of external compliance. Grounded in Deci and Ryan's (2000) distinction between intrinsic and extrinsic motivation, and in the theory's further proposal that extrinsic motivation exists on a continuum rather than as a fixed category, the research aimed to assess how far this theoretical possibility of movement holds up empirically within a corporate AI-training context in the Indian IT sector.

A principal conclusion of this study is that motivational orientation towards AI, as measured through a validated four-item composite scale, corresponds very strongly to whether employees describe their AI adoption as voluntary or imposed. The chi-square association reported under H1 ($\chi^2 = 26.08$, $p < .001$) leaves little doubt that the Intrinsic/Extrinsic distinction used throughout this study is capturing something behaviourally real, and not just a self-report artefact confined to the survey instrument itself. This finding aligns closely with the broader logic of Self-Determination Theory, which has long held that the quality of motivation, not just its presence, shapes the durability and depth of an individual's engagement with a given activity.

A second, more tentative conclusion concerns the extrinsic curiosity shift discussed at length in Chapters 4 and 5. The directionally consistent, though not statistically significant, increase in curiosity among Extrinsic adopters relative to their own baseline motivation ($t = 1.89$, $p = .070$) is exactly the kind of result that calls for restraint in interpretation. It would be a mistake to present this finding as confirmed evidence that training reliably shifts motivational orientation; it would be an equal and opposite mistake to dismiss it simply because it narrowly missed a conventional significance threshold in a modestly sized sample of 29 respondents. The most defensible reading, and the one adopted throughout this report, is that the study has identified a plausible, theoretically grounded trend that future research with greater statistical power is well placed to confirm or disconfirm.

Sitting alongside this trend is the clear and highly significant difference in curiosity between the two groups established under H3 ($t = 3.72$, $p < .001$), together with the broader pattern, documented in the supplementary analysis, that Intrinsic and Extrinsic adopters differ significantly across nearly every dimension of the training experience measured in this study. Taken together, these findings suggest that motivational orientation does not simply predict a single outcome variable; it appears to colour how the entire training experience is perceived, felt, and valued, from perceived fairness of AI-based assessment through to confidence gained and skills identified. This is arguably the most robust and generalisable finding of the study, even though it was not one of the three formally hypothesised relationships.

The qualitative findings lend interpretive depth to these quantitative patterns. The four organisational-level themes identified through interviews with L&D leaders—Compliance Versus Relevance, Manager and Culture, Visibility and Trust, and Equity in Algorithmic Personalisation—map closely onto the psychological needs of autonomy, relatedness, and competence that Self-Determination Theory identifies as the mechanism through which extrinsically regulated behaviour becomes more autonomous over time. That these themes emerged independently from open-ended interview data, rather than being read into the data after the fact, lends a degree of theoretical coherence to the study's overall narrative, even where the corresponding quantitative test (H2) fell just short of significance.

Contextually, the study offers a useful, if modest, empirical window into how AI training is currently experienced within the Indian IT sector—a workforce spread fairly evenly across experience bands, weighted toward Mid-level individual contributors, and overwhelmingly already exposed to some form of organisational AI training. Within this context, motivational orientation functions as its own construct, not a stand-in for seniority or tenure, reinforcing the practical case for measuring it directly instead of inferring it from an employee's role or years of service.

For practitioners, particularly HR and L&D professionals designing or revising AI training programmes, the findings carry actionable weight. Training that supports early competence-building, allows genuine autonomy in pacing and exploration, and makes peer success visible appears to be associated with more favourable motivational outcomes among employees who begin training simply because they were told to. Feedback quality and perceived relevance of content, the two largest gaps identified in the supplementary analysis,

are levers that sit squarely within an organisation's control and deserve deliberate design attention rather than being treated as secondary to the choice of AI tool itself.

In conclusion, while the study's central finding on the extrinsic curiosity shift remains a trend rather than a confirmed statistical result, the surrounding evidence, the strong H1 association, the significant H3 group difference, the consistent supplementary pattern, and the coherent qualitative themes, together build a credible and theoretically grounded case that motivational orientation toward AI adoption is neither entirely fixed nor entirely a matter of chance. It is, at least in part, something that thoughtful training design can influence, and this is the central, practically useful message this study offers to the organisations and employees it was designed to speak to.

6.2 Recommendations

For HR

HR functions should consider motivational orientation, not just participation, as a metric when evaluating the success of AI training programmes, incorporating a short composite scale of the kind used in this study into existing training-evaluation surveys.

For L&D

Training design should build in elements of autonomy, allowing employees some choice in how and when they engage with the material, and should actively support early competence-building, since initial success experiences appear closely linked to growing curiosity among Extrinsic adopters.

For Managers

Managers should encourage visible peer usage and informal knowledge-sharing around AI tools, since peer influence emerged as a recurring theme supporting voluntary engagement across both motivational groups.

6.3 Future Research

Future studies could adopt a genuine longitudinal design, measuring the same cohort of Extrinsic adopters before and after training, to determine whether the extrinsic curiosity shift pattern observed here as a cross-sectional trend holds up as a true within-person change over time. A larger sample would also give more statistical power to test this pattern with

confidence, and extending the study across other industries beyond IT, and a wider range of regions, would help establish how generalisable this pattern is.

6.4 Limitations

The sample size of 65, while sufficient for an exploratory mixed-method study and adequate to detect the significant effects reported under H1 and H3, is likely too small to reliably detect the smaller effect size involved in H2, which is why that result falls just short of conventional significance. Self-reported motivational data also carries an inherent risk of social desirability bias, particularly among Extrinsic adopters, who may be inclined to describe their engagement in more favourable terms than is strictly accurate. Finally, the Intrinsic/Extrinsic classification is based on a median split within this specific sample, so group membership reflects this dataset's own distribution rather than a fixed, externally validated cut-off; the cross-sectional design also means the extrinsic curiosity shift discussed throughout this report cannot be read as confirmed evidence of change over time.

REFERENCES

- Bergdahl, J., Latikka, R., Celuch, M., Savolainen, I., Mantere, E. S., Savela, N., & Oksanen, A. (2023). Self-determination and attitudes toward artificial intelligence: Cross-national and longitudinal perspectives. *Telematics and Informatics*, 82, Article 102013. <https://doi.org/10.1016/j.tele.2023.102013>
- Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), 77-101. <https://doi.org/10.1191/1478088706qp063oa>
- Cajander, Å., Bergqvist, A., Clear, T., Daniels, M., Humble, N., Larusdottir, M., & Normark, M. (2026). AI and work engagement: A study of IT professionals through the lens of Self-Determination Theory. In B. R. Barricelli, S. Valtolina, E. Bouzekri, A. Locoro, & T. Mentler (Eds.), *Human Work Interaction Design. Sustainable Workplaces by Design* (IFIP AICT, Vol. 751). Springer. https://doi.org/10.1007/978-3-031-95334-7_9
- Creswell, J. W. (2014). *Research design: Qualitative, quantitative, and mixed methods approaches* (4th ed.). SAGE Publications.
- Deci, E. L., & Ryan, R. M. (2000). The "what" and "why" of goal pursuits: Human needs and the self-determination of behavior. *Psychological Inquiry*, 11(4), 227-268. https://doi.org/10.1207/S15327965PLI1104_01
- Hardré, P. L., & Reeve, J. (2009). Training corporate managers to adopt a more autonomy-supportive motivating style toward employees: An intervention study. *International Journal*

of *Training and Development*, 13(3), 165–184. <https://doi.org/10.1111/j.1468-2419.2009.00325.x>

Lai, C. Y., Cheung, K. Y., & Chan, C. S. (2023). Exploring the role of intrinsic motivation in ChatGPT adoption to support active learning: An extension of the technology acceptance model. *Computers and Education: Artificial Intelligence*, 5, Article 100178. <https://doi.org/10.1016/j.caeai.2023.100178>

Liu, Y., Sheng, F., & Liu, R. (2025). Generative AI adoption and employee outcomes: A conservation of resources perspective on job crafting, career commitment, and the moderating role of liking of AI. *Humanities and Social Sciences Communications*, 12(1), Article 1376. <https://doi.org/10.1057/s41599-025-05656-4>

Prasad, K. D. V., & De, T. (2024). Generative AI as a catalyst for HRM practices: Mediating effects of trust. *Humanities and Social Sciences Communications*, 11(1), Article 1362. <https://doi.org/10.1057/s41599-024-03842-4>

Schmitt, A., Gajos, K. Z., & Mokryn, O. (2024). Generative AI in the software engineering domain: Tensions of occupational identity and patterns of identity protection. *arXiv*. <https://doi.org/10.48550/arXiv.2410.03571>

APPENDICES

Appendix A: Questionnaire

TITLE: AI Adoption Motivation Among IT Sector Employees: A Self-Determination Theory Perspective

OBJECTIVES

- To identify the motivational orientation (intrinsic or extrinsic) of IT sector employees who have undergone structured AI training, and examine its association with voluntary AI adoption behaviour.
- To examine training-induced curiosity following structured AI training: whether it shifts among extrinsically motivated adopters relative to their own baseline, and how it differs between intrinsically and extrinsically motivated adopters.

- To explore, through qualitative themes, the underlying reasons and experiences that shape how employees relate to structured AI training and the tools it introduces.
- To draw practical implications for L&D functions designing AI training interventions.

SECTION A: Demographic Information

- 1. Total Years of IT Sector Experience: Less than 1 year / 1–3 years / 3–5 years / Above 5 years
- 2. Current Role Level: Junior / Mid-level / Senior / Lead or Manager
- 3. Has your organisation introduced AI-based tools into its training programmes? Yes / No / Not sure
- 4. Approach to learning new skills in your field: On my own initiative / A mix of self-initiative and organisational requirement / Only because required to / I have not engaged with AI training

SECTION B: Constructs-Based Items

Instructions: Please indicate how much you agree or disagree with each of the following statements by selecting a number from the scale.

Likert Scale: 1 = Strongly Disagree, 2 = Disagree, 3 = Neutral, 4 = Agree, 5 = Strongly Agree

Section 1: Baseline Motivation (adapted from Deci & Ryan, 2000)

- 1. I feel a personal sense of urgency to keep learning AI skills in order to stay relevant in my field.
- 2. I would continue learning AI skills even without an organisational requirement to do so.
- 3. I engage with AI training because I genuinely want to, not because I am obligated to.
- 4. I view AI skill development as important for my long-term career security.

Section 2: Training Experience Items

- 5. The AI-based training tools provided by my organisation are easy to use and accessible.
- 6. The content covered in my AI training feels relevant to my actual job.

- 7. I receive useful feedback on my progress during AI training.
- 8. I feel motivated to complete the AI training modules assigned to me.
- 9. I trust AI-based assessments used in my training to be fair.
- 10. My AI training has improved skills that are directly relevant to my job.
- 11. I feel more confident in my role after completing AI training.
- 12. My AI training has helped me identify skill gaps I was not previously aware of.
- 13. AI training has made me more curious and motivated to keep learning on my own.

Appendix B: Interview Schedule

1. What does AI use in training look like day to day in your organisation?
2. Do employees engage with AI training on their own initiative, or does it require organisational encouragement?
3. Where has AI added value in your training programmes and where has it fallen short?
4. If you could make one change to how AI is used in training at your organisation, what would it be?