



**INSTITUTO SUPERIOR DE TECNOLOGIAS AVANÇADAS DE
LISBOA**

Mestrado em Informática

**Dissertação para obtenção do Grau de Mestre em
Informática**

**The Impact of Artificial Intelligence on Employment:
Individuals' Perceptions**

Daniel Castanheira Lourenço Oliveira

Presidente: Professor Doutor Paulo André Reis Duarte Branco

Arguente: Professor Doutor Pedro Ramos dos Santos Brandão

Orientadora: Professora Doutora Carla Sofia Rocha da Silva

Lisboa

Novembro 2025

ISTEC
INSTITUTO SUPERIOR DE TECNOLOGIAS AVANÇADAS DE LISBOA

Campus Académico do Lumiar, Lisboa

Dissertação
Mestrado em Informática

Daniel Castanheira Lourenço Oliveira

Dissertação apresentada para cumprimento dos requisitos necessários à obtenção do grau de Mestre em Informática, realizada sob a orientação científica da Professora Doutora Carla Sofia Rocha da Silva.

Lisboa, 2025

Resumo

A rápida difusão da Inteligência Artificial (IA) nos contextos de trabalho tem reconfigurado tarefas e reavivado o debate sobre implicações para o emprego. Esta dissertação analisa as percepções de profissionais acerca do impacto da IA em quatro dimensões: benefícios ao nível da tarefa, riscos ao nível do posto de trabalho, necessidades de competências e formação, e considerações éticas. Recorreu-se a um inquérito transversal aplicado a 75 profissionais de vários países, maioritariamente do setor das Tecnologias de Informação (83 %). O questionário integrou 20 itens em escala de *Likert* e quatro questões abertas; os dados foram analisados por estatística descritiva e codificação temática. Dadas a amostragem por conveniência, a concentração setorial e a ausência de estratificação geográfica, os resultados são exploratórios e não generalizáveis.

Os resultados apontam para uma percepção amplamente positiva da produtividade associada à IA: 84 % reportaram ganhos de produtividade e 94 % concordaram que a IA poupa tempo nas tarefas. As preocupações com substituição de funções surgem mitigadas: 28 % referiram ansiedade com substituição, enquanto 56 % discordaram dessa hipótese. O impacto no stress foi modesto (15 %) e 69 % indicaram que a IA lhes liberta tempo para atividades de maior valor acrescentado. Ao nível das competências, 91 % declararam confiança no uso de tecnologias de IA e 84 % referiram estar a desenvolver novas competências; o apoio organizacional foi geralmente positivo (85 %), embora 35 % não tenham recebido formação formal, recorrendo sobretudo a autoaprendizagem.

Emergiu, contudo, um paradoxo: apesar da elevada autoconfiança (~ 91 %), 45 % manifestaram preocupações éticas – sobretudo privacidade e segurança de dados (29 %), transparência algorítmica (20 %) e enviesamentos em decisões mediadas por IA. Ainda assim, as atitudes globais mantiveram-se otimistas: 75 % expressaram entusiasmo face ao aumento do uso de IA no seu setor e 85 % intenção de adoção adicional. Em síntese, os inquiridos veem a IA sobretudo como potenciador de produtividade, mais do que como ameaça direta ao emprego; porém, a integração bem-sucedida dependerá de mitigação de riscos éticos e reforço da formação, garantindo utilização responsável e equitativa.

Palavras-chave: Inteligência Artificial; Emprego; Produtividade; Percepções no local de trabalho; Segurança no emprego; Ética da IA.

Abstract

Artificial intelligence (AI) is rapidly diffusing across workplaces, reshaping tasks and reigniting debates about its implications for employment. This dissertation examines how professionals perceive AI's impact across four key dimensions: task-level benefits, job-level risks, skill and training needs, and ethical considerations. A cross-sectional survey was conducted with 75 professionals from multiple countries, the majority of whom work in the information-technology sector (83%). The questionnaire included 20 Likert-scale items and four open-ended questions, and the resulting data were analyzed using descriptive statistics and thematic coding. Given the convenience sampling, sectoral concentration, and the absence of country-level stratification, the findings should be considered exploratory rather than generalizable.

Results reveal widespread perceptions of AI-enabled productivity: 84% of respondents reported productivity gains, and 94% agreed that AI saves time on tasks. Job-replacement concerns were limited – 28% expressed anxiety that AI might replace their roles, whereas 56% explicitly disagreed. Reported stress impacts were modest (15%), and 69% indicated that AI freed them to focus on higher-value activities. In terms of skills and adaptation, approximately 91% felt confident using AI technologies, and 84% reported actively pursuing new AI-related skills. Organizational support was generally positive (85% cited management encouragement for AI adoption), although 35% indicated they had received no formal employer-provided training; many nonetheless perceived adequate resources for self-learning.

A notable paradox emerged: despite high levels of self-efficacy (~91%), 45% of respondents reported ethical concerns – most frequently related to data privacy and security (29%), algorithmic transparency (20%), and bias in AI-mediated decision-making. Nevertheless, overall attitudes remained positive: 75% expressed enthusiasm about the growing use of AI in their field, and 85% indicated an intention to adopt additional AI tools.

In summary, respondents predominantly viewed AI as a productivity enhancer rather than a direct threat to employment. However, successful integration will depend on addressing ethical risks and bridging training gaps to ensure responsible and equitable use of AI technologies.

Keywords: Artificial Intelligence; Employment; Productivity; Workplace Perceptions; Job Security; AI Ethics.

Acronyms

ADP – Automatic Data Processing (ADP) Research Institute
AI – Artificial Intelligence
ANOVA – Analysis of Variance
CAQDAS – Computer-Assisted Qualitative Data Analysis Software
FATE – Fairness, Accountability, Transparency and Ethics (in AI)
GDPR – General Data Protection Regulation
GPT – Generative Pre-trained Transformer
HR / HRM – Human Resources / Human Resource Management
IA – Inteligência Artificial
IEEE – Institute of Electrical and Electronics Engineers
ILO – International Labour Organization
IT – Information Technology
JD-R – Job Demands–Resources (model)
KPI – Key Performance Indicator
LLM – Large Language Model
MLOps – Machine Learning Operations
NLP – Natural Language Processing
O*NET – Occupational Information Network (U.S.)
OECD – Organisation for Economic Co-operation and Development
PAIR – People + AI Research (Google)
QUAL – Qualitative (strand in mixed-methods design)
QUAN – Quantitative (strand in mixed-methods design)
R&D – Research and Development
ROI – Return on Investment
RPA – Robotic Process Automation
RQ – Research Question
SCT – Social Cognitive Theory
SOP – Standard Operating Procedure
STARA – Smart Technology, Artificial Intelligence, Robotics and Algorithms
TAM – Technology Acceptance Model
TI – Tecnologia da Informação
TTF – Task–Technology Fit
UTAUT – Unified Theory of Acceptance and Use of Technology
WEF – World Economic Forum

Index

Resumo.....	I
Abstract.....	III
Acronyms.....	V
Index.....	VII
List of Tables	XI
1. Introduction.....	1
1.1 Background.....	1
1.2 Motivation	2
1.3 Problem Statement.....	3
1.4 Research Objectives and Questions.....	5
1.5 Research Methodology.....	6
1.6 Significance of the Study.....	7
1.7 Structure of the Dissertation	8
1.7.1 Chapter 1: Introduction	8
1.7.2 Chapter 2: Literature Review	9
1.7.3 Chapter 3: Methodology.....	10
1.7.4 Chapter 4: Results and Discussion	11
1.7.5 Chapter 5: Discussion.....	12
1.7.6 References	13
1.7.7 Appendix A: Survey Instrument – Questionnaire	13
2. Literature Review.....	15
2.1 Theoretical Foundations	15
2.1.1 Technology Acceptance Theories	15
2.1.2 Organizational Change and Learning Theories.....	16
2.1.3 Human–AI Collaboration Frameworks	17
2.2 Empirical Evidence on AI’s Impact	18
2.2.1 Productivity Studies	18
2.2.2 Employment Impact Studies	21
2.2.3 Measurement Approaches in AI Impact Research	24
2.3 Focused Evidence Base Motivating the Research Questions	26
2.3.1 Productivity and Work Quality (RQ1).....	26
2.3.2 Job Insecurity, Displacement Risk, and Technostress (RQ2)	27
2.3.3 Skills, Upskilling, and Organizational Support (RQ3)	27

2.3.4	Ethics, Transparency, and Trust in AI (RQ4)	28
2.4	Workforce Transformation in the AI Era	28
2.4.1	Job Displacement vs. Augmentation.....	28
2.4.2	Skills and Competencies for the AI Era	31
2.5	Ethical and Societal Implications of AI in Employment	33
2.5.1	Algorithmic Bias, Fairness, and Accountability	33
2.5.2	Privacy and Surveillance.....	36
2.5.3	Worker Rights, Well-Being, and Power Dynamics	38
3.	Methodology	43
3.1	Research Design	43
3.1.1	Philosophical stance: Pragmatism.....	43
3.1.2	Time horizon: Cross-sectional survey.....	43
3.1.3	Strategy: Concurrent embedded mixed-methods	44
3.1.4	Alignment to research questions (RQ1–RQ4)	44
3.1.5	Scope of inference and threats to validity	44
3.1.6	Why this design and not alternatives.....	45
3.2	Population and Sampling.....	45
3.2.1	Target population & inclusion criteria	45
3.2.2	Sampling frame & technique	46
3.2.3	Achieved sample & composition	46
3.2.4	Representativeness & external validity	46
3.3	Survey Instrument Development.....	47
3.3.1	Theoretical foundations.....	47
3.3.2	Structure & response formats (by section).....	47
3.3.3	RQ–construct–item mapping table.....	48
3.3.4	Item wording & negatively phrased items	49
3.3.5	Pre-survey and refinement	49
3.3.6	Psychometric scope.....	49
3.4	Data Collection Procedure.....	49
3.4.1	Platform & timeline.....	49
3.4.2	Recruitment channels & message ethics	50
3.4.3	Consent gate & response rules	50
3.4.4	Data handling and cleaning	50
3.4.5	Quality safeguards.....	50
3.4.6	Collection-stage limitations	51
3.5	Data Analysis Approach.....	51
3.5.1	Analytical framework & tools.....	51

3.5.2	Quantitative analysis	52
3.5.3	Qualitative analysis	52
3.5.4	Mixed-methods integration	52
3.6	Ethical Considerations	53
3.6.1	Informed consent & participant rights	53
3.6.2	Privacy, anonymity, and data protection	53
3.6.3	Risk assessment and mitigation	53
3.6.4	Data security and retention.....	54
3.7	Limitations of Methodology	54
3.7.1	Sampling and external validity limitations.....	54
3.7.2	Data collection and measurement limitations	54
3.7.3	Implications and mitigation strategies.....	54
3.8	Linking methodology and empirical results	55
4.	Results and Discussion.....	56
4.1	RQ1 – Perceived Task-Level Benefits	58
4.1.1	Quantitative Results: Productivity and Efficiency Gains.....	58
4.1.2	Qualitative Insights: Employee Explanations and Examples.....	58
4.1.3	Discussion: AI as a Task Augmentation Tool.....	59
4.2	RQ2 – Perceived Job-Level Risks	60
4.2.1	Quantitative Results: Limited Personal Threat Perception	60
4.2.2	Qualitative Insights: Nuanced Concerns About Broader Impacts.....	61
4.2.3	Discussion: Experience and Familiarity as Risk Mitigators	62
4.3	RQ3 – Skill Development and Organizational Support	63
4.3.1	Quantitative Results: High Self-Efficacy Amid Training Gaps.....	63
4.3.2	Qualitative Insights: Calls for Training, Guidelines, and Inclusive Implementation	65
4.3.3	Discussion: Bridging the Support Gap for an AI-Ready Workforce	67
4.4	RQ4 – Ethical Concerns and Future Adoption Intentions	68
4.4.1	Quantitative Results: Caution Meets Optimism.....	68
4.4.2	Qualitative Insights: Transparency, Data Privacy, and Responsible Use	70
4.4.3	Discussion: Balancing Enthusiasm with Ethics and Trust	71
4.5	Overall Summary: A Paradox of Enthusiasm and Caution	73
5.	Discussion	75
5.1	Key Findings	75
5.2	Theoretical Contributions	77
5.2.1	Extending Technology Acceptance Models.....	77

5.2.2	Advancing Automation Anxiety Theory	77
5.2.3	Integrating Sociotechnical Systems Perspectives	78
5.3	Practical Recommendations	79
5.3.1	For Organizations and Managers	79
5.3.2	For Individual Employees	80
5.3.3	For Policymakers	80
5.3.4	Implementation roadmap (time horizons)	81
5.4	Limitations.....	81
5.5	Future Research Directions	82
5.5.1	Methodological Advances.....	82
5.5.2	Sample and Context Extensions	83
5.5.3	Emerging Questions	84
5.6	Overall synthesis and final reflections.....	85
6.	References	87
7.	Appendix A: Survey Instrument – Questionnaire	93

List of Tables

Table 1 - Mapping of Research Questions (RQs) to Constructs and Questionnaire Items (Q12–Q31).	48
Table 2 - “Key indicators at a glance” (Agree/Strongly agree unless noted).....	56

1. Introduction

1.1 Background

Artificial intelligence (AI) has rapidly permeated contemporary workplaces, reshaping how tasks are performed across a wide range of industries. Organizations increasingly deploy AI-enabled tools for data analysis, process automation, decision support, and customer service – propelled by the maturation of scalable, cloud-based AI services. A widely noted indicator of this acceleration is the generative-AI application ChatGPT, which reached an estimated 100 million users within two months of launch – reportedly the fastest user-base growth for a consumer application in history (Hu, 2023).

Such milestones exemplify the unprecedented pace at which AI capabilities are diffusing into everyday work and life. Although the information technology (IT) sector often acts as an early adopter – reflecting relatively high technical readiness and clearer pathways for integrating AI into digital workflows – AI’s uptake has become meaningfully cross-sectoral, influencing professional practice well beyond IT.

Alongside this rapid diffusion, vigorous debates have emerged about AI’s consequences for productivity and employment. Optimists contend that AI systems augment human labor by automating routine and repetitive tasks, thereby allowing employees to reallocate attention to higher-value, creative, and strategic activities (Brynjolfsson et al., 2023). Early empirical studies support this view. In a large-scale field experiment with thousands of customer service agents, access to an AI assistant increased productivity by roughly 14%, with the largest gains accruing to less-experienced workers (Brynjolfsson et al., 2023). Similarly, in a randomized experiment with knowledge-work writing tasks, professionals using ChatGPT completed work approximately 40% faster and produced higher-quality outputs (Noy & Zhang, 2023). Taken together, these findings position AI as a catalyst for efficiency, capable of amplifying human performance in much knowledge-intensive contexts.

Countervailing perspectives, however, emphasize risks. Early task-automation projections warned that AI and computerization could endanger large segments of employment. Frey & Osborne (2017) famously estimated that up to 47% of U.S. jobs were at risk of computerization over coming decades, galvanizing public concern about technological unemployment. Subsequent scholarship offers a more nuanced picture: rather than eliminating entire occupations, AI tends to reconfigure jobs by automating specific tasks while complementing and even creating other forms of work – thus shifting skill demands and elevating the importance of judgment, creativity, and interpersonal capabilities (Autor, 2015). Historical precedents – from mechanization to personal computing – show a recurring pattern in which technologies displace certain tasks even as they foster new roles and raise the skill threshold for many existing ones.

Nevertheless, credible evidence of disruption persists. For example, increased industrial-robot adoption has been linked to statistically significant declines in employment and wages in exposed local labor markets (Acemoglu & Restrepo, 2020). These mixed outcomes – productivity gains coupled with potential displacement – suggest that AI’s impact on work is inherently double-edged, with implications that depend heavily on organizational choices, implementation design, and the distribution of benefits and risks.

Against this backdrop, an employee-centric perspective is essential. Much of the public discourse still rests on macro-level forecasts or technical capability assessments, whereas the day-to-day experiences of workers – their perceived benefits, anxieties, and ethical concerns – remain comparatively underexamined. To address this gap, the present dissertation examines how professionals perceive AI’s expanding role in their work.

The empirical study draws on a cross-sectional survey of 75 professionals across multiple countries; while a large majority of respondents work in the IT sector (83%), the analysis throughout the dissertation is conducted at the level of the entire professional sample, not as an IT-only subgroup analysis. This composition provides an early-adopter vantage point while ensuring that the results are interpreted as the perceptions of professionals more broadly, with the IT concentration acknowledged as a limitation for external validity. As summarized in the Abstract and developed in Chapter 4, respondents report strong task-level benefits (e.g., time savings and perceived productivity gains) alongside notable ethical reservations – particularly around transparency, data privacy, and appropriate reliance – yielding a pattern of “cautious optimism” that frames the study’s research questions and subsequent discussion.

1.2 Motivation

Against this backdrop, the central motivation of this dissertation is to foreground the **employee perspective** on workplace AI – an angle still comparatively underexamined relative to macroeconomic forecasts and technical assessments. While public debate often extrapolates from model capabilities or firm-level adoption rates, far less is known about how professionals actually experience, interpret, and respond to AI in their day-to-day work.

This study responds to that gap by placing workers’ voices at the center of inquiry and by treating their perceptions as critical evidence for understanding AI’s real effects inside organizations.

Importantly, the motivation extends beyond any single occupation or sector. Although the information technology (IT) domain is frequently an early adopter of AI tools – and is well represented in the present sample – AI systems are reshaping tasks and workflows across many professional roles. The dissertation therefore examines professionals in general, using an empirically AI-exposed cohort to surface patterns that are relevant across sectors, while explicitly acknowledging the IT-predominant

composition of respondents as a boundary on generalizability. Throughout the thesis, analyses are conducted at the level of the full professional sample, not as an IT-only subgroup analysis.

The need for an employee-centered perspective is both practical and scholarly. Organizations risk mismanaging AI integration when they overlook how employees actually experience the technology – misjudging readiness, discounting legitimate concerns, or failing to provide the training and safeguards that transform potential into performance. In such cases, professionals may either resist or underutilize AI (e.g., developers declining to rely on coding assistants or analysts distrusting model outputs) or, conversely, over-rely on opaque recommendations without sufficient validation, thereby introducing new quality, ethical, and compliance risks. Both error patterns can erode productivity gains, amplify technostress, and generate avoidable workplace tensions. Consequently, attending to employees' perceptions is essential for fit-for-purpose adoption, effective change management, and the sustained realization of AI's organizational value.

A further motivation arises from mounting evidence of a paradox in lived experience: the same professionals who report tangible task-level benefits (time savings, quality improvements, fewer errors) often voice concerns about transparency, data protection, bias, and appropriate reliance. Rather than treating this as inconsistency to be ironed out, the dissertation treats it as a substantive phenomenon that requires explanation and has direct implications for governance and practice. By documenting both enthusiasm and caution, the research aims to clarify where AI is perceived to help, why reservations persist, and which organizational supports – training, guidelines, human-in-the-loop processes – employees believe would enable responsible, higher-trust use.

In sum, this study is motivated by the conviction that understanding professionals' perceptions is indispensable for credible assessment of AI's workplace impact and for designing interventions that work in practice. Using a cross-country sample of professionals – predominantly but not exclusively from IT – the dissertation seeks to illuminate how employees navigate AI's benefits and risks across four domains that structure the remainder of the work: task-level outcomes, job-level risks and enrichment, skill development and organizational support, and ethics/trust. This motivation aligns with the empirical approach and reporting adopted in Chapters 3–5, where findings are presented for the entire professional cohort while the IT concentration is transparently acknowledged as a scope condition.

1.3 Problem Statement

The rapid integration of artificial intelligence (AI) into workplaces has outpaced systematic evidence about its effects on the workforce – especially at the level of day-to-day employee experience. Despite extensive discussion of model capabilities and firm-level adoption, there remains limited clarity about how professionals actually experience, interpret, and respond to AI across the core domains of their work. In particular, prior accounts often emphasize technology supply (what AI can do) rather

than human reception (how AI is lived and managed), leaving a gap in understanding the micro-level consequences of AI for tasks, jobs, skills, and ethics.

This gap is consequential. Professionals – many of whom already use AI tools – report tangible task-level benefits (e.g., time savings, quality improvements, fewer errors), yet they may simultaneously voice ethical and organizational concerns about transparency, data protection, bias, and appropriate reliance. These apparently contradictory views produce a paradox in lived experience: employees can embrace AI's utility while remaining cautious about its risks. Untangling this paradox requires evidence that distinguishes between (a) near-term task augmentation and (b) broader job-level implications, including perceived security, stress, and enrichment, as well as (c) the readiness conditions that enable responsible use (skills, training, and managerial support). Without such evidence, organizations risk overestimating employee readiness, overlooking legitimate reservations, or failing to provide the guardrails that convert potential into performance.

A second source of ambiguity concerns scope of inference. Many studies focus on specific occupations or report managerial viewpoints; fewer center on professionals as a whole while explicitly acknowledging early-adopter contexts. Because the present study's sample is IT-predominant, there is a risk that readers infer an IT-only focus. However, the analytical aim is broader: to understand how professionals in general perceive AI's effects, using a cohort with high exposure as an informative window into emerging patterns. This requires a problem formulation that (i) treats IT concentration as a boundary on generalizability, and (ii) conducts all analyses at the level of the full professional sample, not as an IT-only subgroup comparison.

Accordingly, the problem this dissertation addresses is the lack of employee-centered, fine-grained evidence on how professionals – working in multiple countries and exposed to AI in their daily roles – perceive AI's impact across four interrelated domains that structure the remainder of the thesis: (1) task-level outcomes, (2) job-level risks and enrichment, (3) skill development and organizational support, and (4) ethics, transparency, and trust.

To address this, the study examines the perceptions of 75 professionals drawn from several countries, 83% of whom work in IT, and reports results for the entire sample. The design is exploratory and intentionally descriptive; given the convenience sampling, sector concentration, and international scope without country-level stratification, findings are positioned as context-bound patterns rather than generalizable estimates.

By centering professionals' voices and separating task-, job-, readiness-, and ethics-related perceptions, the dissertation seeks to clarify where AI is perceived to help, why reservations persist, and which organizational supports (training, guidelines, human-in-the-loop processes) employees believe would enable responsible, higher-trust use. This problem framing aligns with the empirical

strategy and reporting in Chapters 3–5, where analyses are conducted at the aggregate professional level, the IT predominance is disclosed as a scope condition, and the results are interpreted through the lens of cautious optimism – strong perceived task benefits tempered by ethical concerns and a call for fit-for-purpose organizational support.

1.4 Research Objectives and Questions

In light of the above problem, this research aims to explore and clarify how professionals perceive the introduction of AI in their workplace – encompassing its perceived advantages, potential drawbacks, and the contextual factors that influence these perceptions. The study’s objectives are designed to capture both the positive outcomes and negative implications of workplace AI, while also accounting for conditions such as training and ethics that may shape professionals’ views. Grounded in a survey of 75 professionals across multiple countries – 83% of whom work in the IT sector – the study pursues four objectives and corresponding research questions (RQ1–RQ4):

1. Evaluate task-level outcomes of AI use.

Examine how professionals perceive the effects of AI tools on their individual work tasks, including whether AI improves productivity, enhances output quality, and reduces the time required to complete routine activities.

RQ1: How do professionals perceive AI’s effects on their task-level productivity, output quality, and time savings?

2. Assess job-level risks associated with AI.

Determine the extent to which AI integration is linked to professionals’ perceptions of job insecurity and work-related stress – for example, concerns about losing roles or responsibilities to AI, or pressures arising from monitoring and validating AI outputs.

RQ2: To what extent is AI integration associated with perceived job insecurity and work-related stress among professionals?

3. Investigate skill development needs and organizational support.

Explore how AI adoption influences perceived needs for new skills or upskilling and evaluate the role of organizational support – particularly employer-provided AI training – in shaping these perceptions (e.g., whether training reduces anxiety or heightens awareness of skill gaps).

RQ3: How does AI adoption influence perceived upskilling needs among professionals, and how does organizational AI training shape these perceptions?

4. Identify ethical concerns and their influence on adoption intentions.

Identify which ethical issues (e.g., algorithmic bias, transparency, and data privacy or security) professionals associate with AI at work and analyze how the salience of these concerns relates to enthusiasm for – or resistance to – further AI adoption.

RQ4: What ethical concerns do professionals associate with workplace AI, and how are these concerns related to their intentions to adopt or expand AI tools?

Together, these four objectives and research questions offer a comprehensive framework for examining AI's human impact within organizations, particularly among professionals most familiar with the technology. They also illuminate a central paradox: that employees may simultaneously embrace AI's productivity gains while harboring ethical or employment-related concerns. Recognizing this duality clarifies both the promises and the pressures of AI integration in practice, highlighting the complex balance between technological advancement and human experience.

1.5 Research Methodology

To address the research questions, this study adopted a quantitative cross-sectional survey design implemented through a self-administered online questionnaire. The instrument was developed to capture professionals' perspectives across the key dimensions of interest – including perceptions of AI's effects on productivity and task efficiency, job-related risks or stress, needs for skill development and training availability, and ethical concerns regarding AI use.

The target population consisted of working professionals, with particular emphasis on the IT sector given its role as an early adopter of AI technologies. The questionnaire was distributed online through professional networks and ultimately yielded 75 complete responses – predominantly from the IT sector (83%) – providing insights from early adopters across multiple countries.

The data collected were analyzed using both quantitative and qualitative methods. Responses to the closed-ended Likert-scale items were examined using descriptive statistics to identify overarching trends and patterns in participants' perceptions. In parallel, open-ended responses were subjected to thematic coding to uncover recurring themes and generate deeper interpretive insights.

This mixed-methods approach allowed the study to quantify broad attitudes and experiences while simultaneously capturing the more nuanced reasoning expressed in participants' own words. Together, these complementary analyses provide a comprehensive and empirically grounded understanding of how professionals perceive AI's impact within their work environments.

1.6 Significance of the Study

From a practical perspective, the findings of this study can help organizations and policymakers implement artificial intelligence (AI) in ways that are responsible, effective, and human-centered. By illuminating professionals' lived experiences, concerns, and support needs, the research provides empirical evidence to guide adoption strategies that align with the realities of the workforce.

The international scope of the sample yields cross-border insights that are especially relevant for multinational organizations seeking consistent approaches to AI deployment. For example, if persistent fears or misconceptions are observed even among technically proficient professionals, managers can address them through targeted communication, role-specific training programs, and support initiatives (e.g., human-in-the-loop workflows, clear usage guidelines, and data-governance guardrails). Such actions can improve readiness, reduce preventable errors, and foster trust, thereby converting technical potential into durable performance gains.

Understanding how early adopters navigate AI integration also offers lessons for sectors beginning similar transformations. Incorporating the perspectives of those at the forefront of adoption increases the likelihood that AI initiatives deliver productivity and quality improvements without unintended adverse effects on morale, job security, or employee well-being. In practice, this means aligning AI rollout plans with employees' perceptions of value and risk, building the organizational supports (skills development, oversight mechanisms, and feedback channels) that employees themselves identify as necessary for responsible use.

From an academic perspective, the study advances understanding of AI's impact at the micro (individual) level. Much of the existing literature remains conceptual or emphasizes macroeconomic and firm-level outcomes; by contrast, this research contributes systematic, employee-centered evidence on the human side of AI integration. By examining professionals across multiple countries – predominantly but not exclusively from the IT sector – the study offers a vantage point from participants who are both creators and users of AI technologies. The identification of a paradox in lived experience – whereby professionals report tangible task-level benefits while simultaneously expressing ethical and organizational concerns – adds nuance to prevailing narratives of AI acceptance and informs debates in technology adoption, sociotechnical systems, and responsible AI.

By capturing and analyzing professionals' perceptions across four interrelated domains – task-level outcomes, job-level risks and enrichment, skill development and organizational support, and ethics/trust – the study introduces a structured lens for assessing AI's workplace effects. In doing so, it fills a documented gap in the literature and opens avenues for future research on how human factors mediate the success of technological innovation within organizations.

1.7 Structure of the Dissertation

This dissertation is organized to guide the reader from context and motivation through a focused review of prior knowledge, into the study's design, results, and implications. Each chapter serves a distinct purpose, and together they form a coherent line of inquiry: (i) establish why the topic matters and what is being asked; (ii) position the work within existing theory and evidence; (iii) describe how the research was conducted; (iv) present and interpret the findings in relation to the research questions (RQs); and (v) synthesize the study's contributions, limitations, and directions for future research.

The following section provides a detailed roadmap of the full document – including the front matter, five core chapters, references, and appendix – to help different audiences (researchers, practitioners, and policy-oriented readers) navigate the structure and evaluate the work's rigor and relevance.

1.7.1 Chapter 1: Introduction

Chapter 1 frames the problem and defines the scope of the inquiry. Section 1.1 situates artificial intelligence (AI) in contemporary workplaces and outlines both its promise – productivity, quality, and learning spillovers – and its challenges – task displacement, stress, and inequality. Section 1.2 clarifies the motivation for an employee-centric perspective, particularly among professionals who are both creators and users of AI tools, while Section 1.3 articulates the problem: organizational AI adoption has outpaced a detailed understanding of its human-level effects. Section 1.4 presents four research objectives and corresponding questions:

- **RQ1:** perceived task-level outcomes (productivity, time savings, work quality); **RQ2:** perceived job-level risks (insecurity, stress) and enrichment; **RQ3:** skills, upskilling, and organizational support; and **RQ4:** ethics, transparency, trust, and adoption intentions.

Section 1.5 summarizes the research methodology – an exploratory cross-sectional survey analyzed using descriptive statistics and thematic coding – while Section 1.6 explains the study's practical and academic significance.

The present Section 1.7 concludes the chapter by showing how each subsequent chapter addresses the RQs and where supporting materials can be found (e.g., the questionnaire in Appendix A). Readers seeking a concise overview may combine Chapter 1 with the synthesis in Chapter 5, while those focused on methods and instruments should read Chapter 3 alongside Appendix A.

1.7.2 Chapter 2: Literature Review

Chapter 2 establishes the theoretical and empirical foundations of the dissertation and justifies the constructs measured later in the survey instrument.

- **Section 2.1 – Theoretical Foundations.** This section reviews technology acceptance theories (TAM and UTAUT), performance alignment through the Task–Technology Fit model, organizational learning and adaptation theories (e.g., Social Cognitive Theory and JD-R), and emerging human–AI teaming frameworks (trust calibration, explainability, and division of labor). Collectively, these perspectives motivate the measurement of perceived usefulness, task fit, trust, transparency, and self-efficacy in the survey.
- **Section 2.2 – Empirical Evidence on AI’s Impact.** The review synthesizes field and experimental evidence on productivity effects – such as faster task completion and quality gains – contrasts augmentation and displacement dynamics, and summarizes exposure analyses indicating uneven, task-level automation rather than wholesale job loss in most occupations. It also highlights measurement approaches – including controlled experiments, exposure indices, and large-scale surveys – that inform the design of the study’s descriptive indicators.
- **Section 2.3 – Focused Evidence Base Motivating the RQs.** This section narrows the lens, mapping specific bodies of literature to RQ1–RQ4 (e.g., productivity studies for RQ1; technostress and insecurity for RQ2; training and organizational support for RQ3; and ethics and explainability for RQ4). This alignment ensures that the instrument operationalizes constructs with demonstrated relevance in prior research.
- **Section 2.4 – Workforce Transformation in the AI Era.** The review discusses how jobs are re-bundled around AI (augmentation vs. displacement), identifies human capabilities that complement AI (analytical, creative, collaborative, and adaptive skills), and anticipates new roles such as AI facilitators and auditors.
- **Section 2.5 – Ethical and Societal Implications.** This section examines algorithmic fairness, accountability, privacy and surveillance, worker autonomy, and power dynamics under algorithmic management – issues that underpin the study’s ethics and trust measures and later interpretation.

Researchers will find conceptual scaffolding for the constructs; practitioners can extract a concise overview of risks, opportunities, and governance priorities; and both groups will see why the study measures what it measures.

1.7.3 Chapter 3: Methodology

Chapter 3 documents the study's methodological design choices with transparency to support appraisal and replication.

- **Section 3.1 – Research Design.** This section outlines the study's pragmatist stance, a cross-sectional time horizon (mid-June to 31 July 2025), and a concurrent embedded mixed-methods strategy – a primarily quantitative, self-administered online survey with embedded open-ended prompts. It also details the alignment between the research questions (RQs) and item blocks (e.g., Q12–Q17 for RQ1; Q22–Q25 for RQ2; Q27–Q30 for RQ3; and Q18–Q21 plus Q26/Q31 for RQ4).
- **Section 3.2 – Population and Sampling.** The target population comprises employed adults. Recruitment used convenience and snowball sampling across professional networks, producing $N = 75$ complete responses (83% from the IT sector). The section discusses representativeness and scope of inference.
- **Section 3.3 – Survey Instrument Development.** This section describes theory-driven item construction, including the RQ–construct–item mapping and the inclusion of a “Not sure” option to minimize forced responses. A brief pilot informed refinement. Table 1 summarizes the mapping for ease of reference.
- **Section 3.4 – Data Collection Procedure.** Platform selection, consent gating, recruitment messages, data handling and cleaning, and quality safeguards – such as required closed-ended items and negatively phrased statements to reduce acquiescence bias – are specified.
- **Section 3.5 – Data Analysis Approach.** Quantitative analysis is descriptive (frequencies, percentages, and item means excluding “Not sure” responses). Qualitative analysis employs inductive thematic coding, and integration follows a convergent-triangulation logic in Chapter 4 (QUAN → QUAL → integrated discussion for each RQ).
- **Section 3.6 – Ethical Considerations.** This section documents informed consent procedures, anonymity assurances (no personal identifiers stored), risk assessment and mitigation (topic sensitivity addressed through neutral wording and optional text responses), and data retention/deletion policies.
- **Section 3.7 – Limitations of Methodology.** The chapter concludes with a candid discussion of sampling constraints, self-report and cross-sectional limitations, and design trade-offs – setting appropriate expectations for the interpretation of Chapter 4.

Method-minded readers can assess internal and external validity boundaries, while practitioners gain clarity on how the “**voice of the employee**” was captured and analyzed.

1.7.4 Chapter 4: Results and Discussion

Chapter 4 integrates results with interpretation, structured explicitly around the research questions. Each RQ is addressed in three parts: (i) quantitative results, (ii) qualitative insights, and (iii) an integrated discussion connecting the findings back to the literature and theoretical frameworks introduced in Chapter 2. This design allows immediate sense-making of the evidence rather than a deferred discussion.

- **Section 4.1 (RQ1)** – Reports perceived task-level benefits (productivity, quality, and time savings) and uses open-ended examples to illustrate how AI offloads repetitive tasks and reduces errors, interpreting these patterns through TAM and TTF (usefulness and fit) and recent field evidence.
- **Section 4.2 (RQ2)** – Details perceived job-level risks (insecurity and technostress) alongside enrichment through higher-value task focus. The discussion examines how familiarity and self-efficacy may buffer anxiety while recognizing the broader labor-market concerns respondents noted.
- **Section 4.3 (RQ3)** – Presents self-efficacy and upskilling intent against a backdrop of uneven formal training and organizational support. Qualitative suggestions – including training, guardrails, and inclusive change management – are analyzed through the JD-R framework and organizational-learning theories.
- **Section 4.4 (RQ4)** – Highlights the “cautious optimism” characterizing respondents: strong enthusiasm and adoption intent tempered by concerns about transparency, privacy, and fairness. The discussion links explainability, governance, and trust calibration to sustainable use.
- **Section 4.5** – Concludes the chapter with an overall synthesis, emphasizing the confidence–concern paradox as a defining feature of early adopters’ perceptions.

Readers seeking actionable insights can refer directly to the relevant RQ section(s). The integrated structure ensures that each set of findings is immediately contextualized, minimizing the cognitive load of cross-referencing across chapters.

1.7.5 Chapter 5: Discussion

This chapter discusses the main findings in light of the literature, outlines theoretical contributions, practical recommendations, limitations, and future research directions, and provides an overall synthesis and final reflections.

- **Section 5.1 – Key Findings.** Presents a concise summary of answers to RQ1–RQ4, emphasizing the prevalence of perceived task-level gains, limited personal threat perceptions, strong self-efficacy and upskilling motivation amid training gaps, and the ethical concerns that temper – but do not derail – adoption intentions.
- **Section 5.2 – Theoretical Contributions.** Explains how the findings extend technology-acceptance models for intelligent systems (by incorporating trust, transparency, and technostress), refine automation-anxiety theory (the coexistence of enthusiasm and concern), and reinforce sociotechnical perspectives on AI integration.
- **Section 5.3 – Practical Recommendations.** Provides guidance organized by stakeholder group: **Organizations and managers:** invest in continuous training, establish ethical AI frameworks, and promote human–AI collaboration; **Employees:** embrace lifelong learning, engage constructively with rollouts, and advocate for ethical practice; **Policymakers:** support workforce development, set guardrails for responsible AI, and fund independent impact research.
- **Section 5.4 – Limitations.** Restates the study’s boundaries – self-report data, cross-sectional design, non-probability sampling, and an evolving context – to avoid overgeneralization.
- **Section 5.5 – Future Research Directions.** Proposes longitudinal and experimental designs, broader and comparative sampling (across sectors, firm sizes, and countries), and new lines of inquiry addressing long-term career trajectories, organizational outcomes under varying implementation strategies, and societal-level impacts.
- **Section 5.6 – Final Reflections.** Offers a closing reflection underscoring the human-centered imperative: successful AI integration depends as much on leadership, ethics, and learning cultures as on technical capability.

Executives and practitioners pressed for time can pair Section 5.1 (*what we found*) with Section 5.3 (*what to do*), while scholars can focus on Section 5.2 and Section 5.5 for theoretical and methodological agendas.

1.7.6 References

A comprehensive reference list follows Chapter 5. It consolidates the theoretical, empirical, and methodological sources cited throughout the dissertation – spanning foundational acceptance models and sociotechnical theories to recent field experiments and governance frameworks. This section enables readers to verify claims, trace conceptual lineages, and pursue deeper exploration of specific topics such as explainability, algorithmic auditing, and the JD-R model.

1.7.7 Appendix A: Survey Instrument – Questionnaire

Appendix A reproduces the full instrument used in the study, including consent language, closed-ended items (Sections A–C), and open-ended prompts (Section D). It documents response options, the inclusion of a “Not sure” category on every Likert item, and the precise wording for constructs associated with each research question (e.g., Q12–Q16 for task-level outcomes; Q17–Q21 plus Q26 for fit, trust, transparency, and ethics; Q22–Q25 for job-level risks and enrichment; and Q27–Q30 for readiness and organizational support, with Q31 (adoption intention) grouped under RQ4). Providing the complete instrument ensures transparency, replicability, and critical appraisal of measurement choices – for instance, the inclusion of negatively phrased items to reduce acquiescence bias.

2. Literature Review

Advances in artificial intelligence (AI) are rapidly transforming the nature of work, sparking widespread debate about the future of jobs. Optimists emphasize AI's potential to augment human productivity – automating tedious tasks and creating new opportunities – while pessimists warn of job displacement and economic disruption. This literature review provides a comprehensive overview of current knowledge on AI's impact on employment, with a particular focus on individuals' perceptions of these changes. It draws on foundational theories of technology adoption and organizational change, synthesizes empirical evidence on AI's effects on productivity and employment, examines how organizations are implementing AI, and explores the workforce transformations currently underway. Key ethical and societal implications are then discussed, followed by an identification of gaps in the literature and emerging opportunities for research. Through this review, the groundwork is laid for understanding how AI is reshaping work – and how both workers and organizations can navigate this jagged technological frontier.

2.1 Theoretical Foundations

A robust theoretical foundation is essential for analyzing how AI adoption influences both workers and organizations. The relevant theoretical frameworks span individual technology acceptance, organizational change and learning, and emerging models of human–AI collaboration.

2.1.1 Technology Acceptance Theories

The Technology Acceptance Model (TAM), proposed by Davis (1989), is one of the earliest and most influential frameworks for explaining user adoption of technology. The model posits that two key beliefs – **perceived usefulness (PU)**, the extent to which using a technology improves job performance, and **perceived ease of use (PEOU)**, the extent to which its use is free of effort – jointly determine an individual's **intention to adopt** a new technology. Davis demonstrated that these factors significantly explain technology-use behavior, with *perceived usefulness* emerging as the stronger predictor of acceptance (Davis, 1989). TAM has been widely applied to understand employees' adoption of AI-enabled tools, based on the premise that when an AI system is perceived as useful for one's job and not overly difficult to use, individuals are more likely to embrace it.

Venkatesh et al. (2003) integrated the Technology Acceptance Model (TAM) with several complementary frameworks to develop the Unified Theory of Acceptance and Use of Technology (UTAUT). This model identifies four key determinants of technology acceptance: performance expectancy (the anticipated benefit to job performance), effort expectancy (perceived ease of use), social influence (peer or supervisory encouragement or pressure), and facilitating conditions (organizational and technical infrastructure support). UTAUT also incorporates moderating variables,

such as age, gender, experience, and voluntariness of use, to explain differences in adoption behavior across user groups.

In the context of AI adoption, the UTAUT framework suggests that an employee's decision to use AI tools depends not only on the technology's perceived usefulness but also on the social and organizational environment surrounding its implementation. Adoption is influenced by whether leaders and colleagues encourage its use and whether adequate training and IT support are available. For instance, if an organization introduces an AI coding assistant, software developers are more likely to adopt it when they believe it enhances productivity (performance expectancy), find it user-friendly (effort expectancy), observe peers successfully using it (social influence), and have access to sufficient documentation or technical assistance (facilitating conditions).

In UTAUT terms, this study's instrument captures performance expectancy through items Q12–Q16 and facilitating conditions through items Q28–Q29. Social influence and effort expectancy were not directly measured and are therefore acknowledged as scope limitations of the design.

Goodhue & Thompson (1995) introduced the Task–Technology Fit (TTF) theory as a complementary lens for understanding technology adoption and performance. The theory posits that a technology will positively influence individual performance only when its capabilities closely align with the tasks the user must perform – meaning the technology effectively supports the user's task requirements. In essence, even a highly sophisticated AI tool must fit the job's specific needs to be genuinely effective.

For example, an AI-powered customer service chatbot can significantly enhance an agent's performance if it is well suited to handling the types of inquiries the agent typically encounters. Conversely, the same tool could hinder performance if applied to tasks beyond its intended scope. TTF therefore underscores the importance of evaluating AI–job alignment, ensuring that AI systems are implemented in contexts where their strengths – such as rapid data processing, pattern recognition, and language generation – complement human tasks rather than misalign with them. Without such fit, both adoption and performance outcomes are likely to remain limited.

This study examines these concepts through items Q12–Q16 (perceived usefulness, quality, and time gains) and Q17 (task–technology fit).

2.1.2 Organizational Change and Learning Theories

Bandura (1977) provides key insights into how individuals learn and adapt to new technologies such as artificial intelligence (AI). A central construct of SCT is self-efficacy – an individual's belief in their capability to successfully perform a task or master a new skill. Employees with higher AI-related self-efficacy (confidence in using AI tools) are more likely to experiment with, adopt, and effectively integrate AI into their workflows.

SCT also highlights learning through observation and social experience: in workplace settings, observing peers successfully collaborating with AI can increase others' confidence and willingness to engage, creating a virtuous cycle of learning and adoption. Conversely, low self-efficacy or high anxiety can hinder acceptance – when employees fear they cannot keep pace with AI or worry about being replaced, they may resist engaging with the technology. These dynamics underscore the importance of training, peer modeling, and supportive leadership to build confidence and sustain positive attitudes toward AI use.

Recent research supports these dynamics. Kim & Lee, (2024) found that AI adoption can increase employee stress and burnout unless employees possess high self-efficacy for learning AI, which buffers the stress effect. In practice, organizations should cultivate supportive learning environments that strengthen employees' belief in their ability to work effectively with AI – thereby reducing technostress and improving adaptation.

The adoption and sustained use of new technologies also depend on self-efficacy and access to organizational resources, such as training, managerial support, and adequate infrastructure. The JD-R model (Bakker & Demerouti, 2017) posits that job resources buffer strain and foster engagement, whereas job demands – such as workload, complexity, or technological pressure – contribute to burnout.

In the context of AI adoption, organizational training and managerial support operate as critical resources that enhance employee confidence, reduce uncertainty, and mitigate technostress. Consistent with Social Cognitive Theory (SCT), self-efficacy is represented in this study by item Q27, while related organizational supports are captured by items Q28–Q29. Observational learning, although central to SCT, was not directly measured in the present instrument.

2.1.3 Human–AI Collaboration Frameworks

Classical technology-adoption theories generally conceptualize technology as a tool and the human as its user. However, the emergence of advanced artificial intelligence (AI) – capable of functioning as a semi-autonomous agent – necessitates new frameworks that capture the dynamics of human–AI interaction and collaboration. Researchers and industry practitioners are increasingly developing models that describe how humans and AI systems work together, share decision-making, and divide tasks to optimize collective performance. Several emerging concepts are particularly noteworthy:

A recurring theme in human–AI collaboration is *trust calibration* – users must develop an appropriate level of trust in AI systems. Too little trust can result in underutilizing a valuable tool, whereas too much trust increases the risk of misuse or error – for example, accepting AI outputs without verification, which can be problematic if the AI occasionally “hallucinates” false information.

Frameworks for human–AI interaction emphasize the importance of transparency and explainability to support calibrated trust.

Another critical dimension of collaboration involves determining the division of labor between humans and AI systems. Early frameworks propose that tasks should be allocated according to which agent – human or machine – performs them more effectively, echoing the long-standing principle that *humans are better at X, machines are better at Y*. For instance, creative strategy, empathetic customer service, and complex problem definition typically remain human-led, whereas data analysis, routine optimization, and pattern recognition are better suited to AI.

Many organizations now design workflows in which AI performs the initial drafts or analyses and humans make the final decisions – a *human-in-the-loop* approach that ensures oversight and accountability. The Task–Technology Fit (TTF) concept extends to this context: maximizing AI’s value requires assigning it tasks aligned with its strengths while redesigning human roles to emphasize the areas AI cannot replicate, such as high-level judgment, contextual reasoning, and emotional intelligence (Goodhue & Thompson, 1995).

In summary, the theoretical foundation for this study spans classic models of individual technology adoption (TAM and UTAUT) and performance impact (TTF) through broader theories of innovation diffusion and emerging paradigms for collaboration with “smart” machines. Collectively, these lenses suggest that successful integration of AI into the workplace requires more than simply introducing new technologies – it demands alignment between technology and users’ tasks and skills, cultivation of confidence and trust, provision of social and organizational support, and a fundamental rethinking of workflows to optimally combine human and AI strengths.

Trust, transparency, and comprehensibility are captured by items Q18–Q21, which address RQ4.

2.2 Empirical Evidence on AI’s Impact

A growing body of empirical research – spanning controlled experiments, industry case studies, and macro-level analyses – has begun to quantify how AI influences productivity, job performance, and employment patterns. This section reviews the evidence in three parts: (1) productivity and efficiency gains at the task or firm level, (2) impacts on jobs and employment across occupations and industries, and (3) the key methodological approaches used to measure AI’s effects.

2.2.1 Productivity Studies

Early experimental evidence demonstrates notable productivity gains when workers are provided with AI assistance – particularly in knowledge-intensive tasks. Noy & Zhang, (2023) conducted a randomized controlled trial involving 453 college-educated professionals (e.g., marketers, grant writers, and consultants) who completed writing tasks representative of their jobs – drafting

emails, letters, and reports – either with or without access to *ChatGPT*. The results were striking on average, participants using ChatGPT completed tasks 40 % faster, and the quality of their output, as rated by blind evaluators, was 18 % higher.

Importantly, AI assistance benefited lower-skilled writers the most – narrowing the performance gap between weaker and stronger participants as *ChatGPT* acted as an equalizing tool that elevated those who struggled most. These findings suggest that generative AI can enhance both efficiency and equity in performance by supporting less experienced workers. However, the study also cautioned that for tasks requiring factual accuracy and contextual understanding – which *ChatGPT* lacks unless provided – human judgment remains essential. This reinforces the view of AI as a *complementary* rather than a replacement technology for complex knowledge work.

A similar large-scale experiment examined AI's effects on realistic business tasks. In a pre-registered field study, 758 consultants were assigned 18 different knowledge-intensive tasks (e.g., strategic analysis, creative slide writing) either with or without *GPT-4* assistance (Dell'Acqua et al., 2023). The findings closely mirrored the laboratory results: across tasks, consultants using AI completed 12.2 % more tasks on average, worked 25.1 % faster, and produced outputs rated 40 % higher in quality by expert evaluators.

Both high-performing and lower-performing individuals benefited, though the relative gains were largest among those initially less experienced – participants below the median performance increased their output by 43% with AI, whereas those above the median improved by 17%. This “productivity paradox” – where novices gain the most – likely occurs because the AI system, trained on vast amounts of domain knowledge, provides less experienced workers with tacit expertise that normally develops over years of practice. Interestingly, the study also included a task deliberately designed to fall outside AI's capability frontier; on that task, the AI-assisted group underperformed, being 19 percentage points less likely to reach the correct solution.

These results illustrate that when AI is applied to problems beyond its competence, it can mislead users and diminish performance. This finding underscores the importance of recognizing the boundaries of current AI capabilities – human experts must discern when reliance on AI is inappropriate and exercise critical oversight. Overall, the evidence demonstrates that in many knowledge-intensive tasks, AI augmentation can deliver substantial productivity gains, supporting theoretical perspectives such as Task–Technology Fit (TTF) – where task selection aligns with AI's strengths – and reinforcing the ongoing necessity of human judgment and oversight in areas that remain beyond AI's reliable scope.

Beyond short-term experiments, evidence from real-world business implementations is now emerging. Brynjolfsson et al., (2023) analyzed the rollout of a generative-AI assistance tool in a large customer contact center at a Fortune 500 software firm. The AI model was integrated directly into the workflow of support agents, providing suggested responses and recommendations during live chat

interactions. The study – aptly titled *Generative AI at Work* – found that AI assistance increased customer service agents’ productivity by an average of 14%. The largest gains occurred among junior and less experienced agents: the AI system effectively disseminated the expertise of elite performers to newer agents by suggesting proven solutions and conversational strategies. As a result, these agents not only handled more inquiries per hour but also achieved higher customer satisfaction scores on AI-assisted interactions (Brynjolfsson et al., 2023).

The findings suggest that the AI tool did more than simply accelerate task completion – it also enhanced service quality by standardizing best practices across agents. The qualitative analysis further indicated that the AI functioned as a real-time coach, capturing tacit organizational knowledge – such as common troubleshooting procedures, polite phrasing, and opportunities for upselling – and delivering this expertise to agents during live interactions. Experienced agents, who already possessed such knowledge, exhibited smaller performance gains, whereas novice agents improved substantially, thereby narrowing the organization’s internal skill gap.

Overall, the empirical evidence consistently portrays artificial intelligence (AI) as a powerful productivity enhancer in tasks involving information processing, writing, coding, and customer interaction. Additional studies report comparable effects across diverse professional domains. For example, Ziegler et al., (2022) observed that software developers using AI coding assistants such as GitHub Copilot – powered by an earlier GPT model – completed coding tasks up to twice as fast as those without AI support. Similarly, Korinek (2023) estimated that economists employing large language models for analytical work could be 10–20 % more productive than their counterparts working without AI tools.

These improvements are not uniform across all types of work – they tend to be most pronounced in tasks that can be well defined and partially automated by current AI capabilities, such as drafting text, writing code, or responding to knowledge-based queries. Creative, ambiguous, or physically grounded tasks show far smaller gains. Even in fields like healthcare, early evidence suggests that AI can save physicians time on documentation and offer diagnostic recommendations, but it cannot – and should not – make independent clinical judgments.

It is important to note that productivity gains from AI often depend on complementary organizational changes to be fully realized. The call-center study discussed earlier aligns with the concept of the productivity J-curve: initially, AI tools may be integrated into existing processes, yielding only modest improvements, while larger benefits emerge when workflows are redesigned to fully leverage AI capabilities – and when employees are trained to collaborate effectively with these systems. For instance, if report writing is accelerated by ChatGPT, organizations may reallocate human effort toward higher-value analytical work rather than simply increasing the quantity of reports produced.

This pattern echoes Brynjolfsson’s earlier finding that general-purpose technologies require complementary investments – in skills, process redesign, and organizational learning – to translate technological potential into measurable productivity growth (Brynjolfsson et al., 2018). In summary, the empirical evidence is highly encouraging: generative AI and related tools can substantially enhance worker productivity, particularly for routine cognitive tasks. However, realizing these gains at scale depends on both the breadth of adoption and the depth of organizational adaptation – themes explored further in Section 2.3.

2.2.2 Employment Impact Studies

A central question concerns how AI will influence overall employment – specifically, which jobs will be created, transformed, or eliminated. Recent studies have sought to quantify AI’s impact on occupations and job tasks across the economy, with particular attention to the rise of powerful generative-AI systems such as large language models. Although it remains too early for comprehensive labor-market effects to appear in official statistics, researchers are increasingly employing task-level analyses and occupational data to estimate the degree of exposure to AI.

One influential approach assesses what proportion of tasks within each occupation could be performed – or significantly assisted – by current AI technologies. Eloundou et al. (2023) introduced a rubric to evaluate large-language-model (LLM) capability alignment with job tasks, combining judgments from both human subject-matter experts and *GPT-4* itself. Their analysis estimated that approximately 80 percent of the U.S. workforce could have at least 10 percent of their work tasks affected by LLMs, and about 19 percent of workers might see 50 percent or more of their tasks impacted by LLM-powered software.

In practical terms, AI is relevant to a broad share of occupations – roughly four out of five workers will experience some degree of impact – but only about one-fifth are in roles where half or more of their activities are highly susceptible to automation or augmentation by current generative-AI tools. These effects span wage levels: intriguingly, some higher-wage, highly educated occupations show greater exposure because they involve intensive information processing (e.g., analysts, programmers), whereas many lower-wage manual jobs are less directly affected. This challenges the long-held assumption that only routine, low-skill work is at risk; LLMs represent a form of *cognitive automation* that increasingly influences white-collar roles (Eloundou et al., 2023).

Felten et al. (2023) expanded this line of research by developing an *AI Occupational Exposure Index* and found that some of the most exposed occupations to language-modeling AI include telemarketers and various postsecondary teaching roles (e.g., in languages and history). The industries with the highest predicted exposure to generative AI include legal services – due to intensive document and research work – and financial services, particularly securities and investment. Interestingly, they also observed a positive correlation between an occupation’s average wage and its exposure to AI,

suggesting that higher-paid knowledge roles contain more tasks that AI can augment or perform, while many low-paid service jobs (e.g., cleaning, food preparation, and other manual-dexterity roles) remain relatively unaffected for now.

These findings raise complex distributional questions: if AI primarily enhances high-skill workers' productivity, it could exacerbate income disparities unless reskilling initiatives and complementary job creation offset these effects.

A central debate concerns whether AI will primarily *augment* human workers – making them more productive in existing roles – or *automate* entire jobs, potentially leading to displacement. The emerging consensus from recent analyses is that augmentation remains the dominant effect under current AI capabilities, although certain job categories face more significant automation risks.

A comprehensive global analysis by the International Labour Organization (ILO) (Gmyrek et al., 2023) examined the exposure of occupational tasks to generative AI across countries. Even under an “upper-bound” scenario, the authors estimated that no more than approximately 5.5 percent of total employment in high-income countries is potentially exposed to full automation by generative AI, while in low-income countries the figure is roughly 0.4 percent – largely reflecting occupational differences. The only broad category identified as highly exposed was clerical work, for which the study estimated that about 24 percent of tasks are highly exposed to being automated by AI, with another 58 percent moderately exposed. For nearly all other occupational groups – professional, technical, and service roles – fewer than 5 percent of tasks were deemed highly automatable. In contrast, medium-level exposure (tasks that AI could perform with human oversight or accelerate through assistance) was more common, reaching up to 25 percent of tasks in some occupations.

These findings support the view that most jobs will experience *partial task automation* – in other words, augmentation rather than full substitution. As the ILO report summarizes, “the most important impact of the technology is likely to be augmenting work – automating some tasks within an occupation while leaving time for other duties – as opposed to fully automating occupations” (Gmyrek et al., 2023). In practical terms, a marketing analyst's job might evolve so that AI generates preliminary campaign reports or first drafts, allowing the analyst to focus on strategy and creative design. Many current examples reflect this pattern: AI “co-pilots” assist doctors in drafting notes (while doctors continue to diagnose and treat), and AI research tools help lawyers with initial case preparation (while lawyers still argue and decide cases).

Of course, some jobs will be eliminated. Technology-driven redundancy tends to affect roles with a high proportion of automatable tasks and limited requirements for human judgment or interpersonal skill. Earlier automation waves saw secretarial pools contract with the rise of personal computers, and assembly-line employment decline following the introduction of industrial robots. With

AI, positions such as telemarketing or entry-level drafting and revision roles – including copywriting, transcription, or basic programming – may face reductions as these functions are increasingly handled by AI systems.

At the same time, new categories of employment are emerging – for example, AI model trainers, *prompt engineers* (specialists who craft effective inputs for AI systems), and expanding opportunities in data science and AI system maintenance. Historically, technology has often created as many jobs as it displaces over the long run, though not necessarily for the same individuals – a dynamic emphasized by Autor (2015). The World Economic Forum, (2023) projects significant structural churn in the global labor market, anticipating that by 2027 approximately 69 million new jobs will be created and about 83 million will be eliminated.

Growth is expected in data-analysis, AI, and sustainability-related roles, while clerical and routine positions are projected to decline. Although this short-term forecast implies a net negative balance, it does not represent a catastrophic loss – rather, it highlights the urgent need for large-scale reskilling and worker transition initiatives from declining occupations into expanding ones.

It is also instructive to consider how workers themselves perceive AI's effects on their job security and workload. Large-scale surveys reveal a complex blend of optimism and anxiety. According to the ADP Research Institute's, (2024) global survey of 35,000 employees, 85 percent believe AI will affect their job in some way in the near future – yet they are divided on the nature of that impact. Forty-five percent expect that AI will help them perform their work more effectively (augmentation), while 42 percent worry it will replace aspects of their role (automation or displacement).

Notably, younger workers – particularly members of Generation Z – are the most convinced that AI will transform their work and are actively preparing for this change, yet they also express heightened concern that AI could replace parts of their function. This ambivalence underscores a growing gap between the promise of augmentation emphasized by experts and the fear of automation experienced by many employees. These perceptions matter when workers view AI primarily as a threat, they may resist its implementation or experience diminished morale and increased psychological strain – issues further explored in Section 2.5 on ethical implications.

In summary, empirical employment-impact studies indicate broad but uneven exposure of jobs to AI. Generative AI now affects a wide range of occupations, reshaping the task composition of many roles. In the near term, however, the trend points more toward job transformation than wholesale job loss. The central challenge lies in managing this transition – ensuring workforce upskilling, mobility, and psychological preparedness – so that the potential for augmentation translates into genuine productivity and economic growth without leaving segments of the workforce behind.

The literature highlights the importance of proactive strategies to mitigate displacement – particularly for occupations such as clerical work that are highly automatable – and to promote augmentation by helping workers integrate AI effectively into their roles. Many scholars advocate for forward-looking measures in education, social safety nets, and job redesign, implemented before rather than after large-scale disruption occurs.

This discussion directly motivates the focus of RQ2, which examines perceived job insecurity and stress (items Q22–Q23). Understanding how professionals perceive these risks provides insight into organizational readiness and the remaining gaps in capturing AI’s benefits while safeguarding workers’ well-being.

2.2.3 Measurement Approaches in AI Impact Research

This section aligns the reviewed literatures with the study’s research questions and survey measures (Q12–Q31). Investigating AI’s effects on work is inherently complex, and scholars have employed diverse methodologies to capture changes in productivity and employment. Understanding these approaches is essential for interpreting findings and recognizing the limitations of existing research.

One major approach is experimental – randomly assigning AI tools to some participants and withholding them from others to measure causal effects. Experiments offer high internal validity, as they can directly demonstrate that a specific AI tool caused a measurable productivity to increase under defined conditions. However, they are often short-term and task-specific and may not capture longer-term adaptation effects – for instance, whether productivity gains persist or diminish as the novelty wears off – or the broader disruptions that AI integration may introduce to workflows. Moreover, such studies typically assess *micro-productivity* indicators such as task completion rates, time efficiency, or output quality. Scaling these findings to firm-level or economy-wide productivity requires additional assumptions.

Nonetheless, experimental designs have been crucial in establishing early causal evidence of AI’s tangible benefits and in identifying where human–AI complementarities are strongest – for example, demonstrating that weaker writers or less experienced employees benefit more from AI assistance, while AI systems tend to falter on tasks that require contextual understanding or judgment.

Researchers assess productivity gains in multiple ways. Common indicators include task completion time, output quantity (e.g., number of cases handled, or lines of code written), and output quality (typically evaluated by experts or reflected in customer satisfaction ratings). In Brynjolfsson et al. (2023), for example, productivity was measured as issues resolved per hour by customer service agents and quality through customer-sentiment scores. Similarly, Noy & Zhang (2023) assessed time taken per task and used blind evaluations of output quality.

Quality itself is often multi-dimensional – in writing tasks, for instance, evaluators may rate clarity, persuasiveness, grammatical accuracy, and tone. Some studies also track *error rates* to determine whether AI assistance introduces mistakes, such as factual inaccuracies. A particularly innovative measure concerns *inequality of performance*. Noy & Zhang (2023) calculated the variance in participants’ outputs and found that it narrowed when AI assistance was introduced – a novel way to quantify AI’s equalizing effect on productivity outcomes.

To assess employment impacts, researchers rely extensively on task-analysis frameworks. A pioneering method by Frey & Osborne (2017) involved subject-matter experts labeling occupations as automatable or not. Newer approaches decompose jobs into specific tasks or skills using occupational databases such as O*NET, evaluating which of these tasks fall within AI’s current capabilities.

Felten et al. (2023) advanced this methodology by mapping 52 “abilities” – such as oral comprehension and inductive reasoning – to AI progress metrics and then aggregating these measures at the occupational level. Their study also introduced a novel element: using *GPT-4* to help rate tasks alongside human experts, effectively asking the AI, “*Could you do this task?*” for thousands of job components. This hybrid approach is innovative but raises methodological concerns – AI may overestimate its own competence or lack sufficient context to judge task complexity. Felten et al., (2023) mitigated these risks by combining AI self-assessment with expert human review to finalize exposure scores.

The ILO’s Gmyrek et al., (2023) approach similarly employed *GPT-4*, prompting it directly with task descriptions to classify each task as *automatable*, *augmentable*, or *unaffected*, which were then aggregated by occupation. Each method, however, presents distinct limitations: expert assessments can be subjective and time-consuming; AI-based evaluations may overlook nuance or change as AI systems evolve; and O*NET task inventories may be incomplete or outdated in rapidly evolving job categories.

Despite these methodological differences, most studies converge on a common pattern – very few occupations are fully automatable, yet many contain subsets of tasks that are amenable to AI assistance. It is also important to note that these represent *potential* effects. Actual outcomes depend on how organizations choose to deploy AI. For example, an occupation might be 50 percent automatable, but employers may opt to retain human oversight for quality assurance, ethical accountability, or customer trust reasons – or instead use AI to augment human productivity, such as enabling translators to complete more projects without reducing headcount.

Other researchers examine AI’s impact through surveys that capture adoption rates and perceptions. McKinsey’s annual *State of AI* survey (McKinsey & Company, 2023) reports that approximately 50–55 percent of companies use AI in at least one business function, while also documenting perceived barriers and impacts. Workforce-level surveys provide complementary insights:

the ADP Research Institute, (2024) found that 85 percent of employees expect AI to affect their jobs in some way, as noted earlier in Section 2.2.2.

These data reveal how ready and receptive the workforce is – an important leading indicator of AI’s real productivity potential. If two-thirds of employees have not yet used AI in their daily work, as suggested by Slack’s *Fall 2024 Workforce Index* (2024), then many productivity forecasts remain contingent on adoption that has yet to occur. Comparing management and employee perspectives further illuminates this gap: Slack’s survey reported that while 95 percent of executives were eager to integrate AI, roughly two-thirds of workers had still not used AI tools at work by late 2023 (Slack, 2024). This discrepancy points to an implementation lag and communication gap that quantitative task analyses alone cannot capture.

In sum, researchers employ a mix of quantitative and qualitative approaches: controlled experiments to identify causal effects on narrowly defined tasks, data-science techniques and modeling to estimate broader economic impacts, and surveys to capture the human dimension of AI adoption. Each approach has distinct strengths and limitations. A clear trend in the literature is the growing emphasis on interdisciplinary methods – for example, combining computational modeling with sociological insight – acknowledging that AI’s impact is not purely a technical question of productivity but also a social process of adaptation.

Caution is warranted when interpreting results. Many of the reported productivity gains are highly context-specific, and exposure analyses depend on modeling assumptions that may not hold universally. This literature review finds that more longitudinal and real-world studies are needed to determine whether short-term productivity boosts persist and to observe how job roles evolve in practice – for instance, whether workers genuinely shift toward more advanced tasks following AI adoption, thereby realizing the promise of augmentation.

Additionally, methodological challenges remain, such as accurately measuring the quality of AI-augmented work or accounting for the creation of entirely new tasks and jobs resulting from AI – factors that task-based models often fail to capture.

2.3 Focused Evidence Base Motivating the Research Questions

2.3.1 Productivity and Work Quality (RQ1)

Across knowledge-work settings, recent randomized and quasi-experimental studies demonstrate that access to generative-AI tools can increase individual productivity and, in many cases, improve output quality. In customer support, the introduction of a conversational assistant increased issues resolved per hour by approximately 14 percent on average – with the largest gains among less experienced agents – providing evidence that AI can disseminate organizational know-how and compress learning curves (Brynjolfsson et al., 2023). In a controlled experiment with 453 professionals

performing writing tasks, completion times declined sharply while quality improved when participants could use *ChatGPT* (Noy & Zhang, 2023). A large-scale field experiment involving 758 management consultants likewise found that teams using *GPT-4* completed 12 percent more tasks and achieved 25 percent faster turnaround on appropriate assignments, while also illustrating the “jagged technological frontier” where performance gains depend on task fit (Dell’Acqua et al., 2023).

Together, these studies motivate RQ1 by showing that measurable task-level improvements in efficiency and quality are both plausible and detectable through survey-based perception data.

2.3.2 Job Insecurity, Displacement Risk, and Technostress (RQ2)

While augmentation effects are prominent, credible evidence also documents displacement pressures in specific domains. In U.S. manufacturing, local labor markets with greater exposure to industrial robots experienced declines in both employment and wages, indicating that automation can generate negative distributional consequences (Acemoglu & Restrepo, 2020). At the individual level, *STARA* awareness – the perception that Smart Technology, AI, Robotics, and Algorithms could replace aspects of one’s job – has been linked to lower organizational commitment, reduced career satisfaction, and higher turnover intentions (Brougham & Haar, 2018).

Beyond job insecurity, classic technostress research demonstrates that technology-related stressors are inversely associated with individual productivity and positively associated with role stress (Ragu-Nathan et al., 2008). The Job Demands–Resources (JD-R) framework integrates these mixed outcomes: AI can function as a *resource* (reducing effort and enhancing efficiency) or as a *demand* (increasing surveillance or requiring relearning), with the overall impact depending on the balance between pressures and supports (Bakker & Demerouti, 2017).

Together, these strands of evidence justify RQ2’s focus on perceived job insecurity and stress, capturing both the potential risks and the resource dynamics associated with AI adoption.

2.3.3 Skills, Upskilling, and Organizational Support (RQ3)

AI adoption is reshaping the workforce skill mix toward complementary human capabilities and continuous learning. Macro-level evidence shows increasing returns to social and collaborative skills – which are difficult to automate – alongside rising demand for technical literacies that enable effective human–AI collaboration (Deming, 2017).

Survey data from firms and workers across seven OECD countries further indicate that, while attitudes toward AI’s impact on performance are generally positive, training and facilitation remain pivotal. Employer-provided AI training is consistently associated with better perceived outcomes and reduced fear of job displacement (Lane et al., 2023).

These findings motivate RQ3 by linking perceived skill gaps and organizational support to employees' readiness for AI adoption.

2.3.4 Ethics, Transparency, and Trust in AI (RQ4)

Employee trust in AI is shaped by transparency, explainability, and governance. Foundational work on explainable AI (XAI) demonstrates that localized explanations for model outputs help users determine when to trust or question a prediction (Ribeiro et al., 2016). Broader reviews further argue that designing explanations consistent with how humans naturally form causal reasoning improves both acceptance and appropriate reliance on AI systems (Miller, 2019).

From a governance perspective, algorithmic auditing frameworks operationalize accountability throughout the AI lifecycle and are increasingly recommended for decision domains that affect workers – for example, human resources analytics and performance evaluation (Raji et al., 2020). Collectively, these strands support RQ4 by positing that ethical concerns – such as bias, opacity, and privacy – are closely associated with employees' intentions to adopt AI and their day-to-day willingness to rely on its outputs.

2.4 Workforce Transformation in the AI Era

The integration of artificial intelligence (AI) into the workplace is transforming the workforce in profound and multifaceted ways. This section explores how jobs are being redefined rather than merely eliminated, emphasizing the shift from displacement to augmentation and identifying the new skills and competencies emerging as critical for success. It also considers how career trajectories may evolve within an AI-intensive environment, as organizations adapt their roles, structures, and learning pathways to fully leverage the potential of human–AI collaboration.

2.4.1 Job Displacement vs. Augmentation

As discussed in Section 2.2.2, the prevailing evidence suggests that artificial intelligence (AI) is augmenting many jobs by automating specific tasks rather than fully replacing entire roles. However, the boundary between augmentation and displacement remains fluid and often a matter of degree. In some cases, the cumulative automation of tasks significantly reduces the need for human involvement in a given role, leading to workforce displacement or redeployment. In others, the core functions of the job persist, but the human contribution shifts toward higher-level, analytical, or interpersonal activities, representing true augmentation rather than substitution.

An optimistic interpretation of augmentation posits that AI can relieve workers from repetitive tasks, enabling them to focus on more meaningful, creative, or relational aspects of their roles. For example, lawyers now use AI for first-draft contract analysis: the system scans and flags potential issues within seconds – work that would otherwise occupy a junior associate for hours. The lawyer can then use these insights to craft strategies and conduct negotiations more efficiently. The role does not

disappear; rather, time is reallocated from routine reading to analytical and client-facing work (Remus & Levy, 2016).

Many similar examples are emerging. Journalists employ AI to summarize datasets and suggest narrative angles, freeing more time for interviews and story refinement. Teachers leverage AI to grade routine quizzes, allowing greater focus on mentoring students individually. In these scenarios, AI acts as a *force multiplier*, enhancing productivity and potentially increasing job satisfaction – assuming the tasks freed up are indeed more engaging. Studies of early AI adopters confirm that many employees appreciate having tedious tasks offloaded. For instance, customer support agents reported enjoying AI-generated response suggestions, as they reduced repetitive typing and enabled them to serve more customers or tackle more complex inquiries.

However, augmentation can also be accompanied by increased work intensity. If AI enables an employee to work 40 percent faster, will their employer then expect 40 percent more output? The literature raises concerns about a potential *productivity paradox* for workers: historically, when technologies enhanced productivity, workloads sometimes rose proportionally, or staffing levels were reduced, leaving individual workers under similar – or even greater – pressure. This dynamic can contribute to work intensification and stress (see Section 2.5.3).

The ideal of augmentation is to improve the quality of work life, yet the outcome depends largely on how organizations implement it. For augmentation to truly benefit employees, organizations must consciously reinvest productivity gains into reducing overload or enriching job content – rather than simply increasing expectations or output quotas. On the other side of the spectrum, some roles overlap so extensively with AI's capabilities that displacement becomes a tangible risk. Classic examples include routine content generation – such as basic sports or financial news reports that AI can now produce – as well as data entry, data processing, and certain customer service positions where AI chatbots can handle Tier-1 queries end-to-end. In such scenarios, AI may replace a portion of the workforce.

For instance, if a company deploys an AI chatbot capable of resolving simple customer requests, fewer frontline call-center agents may be required. These agents might then be reassigned to manage complex cases – representing augmentation of those specific tasks – yet the total number of agents could still decline. The World Economic Forum, (2023) projects that clerical and secretarial roles will continue to contract as AI systems assume tasks such as scheduling, basic bookkeeping, and form processing. Job postings already reflect this shift, with fewer listings for data-entry clerks and more for data analysts.

Over time, some occupations may shrink substantially or even disappear. The role of the transcriptionist, once essential for converting recorded audio into text, is rapidly being displaced by AI-

driven speech-recognition systems, while assembly-line inspectors are increasingly being augmented – or replaced – by computer-vision tools capable of detecting defects in real time.

Rather than causing widespread job losses, technological change more often results in the redesigning and evolution of existing roles. Hötte et al. (2023) demonstrate in their systematic review that although new technologies may replace specific tasks, compensating mechanisms frequently emerge that sustain – and in some cases expand – overall employment. What changes most is the *human component* of work.

New categories of work are emerging specifically to support and interface with AI systems – positions that were rare or even non-existent until very recently. One prominent example is the *prompt engineer*, now recognized as a specialist who designs and optimizes inputs for generative-AI models to achieve desired outputs. Vu & Oppenlaender, (2025) find that prompt engineers require a blend of technical AI knowledge, prompt-design expertise, creative problem-solving ability, and strong communication skills – a competency mix that differs markedly from other established roles.

In parallel, a growing body of empirical research examines how AI innovation affects employment – specifically, whether it displaces workers by automating tasks or augments work by supporting humans and generating new tasks. For instance, Marguerit et al., (2025) in *Displacement or Augmentation? The Effect of AI* find that AI innovations related to creativity, learning, and engagement tend to augment human labor, whereas those centered on perception are more likely to displace it. Furthermore, Marguerit et al., (2025) distinguish between *automation-AI* and *augmentation-AI*, showing that while automation tends to reduce employment or wages in lower-skilled occupations, augmentation contributes to the creation of new jobs and higher compensation – particularly among high-skilled roles.

A key determinant of whether displacement or augmentation dominates is organizational strategy and policy choice. Companies can deploy AI primarily as a cost-reduction tool – a path that often results in workforce displacement – or as a means to expand capabilities, drive innovation, and retain human talent, which supports augmentation. Some firms are explicitly pursuing augmentation strategies, retraining employees rather than releasing them when tasks are automated or restructured (Marguerit et al., 2025). The mixed empirical evidence suggests that while some displacement is inevitable, job transformation and widespread augmentation are likely to remain the prevailing pattern across many sectors.

For professionals in information technology, this implies that roles are more likely to evolve than disappear. Routine programming or testing tasks may decline or become automated, but demand will grow for specialists who can build, maintain, monitor, and improve AI systems. In addition, hybrid roles that combine technical expertise, AI oversight, interpretive judgment, and human–AI coordination

are expected to expand. A dynamic approach to career development is therefore essential – anticipating both which sub-tasks are likely to be automated and which new competencies and roles will emerge.

2.4.2 Skills and Competencies for the AI Era

The diffusion of AI technologies across industries is not only transforming specific tasks but also reshaping the broader notion of employability. While the automation of routine activities continues to accelerate, a growing consensus from international research and employer surveys indicates that new bundles of skills are essential for workers to thrive in this changing environment (World Economic Forum, 2023). These skill sets combine technical literacies with *uniquely human* abilities that AI cannot easily replicate – such as complex reasoning, ethical judgment, and socio-emotional intelligence.

The World Economic Forum’s *Future of Jobs Report (2023)* is particularly influential in mapping this shift. It identifies four broad categories of competencies expected to gain prominence over the next five years: cognitive skills, technical skills, social and interpersonal abilities, and self-management capacities. Each category reflects not only employer demand but also a structural response to how AI is reshaping value creation and the division of labor between humans and intelligent systems.

At the apex of future skill demand are higher-order cognitive abilities. The World Economic Forum (2023) ranks analytical thinking as the most critical skill globally. This encompasses problem structuring, critical reasoning, and the ability to extract actionable insights from complex datasets. Although AI systems can perform sophisticated analyses, human oversight remains essential for contextual interpretation, sense-making, and ethical evaluation. Analytical thinking therefore becomes the mechanism through which AI outputs are transformed into credible managerial or strategic decisions.

Closely following is creative thinking, ranked second by employers. Creativity in this context extends beyond artistic expression to include originality, divergent thinking, and the ability to recombine knowledge in novel ways. Crucially, creativity also applies to the use of AI itself. As generative models become mainstream, individuals who can design innovative prompts, imagine new applications, and integrate AI into business processes will gain a comparative advantage. Supporting this view, McKinsey & Company, (2023) found that the most successful firms are those that *re-bundle* AI into innovative workflows – rather than merely substituting existing routines.

While not all workers will become data scientists, a baseline level of AI and big-data literacy is increasingly essential. The World Economic Forum (2023) lists skills in data management, programming, and digital platforms among the fastest-growing areas of demand. For information technology professionals in particular, competencies in AI operations (MLOps), model evaluation, and the design of human-in-the-loop processes are becoming critical.

Even non-technical roles will require familiarity with AI’s capabilities, limitations, and risks – for example, understanding bias in algorithmic outputs or knowing how to verify results. Lane et al. (2023) emphasize that this functional literacy in AI will be a prerequisite for informed participation in digital workplaces

As analytical and technical work becomes increasingly augmented by machines, the human comparative advantage is shifting toward communication, collaboration, and leadership. The World Economic Forum (2023) consistently emphasizes employers’ growing demand for individuals who can mediate between technical outputs and organizational decision-making, articulate AI-driven insights to diverse stakeholders, and lead hybrid teams in which humans and AI systems jointly contribute.

These are often labeled as “soft” skills, yet in practice they are difficult to automate and increasingly central to organizational effectiveness. For instance, while AI can generate a financial forecast, it still requires human persuasion, negotiation, and contextual judgment to align stakeholders around its implications and translate data-driven outputs into actionable business strategies.

Finally, the accelerating pace of technological change places a premium on adaptability, resilience, curiosity, and continuous learning. Workers must not only acquire new skills but also *unlearn* outdated practices. This requires meta-learning capabilities – the ability to learn how to learn – enabling individuals to continuously reskill as technologies evolve. The World Economic Forum (2023) notes that employers increasingly value “resilience, flexibility, and agility,” reflecting both the uncertainty associated with AI-driven transformation and the need for employees who can thrive in dynamic, rapidly changing environments.

For information technology professionals, the shift is particularly pronounced. Routine programming, debugging, and testing tasks are increasingly supported by coding copilots and automated development pipelines. At the same time, demand is rising for professionals who can (i) integrate these tools effectively, (ii) design robust system architectures that incorporate AI components, and (iii) monitor, govern, and explain AI outputs responsibly.

This transformation presents a dual challenge: technical upskilling in AI-adjacent domains and the simultaneous strengthening of transversal skills such as problem-solving, communication, and leadership. As a result, career trajectories in IT are becoming less about linear specialization and more about cultivating *dynamic competence portfolios* – blending deep technical expertise with cognitive agility and human-centered capacities.

The ADP Research Institute (2024) found that younger workers increasingly question whether their education has adequately prepared them for the transformations brought about by AI. Across the literature and industry surveys, there is a growing consensus on the need for large-scale reskilling and upskilling initiatives. In response, many companies have launched training programs – often in

partnership with online learning platforms – to strengthen employees’ data literacy, introduce basic coding skills, and teach effective use of AI tools within their specific domains.

Governments and educational institutions are also adapting. Universities are embedding AI ethics and interdisciplinary AI applications into new curricula, while vocational and professional training programs are adding data-analysis and automation modules. Together, these initiatives represent an emerging ecosystem aimed at closing the AI readiness gap across generations and sectors.

2.5 Ethical and Societal Implications of AI in Employment

The deployment of AI in the workplace raises a wide range of ethical, legal, and societal questions. Even as organizations pursue the productivity and efficiency gains discussed earlier, they must also navigate challenges related to fairness, privacy, transparency, and employee well-being. This section examines these dimensions in greater depth – focusing on frameworks for algorithmic accountability, the evolving regulatory landscape, and the implications for worker rights, autonomy, and psychological health.

2.5.1 Algorithmic Bias, Fairness, and Accountability

One of the most widely documented risks of AI systems particularly those based on machine-learning techniques – is their tendency to replicate or even amplify historical biases embedded in training data. In employment contexts, this risk has significant implications for fairness, equal opportunity, and compliance with anti-discrimination law. As AI tools become more prevalent in recruitment and human-resources processes – such as résumé screening, video-interview analysis, and employee-performance evaluation – biased outputs can translate directly into unequal treatment of candidates and workers.

A widely cited case illustrates the risks involved: Amazon’s experimental AI hiring system, trained on historical résumés dominated by male applicants, learned to systematically downgrade female candidates (Dastin, 2018). Although the system was ultimately discontinued, the episode underscores a structural danger – when past inequalities are encoded in training data, AI systems can “learn” and perpetuate them at scale. Similar concerns have been raised about automated résumé parsers that favor certain universities or ZIP codes, as well as facial-analysis tools that perform less accurately on women and people of color (Buolamwini & Gebru, 2018).

To mitigate such risks, the principle of *algorithmic accountability* has gained increasing prominence. This concept emphasizes that automated decisions must be explainable, auditable, and contestable. A notable framework in this regard is Microsoft Research’s FATE initiative – *Fairness, Accountability, Transparency, and Ethics* – which advocates embedding fairness and transparency throughout the entire AI lifecycle. Complementing this perspective, Lepri et al., (2019) outline the conceptual and technical foundations necessary to ensure that algorithmic decision-making processes

are fair, transparent, and accountable. Similarly, Selbst et al., (2021) highlight the importance of *algorithmic impact assessments* as institutional mechanisms that reinforce accountability both before and after system deployment.

Practical mechanisms for promoting algorithmic accountability include: **Auditing training datasets** to detect and correct representational imbalances – for example, the underrepresentation of women or minority groups. **Testing model outputs** to evaluate potential disparate impacts on protected groups prior to deployment. **Applying fairness-aware algorithms**, such as reweighting techniques, decision-threshold adjustments, or constraint-based optimization, to reduce disparities across demographic categories. **Providing meaningful explanations** to affected individuals, particularly in high-stakes contexts such as hiring, ensuring that rejection decisions are understandable rather than dismissed with opaque justifications such as “the AI said so.”

In organizational practice, an increasing number of companies are establishing *Algorithmic Ethics Committees* or equivalent oversight structures. These bodies often include legal experts, ethicists, technical specialists, and employee representatives, bringing diverse perspectives to the review of AI systems used for recruitment, promotion, and termination decisions (European Commission, 2021). Such governance mechanisms signify a growing recognition that algorithmic decision-making is not merely a technical concern but also a matter of organizational justice and societal accountability.

Another essential dimension of accountability involves the assignment of responsibility when AI participates in decision-making. For instance, if an AI-based scheduling system consistently allocates fewer shifts to older workers – raising concerns of age discrimination – who bears accountability? Organizations cannot deflect responsibility onto the algorithm; both regulatory frameworks and ethical norms establish that the company remains accountable for its tools and outcomes. This underscores the need for proactive monitoring of AI systems and the establishment of recourse mechanisms that allow employees or applicants to appeal or request human review of automated decisions.

The European Union’s *General Data Protection Regulation* (GDPR) provides individuals with rights concerning automated decision-making – including the right to human intervention and to receive meaningful information about the logic involved – often interpreted as a limited *right to explanation* (European Commission, 2021).

Transparency is a critical component of responsible AI governance. Many AI models operate as “*black boxes*” but in employment contexts, excessive opacity can undermine trust and hinder effective oversight. Techniques for *explainable AI* (XAI) are advancing – for example, the use of simplified, rule-based surrogate models that approximate the behavior of complex algorithms or

highlight the most influential factors in a particular decision (e.g., “this résumé was ranked lower because it lacked X skill, which was prioritized”).

A persistent trade-off exists between accuracy and interpretability. Highly accurate models, such as deep neural networks, are often less transparent, while more interpretable models may sacrifice some predictive precision. In workplace applications, organizations may need to favor explainability – opting for models that can be understood and justified to audit committees, regulators, or rejected candidates – to preserve fairness, accountability, and stakeholder trust.

Beyond technical fixes, bias in AI is fundamentally a sociotechnical challenge. Addressing requires diversifying the teams that develop AI systems to reduce blind spots, involving domain experts and affected stakeholders in model design, and continuously revisiting models as data and social contexts evolve. For instance, a performance-evaluation algorithm trained on historical company data may inadvertently reflect past inequities – for example, if certain groups were historically under-promoted. Simply codifying those patterns would perpetuate bias. Instead, organizations must decide which metrics and values their AI systems should optimize for – potentially reweighting or adjusting training data to remove historical distortions – thereby embedding ethical judgment directly into the model’s objectives.

Frameworks such as the IEEE’s *Ethically Aligned Design* provide high-level guidance for such practices, emphasizing principles like prioritizing human well-being, ensuring representative datasets, and maintaining auditability of AI decisions (Bengesi et al., 2023). Similarly, the *Partnership on AI* – a multi-stakeholder industry consortium – has issued guidelines for fair AI in hiring, recommending the avoidance of proxies that strongly correlate with protected attributes and emphasizing the need for meaningful human involvement in final decision-making.

Legal accountability is also expanding. In the United States, the Equal Employment Opportunity Commission (EEOC) has intensified its scrutiny of AI-based hiring tools for potential discrimination, while several jurisdictions – such as New York City through its *Bias Audit Law* – now mandate that AI hiring systems undergo independent annual bias audits. These developments are driving the emergence of standardized audit protocols analogous to financial audits but focused on algorithms and automated decision systems U.S. Equal Employment Opportunity Commission.

In summary, ensuring fairness in AI-driven employment decisions has become a top ethical and regulatory priority. The literature and emerging practice emphasize proactive measures: rigorous testing for bias, the integration of fairness criteria into model design, sustained human oversight, and the implementation of transparent systems that are understandable to those affected. Organizations deploying AI without these safeguards not only risk harm to individuals and reputational damage but may also face legal repercussions as regulatory frameworks evolve. Thus, while AI offers significant

efficiency gains, it also imposes a responsibility to uphold equity and accountability – a challenge increasingly addressed through the combined efforts of technical innovation, governance, and policy.

2.5.2 Privacy and Surveillance

Another major concern involves how AI technologies may enable intensified worker surveillance and erode privacy. Modern AI systems, when combined with pervasive sensors and large-scale data collection, allow employers to monitor employees in unprecedented detail – tracking not only keystrokes or email activity but also analyzing video feeds, capturing voice tone during calls, and even inferring physiological indicators such as attention or mood through webcams and wearable devices.

The term *algorithmic management* describes systems in which AI algorithms allocate work, evaluate performance, and even discipline employees – a model increasingly prevalent in gig-economy platforms such as Uber and in warehouse operations at companies like Amazon (Lee et al., 2015). These systems often optimize for efficiency but can leave workers feeling managed by an unaccountable and opaque system rather than by human supervisors. For example, ride-share drivers have long reported being “deactivated” – effectively terminated – by algorithms with little or no explanation. Similarly, warehouse employees may be monitored by AI systems that track picking speed and issue automated warnings or reduce shifts if productivity targets are not met, creating a high-pressure work environment.

The ethical concern is that such pervasive surveillance and control can infringe on workers’ rights and dignity. Continuous monitoring may lead to stress, anxiety, and a sense of lost autonomy. Privacy International and other NGOs have documented instances of intrusive digital oversight – for example, AI-enabled camera systems that calculate “focus scores” based on how frequently employees look at their screens, or remote-work software that captures periodic screenshots of workers’ computers. These practices raise fundamental ethical and legal questions: where is the line between legitimate productivity oversight and privacy invasion? And to what extent do employees meaningfully consent to such monitoring, or have the ability to opt out?

There is growing momentum for regulation to protect workers from excessive surveillance and data exploitation. The European Commission’s proposed *Platform Work Directive* seeks to regulate algorithmic management in gig work by granting platform workers stronger rights to transparency and human oversight. More broadly, the European Union’s *General Data Protection Regulation* (GDPR) also applies to employee data, and several legal cases have challenged practices such as continuous webcam monitoring as forms of excessive data processing. In the United States, although there is no comprehensive federal privacy law, several states – notably California – have enacted strong privacy protections that may extend to employee data and workplace monitoring.

AI systems thrive on data, and employers possess vast quantities of information about their workers. Ensuring data privacy therefore requires limiting the collection and use of personal data to

what is necessary, supported by robust safeguards. When companies begin using AI to analyze employee communications – such as emails or chat messages – to infer “sentiment” or “engagement,” they risk crossing ethical boundaries, especially when such analysis encroaches upon personal correspondence.

Biometric data is particularly sensitive. Facial recognition systems used for time tracking or AI-driven camera analytics designed to detect “unsafe behavior” in factories raise acute privacy and consent concerns. Amazon, for example, has patented a concept for an AI-powered warehouse bracelet capable of tracking workers’ hand movements in real time – a stark illustration of the trade-off between efficiency and privacy.

From both an ethical and legal standpoint, employees should be clearly informed about what monitoring occurs and how AI systems are used in their evaluation. In practice, however, workplace power dynamics often undermine meaningful consent – employees may “agree” to monitoring by default simply because they need the job. This asymmetry is why labor organizations increasingly advocate for enforceable legal protections rather than reliance on company policies alone.

An emerging concept in this context is *privacy by design* for workforce AI systems – the principle that privacy safeguards should be integrated into the design and operation of monitoring technologies from the outset. This means developing AI tools that minimize intrusiveness by aggregating rather than individualizing data wherever possible, anonymizing information when feasible, and flagging only genuinely concerning patterns instead of engaging in continuous, comprehensive behavior tracking.

Excessive surveillance through AI systems can significantly erode trust between employees and employers. When workers feel that every action is being monitored and analyzed, they may disengage or develop an adversarial attitude toward management. Research suggests that such monitoring can even backfire on productivity, as employees attempt to circumvent tracking systems or experience declining morale due to perceived distrust.

The ethical imperative is to balance legitimate business interests with respect for employee privacy. AI systems should be evaluated not only for their accuracy or return on investment but also for their impact on workplace climate and individual rights. This may mean choosing not to implement certain monitoring technologies when the benefits are marginal and the privacy intrusion substantial. It also requires granting employees’ visibility into the data collected about them and, where appropriate, the ability to correct or contextualize it – for example, when an AI system erroneously classifies a worker as “idle” for not typing while they were actually reading or analyzing a document.

Leading organizations in this space have begun implementing responsible policies such as prohibiting webcam monitoring during remote work unless absolutely necessary, configuring AI tools

to identify aggregate behavioral trends rather than individual “infractions,” and ensuring that any automated employment action – such as termination – is subject to human review.

In sum, while AI can enhance workforce management and operational efficiency, it must be deployed in ways that preserve human dignity, privacy, and trust. Otherwise, as human rights organizations caution, the workplace risks resembling a “digital sweatshop,” where workers are managed by algorithmic systems rather than human judgment. The literature therefore recommends a cautious approach, involving employee feedback and participatory design when introducing AI-enabled monitoring systems. The next subsection continues this discussion by exploring psychological effects and power dynamics in greater depth.

2.5.3 Worker Rights, Well-Being, and Power Dynamics

The integration of AI into the workplace raises not only technical and privacy concerns but also profound implications for employee well-being – including mental health, stress, and job satisfaction – as well as for the broader balance of power between workers and employers.

Rapid technological change can generate what researchers term *technostress* – the stress experienced from continuous adaptation to new technologies or from being perpetually connected through digital systems. The rise of AI amplifies this effect: employees may experience anxiety about being replaced, fear of inadequacy when struggling to keep pace with AI tools, or cognitive fatigue from overseeing algorithmic outputs. This phenomenon, sometimes described as the *operator’s paradox*, reflects the mental strain of supervising an AI that performs much of the work while still requiring human vigilance to detect errors (Camerini et al., 2022).

Empirical findings suggest mixed psychological outcomes. Kim & Lee, (2024), for example, found that AI adoption increased job-related stress, which in turn elevated burnout levels – except among employees with high self-efficacy for learning AI, which buffered against stress. This indicates that without adequate organizational support, training, and role clarity, the introduction of AI may exacerbate employee stress and negatively affect mental health.

Another critical aspect concerns the balance between trust and over-reliance on AI. When employees lack trust in AI systems, they may experience stress or frustration when required to use tools they perceive as unreliable. For example, a physician who distrusts an AI diagnostic assistant might feel compelled to double-check every output, effectively increasing cognitive load and workload (Dzindolet et al., 2003). Conversely, excessive trust can be equally problematic: if an employee over-relies on AI and an error occurs, the result may be moral injury, reputational harm, or even legal consequences. Achieving the appropriate balance of calibrated trust is therefore essential for psychological well-being – a balance best fostered through training, experience, and clear organizational communication about AI’s role and limitations.

Awareness that AI could potentially render certain roles redundant can create a pervasive sense of insecurity. If left unaddressed, this anxiety can erode morale, loyalty, and engagement, as employees begin to question the value of their contributions (“Why put in extra effort if a machine will soon take over?”). On the other hand, transparent communication from leadership that AI is intended to *assist* rather than replace – supported by concrete actions such as reskilling programs instead of layoffs – can alleviate these fears. A supportive, trust-based culture can transform AI from a perceived threat into a tool of empowerment: when workers see that AI helps them perform better and that their employer invests in their growth, job satisfaction and commitment tend to rise (Brougham & Haar, 2018). However, realizing this outcome requires intentional culture building and continuous reinforcement from management.

As discussed earlier in relation to algorithmic management, AI-driven decision systems can significantly diminish workers’ sense of autonomy – a central determinant of job satisfaction and psychological well-being, as outlined in *self-determination theory*. When an AI system schedules tasks, dictates workflows, or even recommends when to take breaks, employees may feel they have little control over their workday. Humans generally resist micromanagement, and being directed by an algorithm can elicit similar frustration – compounded by the fact that one cannot reason with or negotiate with the system unless meaningful feedback channels exist. Persistent loss of control can foster learned helplessness, disengagement, or resentment toward the organization.

To mitigate these effects, some companies design AI systems that provide recommendations employees can override or modify, rather than issuing rigid commands. For example, Uber introduced features allowing drivers greater visibility and discretion – such as viewing ride destinations in advance – in response to widespread concerns about feeling “enslaved to the algorithm.” Such adjustments illustrate how incorporating human agency into AI-supported workflows can preserve autonomy, enhance trust, and reduce psychological strain.

AI is reshaping the distribution of power within the workplace. Employers who leverage AI to measure, predict, and optimize nearly every aspect of labor gain significant informational and managerial advantage. Conversely, highly skilled employees who learn to harness AI effectively may see their individual influence increase – for example, a data analyst who masters AI-driven tools could become indispensable to organizational decision-making (Kellogg et al., 2020). Yet, in aggregate, AI tends to reinforce existing hierarchies by concentrating control among those who design, deploy, and own the systems – typically senior management or technology vendors.

Labor unions and worker advocates are increasingly recognizing these dynamics and are beginning to incorporate AI-related provisions into collective-bargaining agendas. Some unions now demand that any AI-based monitoring or policy decisions be subject to negotiation rather than unilateral implementation. This recalls earlier industrial movements that fought to regulate mechanization and

assembly-line speeds. Future labor demands may similarly include new rights tailored to algorithmic management – such as a *right to disconnect* extended into a *right not to be constantly algorithmically monitored* or a *right to explanation* for AI-mediated employment decisions.

A key approach to mitigating power imbalances is the active involvement of employees in AI implementation through *participatory design*. When workers are included in decisions about how AI systems are selected, configured, and deployed, those systems are more likely to align with real workplace needs and gain user acceptance. For example, if nurses participate in selecting and tuning an AI scheduling system for their hospital, the resulting tool is more apt to reflect on-the-ground realities and to be viewed as supportive rather than coercive.

Such participation also gives employees a voice in shaping how technology affects their work – transforming AI from something done *to* them into something developed *with* them. Research from the Stanford Institute for Human-Centered Artificial Intelligence (HAI) and similar organizations indicates that technology adoption proceeds more smoothly when end-users have meaningful input and feel a sense of ownership over AI tools (Shneiderman, 2020).

Ethical implications of AI adoption are closely tied to leadership and organizational culture. In the AI era, *ethical leadership* involves transparency about how AI is used, establishing clear boundaries for its application, and prioritizing long-term human well-being alongside short-term productivity gains. Leaders must ask critical questions such as: *Will this AI improve employees' work lives or simply extract more output from them? How can it be implemented as a positive-sum enhancement rather than a zero-sum trade-off?*

At the policy level, the concept of *algorithmic rights* for workers is beginning to emerge. Such rights could include protection from fully automated employment decisions – echoing but extending the GDPR's stance to explicitly cover workplace contexts – as well as the right to access personal data collected at work, the right to fair and bias-audited AI systems, and potentially even a *right to meaningful work*. The latter, while more philosophical, reflects a growing concern that if AI automates all complex and stimulating aspects of a role, the remaining human tasks may become less fulfilling (Dexe et al., 2022). This line of inquiry highlights the need to balance efficiency with human purpose and dignity in AI-enabled workplaces.

The introduction of AI also reshapes how teams' function and collaborate. Questions of accountability and recognition arise: if an AI system generates a recommendation, who deserves credit? Conversely, if it fails, who bears responsibility? Team dynamics can shift when some members effectively leverage AI tools while others do not or cannot, potentially creating an *AI gap* between employees. Managers may inadvertently favor high adopters, alienating those slower to adapt.

Managing this requires inclusive training, peer learning (for example, pairing power users with novices), and fair evaluation criteria that consider differential access to or proficiency with AI tools.

It is also essential to ensure that AI integration does not inadvertently erode collaboration. If each team member follows AI-generated directions independently, interpersonal communication and cooperative problem-solving may decline – undermining the collective intelligence that effective teams rely on.

3. Methodology

3.1 Research Design

3.1.1 Philosophical stance: Pragmatism

This study adopts a pragmatist philosophical stance. In the mixed-methods and applied social-science literature, pragmatism has been proposed as a “new guiding paradigm” for research that combines qualitative and quantitative methods, because it focuses on solving practical problems and treating methods as tools rather than as extensions of fixed epistemological camps (Morgan, 2007). From this perspective, pragmatism is problem-oriented, pluralistic, and method-agnostic: research methods are selected for their practical usefulness in answering the research questions rather than for loyalty to any single metaphysical or epistemological tradition.

In the present inquiry how workers, particularly IT professionals, perceive the effects of artificial intelligence (AI) on their work a pragmatist stance provides the rationale for combining quantitative and qualitative evidence in a single design. It supports the use of a descriptive, cross-sectional survey as the primary strand, complemented by embedded open-ended questions that capture participants’ own explanations, concerns, and suggestions in their own words.

This stance is also consistent with the empirical state of the art reviewed in Chapter 2. Section 2.2.3 showed that contemporary research on AI and employment relies on multiple, complementary methodological traditions – including randomized experiments, task-exposure models, and large-scale worker surveys. Rather than privileging any single tradition, the present dissertation explicitly positions itself within this pluralistic landscape and adopts the methods that are most appropriate for mapping individual-level perceptions across the four research questions (RQ1–RQ4) in a rapidly evolving context.

Consistent with this orientation, the overall design is descriptive and exploratory rather than causal or predictive. The emphasis is on mapping current perceptions and identifying patterns that can inform future, more generalizable studies.

3.1.2 Time horizon: Cross-sectional survey

This research adopts a cross-sectional design, capturing a snapshot of participants’ perceptions during the fieldwork period from mid-June to 31 July 2025 – a time marked by accelerated diffusion of generative AI into knowledge work.

The design provides insight into contemporary attitudes and experiences but does not track changes over time. Accordingly, all inferences are explicitly bounded to mid-2025 and interpreted as context-specific rather than longitudinal trends.

3.1.3 Strategy: Concurrent embedded mixed-methods

This study employs a concurrent embedded mixed methods strategy, implemented through a single online questionnaire. The primary quantitative component consists of Sections A–C (Q1–Q31). Each item is on a Likert scale and includes an explicit “I don’t know / Not sure” option; all items were mandatory for submission. The embedded qualitative component is Section D (Q32–Q35), comprising optional open-ended questions that enabled respondents to elaborate on perceived benefits, risks, and suggestions in their own words.

Both strands were collected simultaneously and analyzed descriptively. Integration occurred during interpretation, using triangulation to assess convergence, complementarity, and divergence. The design can be summarized textually as follows: one Microsoft Forms instrument; fieldwork conducted from mid-June to 31 July 2025 with two reminder waves; quantitative (QUAN) analysis performed in Microsoft Excel; qualitative (QUAL) data analyzed through inductive thematic coding in Excel; and de-identification procedures applied, with no personal identifiers stored.

3.1.4 Alignment to research questions (RQ1–RQ4)

Instrument blocks and items were designed to map directly to the four research questions:

- **RQ1 – Task-level outcomes** (productivity, quality, time): Q12–Q17.
- **RQ2 – Job-level risks and enrichment**: Q22–Q25 (with Q22 and Q23 negatively phrased).
- **RQ3 – Skills, training & organizational support** (readiness): Q27–Q30.
- **RQ4 – Fit, trust, transparency, ethics, and adoption intentions**: Q18–Q21, Q26, Q31.

The detailed RQ construct item mapping is shown in Table 1 of this chapter and is consistent with the questionnaire reproduced in Appendix A.

3.1.5 Scope of inference and threats to validity

This study is based on a non-probabilistic, cross-sectional survey relying on self-reported perceptions. As such, several validity limitations apply.

First, internal validity is limited: the design does not support causal inference. Second, external validity is constrained by both the sectoral concentration of the sample (approximately 83% working in IT/technology roles) and the absence of a known sampling denominator, which precludes generalization to a broader population. Third, common method bias and measurement error remain possible despite procedural safeguards, including neutral item wording, anonymity assurances, provision of a “I don’t know / Not sure” option, and the inclusion of a small set of negatively phrased items. Accordingly, the findings are positioned as an **exploratory and descriptive mapping** of perceptions among early

adopters of AI in the workplace. Results should be interpreted as **context-specific patterns**, not as population-level estimates.

3.1.6 Why this design and not alternatives

This study employs a cross-sectional, concurrent embedded mixed-methods survey because it best aligns with the exploratory aim: mapping perceptions across the four research questions (RQ1–RQ4) within a defined fieldwork window.

Alternative designs were considered but found unsuitable. A **longitudinal (panel) study** would strengthen causal inference by tracking changes in perceptions over time but was infeasible within the time constraints of a master’s dissertation and the need for timely insights.

A field experiment or randomized trial could isolate causal effects of specific AI tools, yet would require organizational partnerships, randomization, and resource-intensive implementation beyond the project’s scope. A **qualitative-only design** (e.g., interviews or ethnography) would provide rich depth but limit breadth and comparability across constructs, while a **quantitative-only design** would miss the explanatory nuance offered by open-ended accounts.

The chosen design therefore makes a deliberate trade-off: it privileges breadth over depth while embedding qualitative prompts to enrich interpretation. It emphasizes descriptive validity and transparency - through item-level reporting and an explicit “I don’t know / Not sure” response option - over psychometric scale validation or inferential modeling, which are more appropriate for future, larger-scale research. This rationale is consistent with the study’s **pragmatist stance** and exploratory objectives.

3.2 Population and Sampling

3.2.1 Target population & inclusion criteria

The target population comprised adults aged 18 years or older who were currently or recently employed in any sector or occupation. These eligibility requirements were clearly stated on the Participant Information and Consent page that preceded the questionnaire. Entry into the survey’s first section (Section A) was conditional on respondents affirming both criteria – that they met the age threshold and that they were currently or recently employed – and on providing informed consent to participate. Only those who selected “*Yes – I consent*” were able to proceed to the questionnaire.

Participation was voluntary and anonymous, with no personal identifiers collected. The consent form informed participants of their right to withdraw at any time, the estimated completion time (~ 10 minutes), and the data-retention policy (secure storage for five years, until 31 December 2029, after which all data will be permanently deleted).

These procedures ensured that all respondents were appropriately screened and that participation complied with the ethical standards approved for this study.

3.2.2 Sampling frame & technique

Given the absence of a comprehensive international registry of employed adults, the study employed a non-probability sampling strategy that combined elements of convenience and snowball sampling. Initial invitations were distributed through the researcher's professional networks – including LinkedIn posts, direct messages, and targeted email outreach – as well as through organizational forwarding by colleagues and contacts. Recipients were encouraged to share the survey link with other eligible participants who met the inclusion criteria.

This approach prioritized feasibility and geographic reach across diverse professional roles within the available fieldwork period. While non-probability sampling limits statistical generalizability, it is appropriate for exploratory research aiming to capture perceptions from early adopters of AI-enabled tools across multiple sectors and countries.

3.2.3 Achieved sample & composition

The final analytic dataset comprised 75 complete responses that met all inclusion criteria. The majority of participants (~ 83 %) reported working in information technology or related technical roles, reflecting an early-adopter context characterized by relatively high exposure to AI systems. Consistent with the study's focus on experienced users, AI adoption was widespread among respondents: approximately 92 % indicated that they were currently using AI tools as part of their day-to-day work activities.

While the resulting sample is not statistically representative of the broader workforce, its composition aligns with the research objective of exploring perceptions among professionals with direct, practical experience of AI integration. This profile provides valuable insight into how employees in AI-intensive environments perceive productivity benefits, skill implications, and ethical considerations.

3.2.4 Representativeness & external validity

Given the non-probabilistic sampling design, the achieved sample is not statistically representative of the broader workforce, and a formal response rate cannot be calculated due to the absence of a defined sampling denominator. The sectoral concentration – with approximately 83 % of respondents employed in IT or technology-related roles – further constrains external validity and limits generalization to non-IT or late-adopting contexts.

Accordingly, the study's findings should be regarded as context-bound to mid-2025, reflecting perceptions among early adopters of AI tools across diverse professional and geographic settings.

Rather than population-level estimates, the results provide exploratory, indicative patterns that inform theory building and guide future, larger-scale research on workforce adaptation to AI.

3.3 Survey Instrument Development

3.3.1 Theoretical foundations

The survey instrument was theory-driven, with item content grounded in established models of technology adoption and work design. Specifically, the instrument drew upon the Technology Acceptance Model (TAM) and its extensions (UTAUT) to capture perceived usefulness and intention to adopt; the Task–Technology Fit (TTF) framework to assess alignment between AI capabilities and job requirements; and the Job Demands–Resources (JD-R) model to frame relationships among self-efficacy, support, and strain.

In addition, an explicit strand addressing trust, transparency, and ethics was incorporated to reflect issues central to human–AI collaboration and responsible AI use. Together, these frameworks guided the clustering of the 20 perception items in Section C of the questionnaire into four conceptual domains: 1) Task-level outcomes; 2) Fit, trust, transparency, and ethics; 3) Job-level risks and enrichment; 4) Readiness, organizational support, and adoption intention.

3.3.2 Structure & response formats (by section)

The questionnaire was implemented in Microsoft Forms and organized into four sections, each with clearly defined response rules and branching logic:

- **Section A (Q1–Q7): Demographics and background.** Collected participants’ age group, gender (including a self-describe option), highest education level, industry sector, primary job function, years of professional experience, and familiarity with AI concepts (rated on a 1–5 scale). All items in this section were mandatory.
- **Section B (Q8–Q11): AI exposure.** Assessed respondents’ current use of AI tools, types of AI systems used (with a multiple-selection list covering generative, predictive, and automation applications), frequency of use, and availability of employer-provided AI training. Question 8 included a respondent-facing definition of “AI tools” to ensure consistent interpretation. All items were mandatory.
- **Section C (Q12–Q31): Perceptions of AI’s impact on work.** Contained 20 Likert-type statements anchored from *Strongly disagree (1)* to *Strongly agree (5)*, with an additional “I don’t know / Not sure” option to avoid forced responses. Each item was mandatory. Items were grouped into four conceptual domains corresponding to the research questions: task-level outcomes, job-level risks and enrichment, skills and organizational support, and ethics/trust/adoption intentions.

- **Section D (Q32–Q35): Open-ended prompts.** Invited participants to provide narrative input on benefits, concerns, implementation suggestions, and additional reflections regarding AI in the workplace. Responses in this section were optional.

The consent page and all section headers reiterated participants' right to withdraw at any time, emphasized the voluntary and anonymous nature of participation, and specified the data-retention period of five years (until 31 December 2029). No personal identifiers were collected, and all data were analyzed in aggregate.

3.3.3 RQ–construct–item mapping table

The 20 Likert-type perception items (Q12–Q31) were mapped to the four research questions (RQs) and their underlying theoretical constructs. Table 1 presents the correspondence between each RQ, its conceptual construct(s), and the specific questionnaire items analyzed.

Table 1 - Mapping of Research Questions (RQs) to Constructs and Questionnaire Items (Q12–Q31).

<i>RQ</i>	<i>Construct(s)</i>	<i>Questionnaire Items</i>
<i>RQ1</i>	Task-level outcomes (productivity, quality, efficiency, and task fit)	Q12 Productivity ↑; Q13 Fewer errors; Q14 Repetitive tasks offloaded; Q15 Work quality ↑; Q16 Time savings; Q17 Task fit
<i>RQ2</i>	Job-level risks and enrichment (stress, anxiety, and opportunity creation)	Q22 Anxious AI may replace (–); Q23 Job more stressful (–); Q24 Higher-value tasks; Q25 Overall job impact positive;
<i>RQ3</i>	Skills, training, and organizational support (readiness)	Q27 Confident with AI; Q28 Management supportive; Q29 Training / resources adequate; Q30 Actively upskilling
<i>RQ4</i>	Trust, transparency, ethics, and adoption intentions	Q18 Trust outputs; Q19 Hard to understand (–); Q20 Transparency; Q21 Ethical concern; Q26 Enthusiastic about AI in field; Q31 Intend to use more AI

Note. Arrows (↑) indicate positive constructs; (–) denotes negatively phrased items; results are reported in the original item polarity (no reverse coding applied).

Full item wordings appear in Appendix A (“Section C: Perceptions of AI’s Impact on Work”).

3.3.4 Item wording & negatively phrased items

To minimize the potential for acquiescence bias – the tendency for respondents to agree with statements regardless of content – a small subset of items was intentionally worded in a negative direction (e.g., *Q19*, *Q22*, *Q23*). Because the analysis was conducted on an item-by-item basis and no composite scales were created, reverse coding was not applied. Each item was therefore analyzed and reported in its original phrasing and polarity, ensuring clarity and transparency in the interpretation of results.

3.3.5 Pre-survey and refinement

Before launching the main survey, a brief pre-survey (pilot pre-test) was administered to a small group of professionals drawn from the target population. The purpose of this pre-inquiry was to assess the clarity of item wording, the overall questionnaire layout, the approximate completion time, and the visibility and interpretability of all response options (including the “I don’t know / Not sure” category). Feedback from this pre-survey confirmed that the average completion time was approximately 10 minutes, in line with the estimate communicated on the participant information and consent page.

On the basis of respondents’ comments, a limited set of refinements was introduced. The most substantive change was the inclusion of a clearer definition of “AI tools” in Question 8, together with minor adjustments to section ordering, headings, and labels to improve readability and logical flow.

This pre-survey thus functioned as a methodological safeguard, ensuring that the final questionnaire used in the main data collection phase was concise, comprehensible, and fit for purpose for an online professional sample.

3.3.6 Psychometric scope

Given the exploratory focus and the modest sample size of the study, no formal psychometric analyses – such as Cronbach’s α reliability testing or factor analysis – were conducted. The research design intentionally prioritized descriptive transparency over scale generalization, focusing instead on a direct, item-level mapping of perceptions to the corresponding research questions (RQs). This approach aligns with the study’s exploratory and theory-informed objectives, emphasizing clarity of interpretation rather than latent construct validation.

3.4 Data Collection Procedure

3.4.1 Platform & timeline

Data were collected using Microsoft Forms between mid-June and 31 July 2025. Two reminder waves were issued – one in late June and another in mid-July – to maximize response rates within the limited fieldwork window. The participant information and consent page, as well as all section headers, reiterated respondents’ rights to withdraw, the voluntary nature of participation, and the response requirements for proceeding through the questionnaire.

3.4.2 Recruitment channels & message ethics

Participants were recruited through professional networks (e.g., LinkedIn posts, direct messages, and professional email outreach), informal organizational forwarding, and snowball sampling. All recruitment messages clearly described the study purpose, estimated completion time (approximately 10 minutes), voluntary participation, and absence of incentives, along with the inclusion criteria (aged 18 or over; currently or recently employed).

Messages also encouraged participants to share the survey link with other eligible contacts, while explicitly emphasizing the need to avoid any form of coercion or undue pressure. This approach ensured ethical recruitment practices consistent with the consent and confidentiality principles detailed in Section 3.6.

3.4.3 Consent gate & response rules

Access to the questionnaire was conditional upon completion of the Participant Information and Consent page, which outlined the voluntary nature of participation, the right to withdraw at any time without penalty, and the data-retention policy (five years, until 31 December 2029). Only participants who explicitly selected “*Yes – I consent*” were permitted to proceed to Section A.

Sections A–C were mandatory, ensuring complete demographic, exposure, and perception data for all respondents. Each Likert-scale item included a “*I don’t know / Not sure*” option to accommodate uncertainty and prevent forced responses. Section D, containing open-ended prompts, was optional, allowing respondents to submit the questionnaire without providing narrative comments.

3.4.4 Data handling and cleaning

Because the questionnaire consisted primarily of closed-ended items, only minimal data cleaning was required. Cleaning procedures included the standardization of free-text responses entered under “*Other (please specify)*” – for example, correcting spelling variations and ensuring consistent case formatting. Categorical and Likert-scale responses were numerically coded for analysis, while the “*I don’t know / Not sure*” category was retained and tabulated separately to preserve interpretive transparency.

The anonymized dataset, which contained no personal identifiers, was stored on a secure, access-controlled drive in accordance with institutional data management and ethical compliance guidelines.

3.4.5 Quality safeguards

Several procedural safeguards were implemented to ensure data quality, ethical compliance, and response reliability. These included:

- Clear eligibility criteria and a consent gate requiring participants to confirm age, employment status, and informed consent before accessing the questionnaire.
- Mandatory responses for Sections A–C to secure complete demographic, exposure, and perception data.
- Provision of an “*I don’t know / Not sure*” response option for every Likert-scale item to prevent forced or speculative answers.
- Concise, neutrally phrased items to minimize interpretation bias and respondent fatigue.
- A limited subset of negatively worded items to reduce straight-lining and acquiescence bias.
- Two reminder waves spaced to improve participation while avoiding over-solicitation.
- Complete anonymization: no IP addresses, names, or other personal identifiers were collected.

Together, these safeguards ensured that the data collection process adhered to ethical research standards and yielded high-quality, reliable responses suitable for descriptive and exploratory analysis.

3.4.6 Collection-stage limitations

Because the survey was administered online and only in English, it may under-represent individuals with limited digital access, lower English proficiency, or restricted availability during the fieldwork period. Additionally, the use of an open-access survey link combined with referral-based (snowball) recruitment precluded calculation of a precise sampling denominator; therefore, a formal response rate could not be established.

These limitations are acknowledged here and are further discussed in Section 3.7, alongside their implications for data representativeness and external validity.

3.5 Data Analysis Approach

3.5.1 Analytical framework & tools

Data analysis followed parallel quantitative and qualitative strands, consistent with the concurrent embedded mixed-methods design adopted for this study. All data processing and analysis were performed in Microsoft Excel to ensure transparency, accessibility, and reproducibility.

The quantitative strand (Likert-scale items, Q12–Q31) was analyzed using descriptive statistics, including frequency and percentage distributions generated through Excel pivot functions.

The qualitative strand (open-ended responses, Q32–Q35) was examined through inductive thematic coding, in which emergent themes were identified and frequency tallies recorded to convey their relative salience.

Integration of the two strands occurred at the interpretation stage, applying a triangulation logic to assess convergence, complementarity, and, where relevant, divergence between quantitative patterns and qualitative insights.

3.5.2 Quantitative analysis

The quantitative strand was descriptive in nature, designed to summarize overall response patterns rather than to test statistical hypotheses. For each Likert-type item (Q12–Q31), the following procedures were applied:

- Frequencies and percentages were calculated for all response categories, including “*I don’t know / Not sure.*” This ensured full transparency in representing both agreement levels and respondent uncertainty.
- Item means (on a 1–5 scale) were computed excluding the “I don’t know / Not sure” responses, which were reported separately to capture instances of uncertainty or inapplicability.
- Medians were reviewed as part of internal consistency checks for potential skewness but were not systematically reported.

No inferential statistical tests (e.g., *t*-tests, ANOVAs), *p*-values, confidence intervals, or multivariate models were undertaken. The objective was to provide a transparent, item-level descriptive summary of participants’ perceptions, consistent with the study’s exploratory and theory-driven design.

3.5.3 Qualitative analysis

Open-ended responses (Q32–Q35) were examined using an inductive thematic analysis approach. The analytic process followed established steps: (1) familiarization with the full set of narrative data; (2) line-by-line initial coding to capture key ideas; (3) grouping of similar codes into categories; (4) development and iterative review of themes; and (5) reporting of thematic patterns with frequency tallies expressed as descriptive percentages relative to the number of respondents who answered each prompt.

Given the modest data volume, all coding and theme aggregation were conducted in Microsoft Excel, which was sufficient for transparency and traceability. Although computer-assisted qualitative data analysis software (CAQDAS) can facilitate multi-coder workflows, it was unnecessary in this single-coder exploratory context. To preserve participant anonymity, only paraphrased and representative excerpts are cited in the Results chapter.

3.5.4 Mixed-methods integration

Integration of the quantitative and qualitative strands followed a convergent triangulation logic. Quantitative trends – for example, broad agreement that AI saves time – were juxtaposed with qualitative elaborations, such as participants’ concrete descriptions of tasks being offloaded or

streamlined. Instances of apparent divergence – for example, the simultaneous expression of high self-efficacy alongside ethical or job-security concerns – were interpreted as meaningful paradoxes rather than inconsistencies to be resolved.

Integration is presented narratively in Chapter 4, with each research question organized in a QUAN → QUAL → integrated discussion sequence to highlight convergence, complementarity, and, where relevant, tension between data types.

Reporting conventions. Percentages are computed over N = 75 respondents. *Agree / Strongly agree* (or *Disagree / Strongly disagree*, where applicable) are counted as endorsements. *I don't know / Not sure* responses are treated as neither agreement nor disagreement. Item means are calculated excluding “I don't know / Not sure” selections, which are tabulated separately to indicate uncertainty or inapplicability.

3.6 Ethical Considerations

3.6.1 Informed consent & participant rights

A Participant Information and Consent page was presented prior to the survey, providing respondents with all essential details required for informed participation. The page described the study's purpose in neutral, non-persuasive terms, the estimated completion time (approximately 10 minutes), and the procedures for data handling, storage, and retention. It also explicitly affirmed the voluntary nature of participation and participants' right to withdraw at any time by discontinuing the survey without any adverse consequences.

Proceeding beyond the consent page to Section A was considered to constitute informed consent, consistent with accepted ethical practice for anonymous, minimal-risk online research.

3.6.2 Privacy, anonymity, and data protection

No names, email addresses, IP addresses, or other direct or indirect identifiers were collected at any stage of the survey. All free-text responses were reviewed to detect any potential self-identifying information, and no such content was retained. The dataset was therefore anonymized at source and stored on password-protected devices with secure cloud backup, in accordance with institutional data-protection and ethical-research guidelines.

3.6.3 Risk assessment and mitigation

The survey's subject matter could potentially evoke mild concern (e.g., anxiety about job replacement) or discomfort (e.g., reflection on ethical issues or workplace monitoring). These risks were mitigated through several design features: 1) Concise, neutrally phrased items to minimize emotional loading; 2) Inclusion of an “I don't know / Not sure” response option to avoid forced judgments; 3) The ability to skip open-ended questions without penalty.

Given the anonymous nature of participation and the absence of any collected identifiers, both psychological and privacy-related risks were assessed as minimal.

3.6.4 Data security and retention

All de-identified datasets will be permanently deleted after five years (31 December 2029), in accordance with the procedures stated in the Participant Information and Consent document.

3.7 Limitations of Methodology

3.7.1 Sampling and external validity limitations

- The use of non-probability sampling (convenience and snowball techniques) precluded the calculation of a formal response rate and limited statistical representativeness.
- The achieved sample ($N = 75$; approximately 83 % IT professionals) reflects a cohort of early adopters, meaning that transferability to other sectors or later stages of AI diffusion is constrained.
- The online recruitment method and English-only administration may have under-represented individuals with limited digital access or lower English proficiency, potentially excluding important perspectives from less-connected or non-English-speaking populations.

3.7.2 Data collection and measurement limitations

- Self-report measures remain vulnerable to social desirability bias, recall inaccuracies, and mood effects, despite anonymity assurances and neutral item wording.
- The cross-sectional design precludes causal inference and does not capture temporal dynamics (e.g., whether training reduces anxiety over time).
- Although theory-informed, the instrument was not psychometrically validated for scale properties; the design explicitly prioritized item-level descriptives over latent constructs.

3.7.3 Implications and mitigation strategies

- The findings are framed as exploratory and context bound. Any generalizations beyond comparable early-adopter contexts should be approached with caution.
- Transparency regarding design decisions (e.g., item-level reporting, explicit inclusion of a “*I don’t know / Not sure*” response category) and acknowledgment of limitations are positioned as quality features that strengthen interpretive credibility rather than as deficiencies.
- Future research (see Chapter 5) may extend sampling breadth, adopt longitudinal designs, and apply psychometric and inferential methods once larger and more representative samples are attainable.

3.8 Linking methodology and empirical results

This section briefly clarifies how the methodological choices detailed in Chapter 3 translate into the empirical results reported in Chapter 4. The aim is to make explicit, at an abstract level, what will be done with the data and how this follows coherently from the research design, population and sampling strategy, and analytical procedures previously described.

First, the pragmatist, concurrent embedded mixed-methods design outlined in Section 3.1 provides the overarching logic for the empirical work. The quantitative strand – a cross-sectional, self-administered online survey of 75 professionals – generates descriptive indicators that map directly onto the four research questions (RQ1–RQ4). The qualitative strand – open-ended responses captured in Q32–Q35 – supplies participants’ own explanations, examples, and concerns, which are analyzed inductively to enrich and problematize the numerical patterns. This dual-strand structure underpins the organization of Chapter 4.

Second, the RQ–construct–item mapping presented in Section 3.3 ensures that each research question is operationalized through a coherent set of survey items. In Chapter 4, the results are therefore reported by RQ rather than by questionnaire section: for each domain – task-level outcomes (RQ1), job-level risks and enrichment (RQ2), skills and organizational support (RQ3), and ethics, transparency, and trust (RQ4) – the analysis begins with descriptive statistics for the relevant Likert-type items, followed by thematic synthesis of the corresponding open-ended responses. This structure reflects the theory-driven instrument design and allows a direct line of sight from constructs to evidence.

Third, the data analysis approach described in Section 3.5 is implemented in a transparent and intentionally non-inferential manner. Quantitative results are limited to frequencies, percentages, and item-level means, consistent with the exploratory purpose and the convenience and snowball sampling strategy outlined in Section 3.2. No causal claims or population-level generalizations are attempted; instead, the chapter focuses on mapping patterns within this specific, predominantly IT-professional sample and on identifying salient regularities and paradoxes in their perceptions of AI at work.

Finally, integration of quantitative and qualitative strands follows the convergent triangulation logic specified in Section 3.5.4. For each research question in Chapter 4, the narrative moves from QUAN (descriptive statistics) to QUAL (illustrative themes and excerpts) and then to an integrated discussion that relates the combined evidence back to the theoretical and empirical literature reviewed in Chapter 2. In this way, the methodological decisions in Chapter 3 not only constrain the scope of the claims that can be made but also provide a coherent framework for interpreting the findings as context-bound, employee-centred insights into AI’s impact on work.

4. Results and Discussion

This chapter presents the findings from the survey of 75 professionals regarding their perceptions of AI's impact on employment.

The results are organized thematically to address each research question systematically, integrating both quantitative survey responses and qualitative insights from open-ended questions.

As established in the methodology (Chapter 3), these findings represent the perspectives of this specific sample, predominantly IT professionals (83%) with high AI exposure (92% current users) and should be interpreted as exploratory insights rather than generalizable conclusions about all workers.

To complement the item-by-item narrative, this section provides a concise tabular overview of the main quantitative findings. It synthesizes agreement levels on the 20 perception items (Q12–Q31), grouped by research question (RQ1–RQ4), so readers can grasp patterns at a glance before diving into the detailed discussion that follows.

The table highlights the most salient indicators (e.g., time savings, productivity gains, self-efficacy) and flags areas of caution (e.g., ethics/trust).

These values reflect the results described later in Chapter 4.

Respondents reported strong task-level benefits from AI (time savings and productivity), limited personal threat perceptions, high self-efficacy and upskilling intent amid uneven formal training access, and a blend of enthusiasm with ethical concerns that temper unconditional trust.

Table 2 - “Key indicators at a glance” (Agree/Strongly agree unless noted)

<i>RQ</i>	<i>Construct / Item (Q#)</i>	<i>Result</i>	<i>Note</i>
<i>RQ1 – Task-level outcomes (productivity, quality, time)</i>	AI saves time (Q16)	94%	Clear efficiency gains.
	More productive (Q12)	84%	Productivity perceived as higher.
	Work quality improved (Q15)	~83%	Quality uplift widely perceived.
	Fewer mistakes (Q13)	~67%	Error reduction, but not universal.
	Repetitive tasks offloaded (Q14)	~60%	Mundane work reduced for many.
	Task–technology fit (Q17)	~73%	Tools broadly match job needs.

<i>RQ2 – Job-level risks & enrichment</i>	Overall job impact positive (Q25)	~80%	Net positive job-level view.
	Freed for higher-value tasks (Q24)	~69%	Enrichment via task reallocation.
	Anxious about replacement (Q22)	~28% agree; 56% disagree	Minority anxiety; majority unafraid.
	Job more stressful (Q23)	~15% agree; 70% disagree	Stress increase limited.
<i>RQ3 – Skills, training & organizational support</i>	Confident working with AI (Q27)	~91%	Very high self-efficacy.
	Management supportive (Q28)	~85%	Encouragement perceived.
	Sufficient training/resources (Q29)	~65%	Gap: formal training uneven.
	Actively upskilling (Q30)	84%	Strong proactive learning.
<i>RQ4 – Fit, trust, transparency, ethics, and adoption intentions</i>	AI decisions transparent (Q20)	~47%	Transparency seen as mixed.
	Concerned about unfair/unethical use (Q21)	~45%	Ethics concerns notable.
	Trust AI outputs (Q18)	~45%	Trust is cautious, not blanket.
	Find AI decisions confusing (Q19)	~35%	Comprehensibility can improve.
	Enthusiastic about AI in field (Q26)	~75%	Cautious enthusiasm overall.
	Intend to use more AI (Q31)	~85%	High adoption intent.

Note. Percentages summarize Agree/Strongly Agree (or Disagree where explicitly noted) as reported in the results sections that follow. See Items Q12–Q31 for wording and detailed distributions.

Chapter Structure: The results are presented in four main sections corresponding to the research questions: task-level benefits (RQ1), job-level risks (RQ2), skill development and organizational support (RQ3), and ethical concerns alongside future adoption intentions (RQ4).

Each section begins with quantitative findings followed by supporting qualitative insights and concludes with an integrated discussion of the implications.

4.1 RQ1 – Perceived Task-Level Benefits

4.1.1 Quantitative Results: Productivity and Efficiency Gains

The survey data showed that respondents overwhelmingly perceived AI as beneficial at the task-level. A strong majority agreed that AI tools have enhanced their productivity and efficiency in day-to-day tasks. Specifically, **84%** of respondents agreed or strongly agreed that using AI had made them more productive in their job (Q12; mean Likert score = 4.31 on a 5-point agreement scale).

Nearly all respondents (**94%**) agreed that AI “*saves me time on certain tasks*” (Q16; mean = 4.49). Many also noted improvements in work quality – 62/75 = 82.7% (~83%) agreed that AI has improved the quality of their work (Q15). Job-level evaluation is summarized in section 4.2.1, **80%** agreed that AI’s overall impact on their job is positive, Q25. Task–technology fit was also assessed: 55/75 respondents agreed that “*the AI systems at my workplace fit well with the tasks I need to accomplish*” (Q17; 73.3%, mean = 3.87).

Several specific task-level benefits stand out from the data. A majority of participants indicated that AI helps **reduce errors** in their work: about **67%** agreed that “*AI has helped reduce mistakes in my work*” (Q13), although roughly one-quarter neither agreed nor disagreed on that point (suggesting some were unsure or had not observed a change). Notably, **60%** of respondents reported that AI tools have **taken over the most repetitive or mundane parts** of their job (Q14), while around one-fifth did not feel this was the case. Nonetheless, Related job-level enrichment is reported under RQ2: 52/75 = 69.3% agreed that AI frees them to focus on higher-value tasks (Q24; see section 4.2.1).

In other words, a significant portion of the sample experienced AI as a means of task enrichment – offloading tedious work and allowing them to engage in more creative or strategic activities. Furthermore, the AI tools in use generally aligned well with respondents’ job needs about **73%** agreed that “*the AI systems at my workplace fit well with the tasks I need to accomplish*” (Q17; mean = 3.87).

This suggests that for most users, the AI solutions available were relevant and useful for their specific duties, likely contributing to the positive impacts reported. Overall, the quantitative results for RQ1 paint a picture of widespread perceived benefits: **respondents saw AI as a tool that boosts their efficiency, accuracy, and work quality in day-to-day tasks.**

4.1.2 Qualitative Insights: Employee Explanations and Examples

Open-ended responses about the biggest benefit or opportunity of AI further illustrate these task-level advantages. Many respondents **emphasized efficiency gains and relief from mundane work** as the primary benefits of AI. One participant noted that AI was effectively eliminating “*boring, time-consuming work*” from their day, giving them more time for “*creative tasks*.” Others echoed that AI helps “*avoid repetitive tasks*” and “*accelerate [their] work*,” underscoring the significant speed-up in workflow.

Participants also appreciated improvements in work quality and decision support. For example, one individual explained that AI can help catch mistakes or suggest optimizations, thereby “*reducing errors*” in their output. Another respondent observed that because AI could handle routine analysis, they were able to “*invest in more strategic work*” instead of being bogged down by basic processing tasks.

These reflections support the quantitative trend: the task-level benefits of AI were strongly felt across the sample. Respondents largely viewed AI as a tool that **augments their capabilities** – automating the drudgery and enabling them to work smarter and focus on higher-level aspects of their job, rather than as a system that would outright replace them in those tasks. This positive outlook on daily task efficiency reinforces the view of AI as a collaborative aid for boosting productivity and creativity in one’s role. Indeed, these day-to-day beneficial experiences likely contribute to the generally favorable attitude toward AI observed later (see RQ4).

Some nuances emerged regarding AI’s task-level impact. A few respondents cautioned that AI’s benefits come with a need for human oversight and understanding of the tool’s limitations. One participant, while acknowledging AI’s usefulness, noted that they remain “*skeptical about how much improvement AI can bring... considering it still can make mistakes and inaccurate statements.*” This highlights that even as AI tools save time and effort, users are aware that the technology is not infallible.

Issues of trust and transparency in AI’s task performance (elaborated further under RQ4) can moderate the otherwise very positive perception of its benefits. In short, the qualitative insights for RQ1 indicate that individuals in this study overwhelmingly perceived AI as delivering significant task-level advantages – especially in terms of efficiency, error reduction, and freeing them for more engaging duties – while also recognizing that these advantages depend on using AI appropriately and understanding its limitations.

4.1.3 Discussion: AI as a Task Augmentation Tool

The overwhelmingly positive task-level perceptions observed here align with the Technology Acceptance Model (TAM) in that high **perceived usefulness** of a technology tends to drive its acceptance and continued use (Davis, 1989). In our case, the mean agreement scores well above 4.0 for productivity and time-saving benefits demonstrate that respondents find AI highly useful in their workflow. This, in turn, likely fuels their ongoing enthusiasm to integrate AI into their tasks.

The finding that 94% report time savings and 84% report higher productivity suggests AI is successfully **augmenting human capabilities** rather than merely automating tasks. This pattern supports the emerging narrative of “**AI as an amplifier**” of human work, as opposed to “AI as a replacement.” Recent workplace studies have shown similar outcomes – for example, a field experiment with customer support agents found that access to a generative AI assistant increased productivity by

14% on average, with gains over 30% for the least experienced workers (Brynjolfsson et al., 2023). Such evidence underscores how AI can disseminate best practices and amplify performance, particularly by handling routine aspects of work and enabling humans to focus on higher-value activities. In our sample, AI's role in taking over mundane tasks and reducing errors is a clear illustration of augmentation: participants experienced AI as a **collaborative tool enhancing their effectiveness**.

It is noteworthy that even as AI handles repetitive work, participants did not report it marginalizing their roles – instead, they credited AI with empowering them to be more creative and strategic. This reflects a **complementary dynamic** consistent with the idea that AI technologies, at least in their current form, primarily serve as assistants that extend human capabilities. Of course, the caveat mentioned by some respondents is important: AI is not perfect, and users must remain vigilant about its outputs. The need for human oversight and quality control means that **task-level trust** in AI is earned, not given unconditionally (a theme revisited under RQ4's ethical concerns). Overall, the evidence from RQ1 suggests that when AI tools are well-aligned with job tasks and perceived as useful, they function as powerful aids that **increase efficiency, reduce drudgery, and improve quality**. These benefits likely contribute to employees' willingness to embrace AI, setting a positive tone that carries into broader attitudes about AI in the workplace.

Our quantitative and qualitative results are consistent with **field** and **experimental** evidence that generative AI can increase throughput and quality, especially on drafting and summarization tasks. The **task-technology fit** appears high for routine text generation and information retrieval, with respondents reporting time-savings (**94%**) and productivity gains (**84%**).

4.2 RQ2 – Perceived Job-Level Risks

4.2.1 Quantitative Results: Limited Personal Threat Perception

In stark contrast to the strong benefits reported for daily tasks, the survey results showed relatively low levels of personal risk or threat perceived from AI in respondents' jobs. Most participants did **not** feel that AI endangers their employment or well-being at work. For instance, a majority (**56%**) disagreed or strongly disagreed with the statement *"I feel anxious or worried that AI could one day replace my job or many of my responsibilities"* (Q22). Only about **28%** of respondents expressed any level of agreement (agreeing or strongly agreeing) that they worry about being replaced by AI, with the remainder being neutral. The mean score for this job-replacement anxiety item was 2.65, well below the midpoint of the 5-point scale – indicating overall disagreement with the notion that AI will make their own roles redundant.

Similarly, most respondents did not report increased stress due to AI's introduction in the workplace. Over **70%** disagreed that *"the introduction of AI in my workplace has made my job more stressful"* (Q23), while only around 15% felt any increase in stress (the rest were neutral). The average rating for the stress item was 2.25 out of 5, which corresponds to a clear tendency to disagree that AI

adds to their job stress. In short, the prevailing sentiment in this sample was that AI has not harmed respondents' job security or working conditions. If anything, as noted earlier, many felt AI has **improved** their jobs (recall that roughly 80% described AI's overall impact as positive in Q25). Consistent with the job-level framing of RQ2, **80%** agreed the overall impact of AI on their job is positive (Q25).

It is useful to view these results in context. Given that the vast majority of participants are actively using AI and seeing positive outcomes (per RQ1), it follows that they might feel more secure rather than threatened – they are incorporating AI as a helpful tool in their work. Indeed, in this study's sample, those who use AI frequently and have high confidence in their ability to work with it (see RQ3) did not report greater fear of job loss; familiarity and self-efficacy with AI appear to buffer anxiety about its impacts. This overall lack of fear is notable because popular discourse often emphasizes the threat of AI “stealing jobs.” In our results, however, outright pessimism about AI-driven unemployment was limited to a minority. For example, only one respondent in the entire sample **strongly agreed** that they worry about being replaced by AI. Instead, most respondents appeared to view AI as a complement to their own skills or as a tool that assists them, rather than as a replacement.

This cautious optimism stands in contrast to some general workforce surveys: for context, a recent Pew Research Center study (Lin & Parker, 2025) found that about half of U.S. workers (52%) are worried about AI's future impact on the workplace, whereas our sample's level of worry was much lower.

In summary, the quantitative data for RQ2 show that direct perceived risks to one's current job are **low for the majority** of individuals in this sample. Most participants felt their roles were safe and even improved by the introduction of AI, and they did not experience increased stress due to AI.

4.2.2 Qualitative Insights: Nuanced Concerns About Broader Impacts

The open-ended responses about **concerns or fears** regarding AI in the workplace provide further insight into RQ2. While the quantitative data showed a majority feeling unthreatened personally, the qualitative comments reveal that job displacement is indeed a salient fear for a notable subset of respondents. Several participants explicitly mentioned concern that AI could render certain roles, or skill sets obsolete in the long run. For example, one respondent warned about “*the extinction of certain areas and jobs that are more procedural and easily replaced by AI.*” Another echoed that sentiment, expressing worry about potential layoffs if companies assume “*concerned about layoffs simply on the base of a possibility that AI can do the job, when it can't yet.*”

These comments reflect a fear that management might **prematurely cut staff** under the expectation that AI will take over those responsibilities, even if the current AI technology is not fully capable of doing so. This nuance highlights that even respondents who were not personally anxious on

a day-to-day basis still recognized a broader risk: organizational decisions (perhaps driven by cost-cutting or hype around AI) could lead to job losses if AI's capabilities are overestimated or misused.

On the other hand, a few respondents took the opposite stance, indicating little to no fear about AI in their own careers. One participant wrote, *"I'm not worried about it – I don't think AI will ever fully replace our role."* This perspective aligns with the majority quantitative finding that most did not anticipate being replaced; it suggests that those well-versed in their field feel human expertise will remain essential despite AI advances. Between these two viewpoints, some respondents pointed out more **indirect or long-term job risks**. For instance, one comment raised the concern that if people rely too heavily on AI, it could *"reduce the number of entry-level jobs and impact overall competence,"* which in the long run might harm career development and the talent pipeline. This is an interesting observation: AI might not eliminate one's own job immediately, but it could shrink the opportunities for newcomers or alter the skill requirements for junior roles, thereby affecting the job landscape over time. Another respondent mentioned the potential *"impact on the economy as more jobs are replaced by AI,"* thus extending the concern to a society-wide scale.

In summary, the qualitative findings for RQ2 illustrate that **job security concerns do exist** within the workforce, even if they were not dominant in this sample. The divergence in view often corresponded with individuals' sense of adaptability: those actively learning and integrating AI into their skill set (as many in this study were) tended to feel secure, whereas those focusing on AI's potential to automate jobs were more likely to express worry.

Notably, even the experienced AI users in our sample acknowledged that organizations should handle AI-driven changes carefully. The comments serve as a reminder that employers need to manage the transition to AI thoughtfully: transparent communication and upskilling opportunities can help dispel unfounded fears, and at the same time ethical workforce planning is needed to avoid premature or unjustified job cuts. Thus, while task-level benefits are clear and personal job-level risks were perceived as minimal by most, there remains an undercurrent of concern about **longer-term employment implications** that merits attention from organizational leaders and policymakers.

4.2.3 Discussion: Experience and Familiarity as Risk Mitigators

The findings for RQ2 suggest that **experience with AI and a sense of personal competence can mitigate fear** of job displacement. Our participants, being largely tech-sector professionals comfortable with AI, exemplify how familiarity breeds confidence rather than insecurity. This aligns with broader technology adoption research: individuals who possess higher self-efficacy and understanding of a new technology are less likely to view it as a threat. In contrast, anxiety about job replacement tends to be higher among those who lack knowledge or feel unprepared for technological change.

The low levels of AI-related anxiety in our sample can thus be seen as a product of both **direct exposure** to AI (92% were current users) and a proactive mindset toward learning (elaborated in RQ3). This dynamic reinforces the idea that **adaptation is key** – workers who engage with AI and update their skills are more inclined to see AI as a tool for empowerment rather than a competitor for their job.

It is important to note that our results do not imply that all fears of AI are unwarranted, but they do challenge some of the more dire predictions when looking at a group actively using AI. Many public opinion polls report much higher levels of concern about AI-driven unemployment than we found, and this discrepancy likely reflects differences in knowledge and context. For example, surveys show substantial portions of workers fearing their jobs will become obsolete due to AI, yet those already leveraging AI may realize the technology’s current limits and the continued value of human expertise.

In our study, even those who voiced concern often framed it as a **conditional risk** – essentially saying “AI could replace jobs if mishandled or if we don’t keep up,” rather than “AI will inevitably replace my job.” This suggests a *cognitive reframing* among informed employees: they view AI as manageable through proper skill development and organizational strategy.

Another notable point is the emphasis some respondents placed on **indirect impacts**, such as the reduction of entry-level roles or erosion of certain skills over time. This indicates that even if people are not personally afraid for their immediate job, they are contemplating the broader workforce evolution. Such reflections connect RQ2 with the ethical and implementation concerns discussed in RQ4 – for instance, if AI is adopted without guidelines (addressed later), it could lead to unfair job cuts or misuse, which in turn feeds job insecurity.

In essence, **ethical AI use and workforce planning** become part of the risk equation. Experienced professionals may feel secure now, but they still expect companies to introduce AI responsibly to ensure *everyone* – including less experienced colleagues – can thrive. This perspective is consistent with a “cautious optimism”: our participants are optimistic because they feel equipped, yet cautious because they recognize outcomes depend on **how AI is implemented**. For organizations, the lesson is that investing in employee readiness and maintaining transparency in AI-driven changes can greatly alleviate fears and foster an environment where AI is viewed as a benefit rather than a threat.

While a minority (**28%**) reported anxiety about job replacement, most respondents (**56%**) disagreed that their role would be replaced aligning with exposure studies that foresee **transformation/augmentation** more than wholesale automation outside clerical domains.

4.3 RQ3 – Skill Development and Organizational Support

4.3.1 Quantitative Results: High Self-Efficacy Amid Training Gaps

RQ3 focuses on how individuals perceive the need for skill development in response to AI and how supported they feel by their organizations in this learning process. The survey results indicate an

overwhelmingly **proactive stance** among respondents regarding upskilling for AI, coupled with generally positive views of organizational support – albeit with some gaps in formal training provision.

A striking approximately **91%** of respondents agreed that they “*feel confident in [their] ability to work with AI technologies*” (Q27), reflecting a high level of self-efficacy (mean score = 4.24). In other words, most participants believed they could learn and adapt to new AI tools as needed. In fact, a large majority were not just confident but **actively taking initiative**: **84%** agreed that “*I am actively seeking to learn new skills (e.g., AI tools) to remain relevant as AI becomes more common in my job*” (Q30; mean = 4.15). Virtually no one strongly disagreed with that statement, indicating that very few respondents were passive or resistant about learning new AI-related skills.

These findings paint a picture of a workforce keenly aware of the importance of continuous learning in the age of AI and ready to invest effort in updating their competencies.

From the organizational side, respondents generally reported that their employers are **encouraging** about AI skill development. About **85%** agreed that “my organization’s management is supportive of employees learning and using AI” (Q28; mean = 4.09). This suggests that many employees feel their leadership encourages AI adoption and innovation, at least in principle. However, when it comes to **concrete resources and training opportunities**, the ratings were slightly less positive. Roughly two-thirds of respondents (around **65%**) agreed that they have “access to sufficient training or resources to learn about AI tools if I need to” (Q29).

This leaves about one-quarter (around 24 %) who disagreed (or strongly disagreed) that they have sufficient training resources, with the remainder neutral or unsure. While 35 % of respondents indicated they had received no formal employer-provided AI training (Q11), 65 % nonetheless agreed that sufficient resources were available for self-directed learning (Q29), indicating that many relied on informal or self-initiated approaches rather than structured programs. The average rating for perceived training access was 3.67, notably lower than the other items in this section – suggesting that while encouragement is high, the practical support in terms of formal training and resources is not uniformly available.

Additional information from the survey reinforces this interpretation: only a small minority of respondents (approximately 15%) reported that they had received extensive AI-related training from their employer, whereas over one-third said they had **no formal training at work** and had learned mostly on their own. Many others indicated they had only minimal or one-off training sessions. In short, respondents felt confident and were personally driven to upskill for AI, and they perceived management to be broadly supportive of these efforts – but **formal training opportunities have lagged behind** employees’ enthusiasm.

These results have important implications. The employees' high confidence and active pursuit of AI skills is a very positive sign; it implies that fear of AI (as noted in RQ2) can be mitigated by knowledge and empowerment. Indeed, in this study those who felt more confident with AI tended to be less worried about negative impacts on their jobs.

Similarly, the strong management encouragement reported (85% supportive) bodes well for creating a culture open to innovation. Yet, the fact that 1 in 4 respondents did not feel they have adequate training resources points to an area for improvement.

While many organizations express support for workforce development in principle, employees often perceive a gap between this stated commitment and the availability of concrete learning initiatives or AI-specific training opportunities. This discrepancy between organizational encouragement and the actual provision of structured training is not uncommon; numerous studies indicate that employees frequently seek more formalized development pathways for emerging technologies than are currently offered. In the present study, although self-directed learning was widespread, many respondents explicitly expressed a desire for greater institutional support.

Enhancing such support mechanisms could facilitate more effective AI integration, as the workforce already demonstrates a strong willingness and motivation to engage in skill development.

4.3.2 Qualitative Insights: Calls for Training, Guidelines, and Inclusive Implementation

The open-ended suggestions that participants offered for how organizations should handle AI's introduction shed light on the kinds of support employees value. A dominant theme in these responses was the call for more **training and education** on AI. Participants repeatedly recommended that their organizations invest in comprehensive training programs, continuous learning opportunities, and AI literacy initiatives to help the workforce adapt. For example, one respondent argued that AI is evolving so rapidly that *"it's insufficient to offer only one introductory course"*; instead, organizations must *"invest more in employee education regarding AI."*

This sentiment that a one-off training session is not enough, and ongoing upskilling is necessary was commonly expressed. Others suggested very practical steps: *"Create workshops or courses within the company and provide the necessary tools,"* one participant wrote, while another advised managers to *"promote the use of AI and provide training on it."* There was also an emphasis on ensuring a baseline understanding across the organization: one respondent stated that *"everyone who uses AI tools should have mandatory training,"* emphasizing that employees need to understand the tools they use and be aware of their limitations. Taken together, these comments strongly indicate that employees want their organizations to take a more active role in building **AI-related skills**, beyond just encouraging them in words.

The workforce is asking for structured learning opportunities – such as workshops, courses, and hands-on training sessions – to properly equip them for working with AI.

In addition to training, clear guidelines and policies for AI use emerged as a notable theme. Many respondents felt that companies should provide explicit guidance on how AI should be implemented ethically and effectively. One individual emphasized the need for *“clear ethical guidelines,”* observing that AI has *“grown so fast that rules and regulations couldn’t keep up... companies should set their own guidelines”*. Similarly, others spoke about establishing guardrails for AI use. For instance, a participant urged management to *“make clear guidelines on how to use AI with guardrails,”* cautioning against a free-for-all adoption where *“every individual adopts their own tools”* without oversight.

The concern here is that without proper guidelines, AI might be misused – potentially leading to issues like data leaks or inconsistent practices. Indeed, a few respondents specifically mentioned the need to train employees on what **not** to do with AI, such as avoiding feeding confidential information into external AI systems (to prevent risking the company’s intellectual property). These suggestions indicate that employees are not only interested in learning how to use AI but also want to know the boundaries and best practices to use it responsibly and securely.

Another recurring suggestion was to **involve employees in the AI adoption process** and change management. Rather than having AI tools imposed from the top down, respondents advocated for a more inclusive approach. *“The integration of AI should not be imposed but added with initiatives that allow employees to see the benefits of it,”* one participant wrote, highlighting the importance of demonstrating value and getting buy-in from staff. Others echoed that management should *“help employees navigate the change so they can extract the best of the technology”* and *“involve employees in AI initiatives”*.

This implies that organizations should communicate openly about AI projects, solicit feedback from employees, and perhaps allow interested staff to pilot new AI tools. By doing so, employees can feel a sense of ownership and become champions of AI rather than feeling that AI is something being forced upon them. Some respondents pointed out that when employees understand the benefits of AI firsthand, they are more likely to embrace it. Notably, a few participants had no suggestions for improvement because, in their view, *“my organization is really proactive with AI, and supports its utilization and exploration... I don’t have any suggestion [for improvement].”* However, such positive cases were the exception; most respondents saw room for their organizations to do more in terms of training, guidelines, and inclusive practices.

Overall, the qualitative insights for RQ3 show that employees are clearly signaling what they need: **continuous training, clear usage policies, and involvement in AI rollouts**. These measures

were seen as key to helping staff adapt and feel comfortable with AI. The emphasis on ethics and guardrails also connects with RQ4's focus on concerns – employees want to be sure that as they adopt AI, they do so in a way that is safe, fair, and aligned with company values.

4.3.3 Discussion: Bridging the Support Gap for an AI-Ready Workforce

The findings for RQ3 reveal a workforce that is **eager to develop AI-related skills** and generally positive about organizational encouragement, but also one that clearly wants more tangible support to accompany the introduction of AI.

On one hand, the high levels of self-efficacy and proactive learning (over 80% actively upskilling) reflect an empowered employee base that is taking charge of its own development. On the other hand, the calls for better training and guidelines highlight a gap between **enthusiasm and infrastructure**. In many ways, this situation is a testament to employees' foresight: they are preparing themselves for AI's rise, yet they recognize that without institutional backing – in the form of structured training programs and clear policies – the full benefits of AI might not be realized or could even backfire (through misuse or uneven adoption).

The scenario described by our participants is echoed in broader trends. For example, while a majority of workers see AI affecting their jobs, only about one-third say their employer provides AI training. This lag in formal training despite increased AI use is a common challenge, and our data suggest it is present even in tech-forward workplaces. Bridging this gap is crucial: organizations that **invest in upskilling** are likely to see smoother AI integration and a more confident workforce.

The strong desire for clear guidelines and ethical guardrails indicates that employees understand the dual nature of AI: it offers great opportunities but also comes with risks that must be managed. By asking for guidelines, employees are effectively asking leadership to **set the boundaries and expectations** for AI use. This is a positive sign – it shows that employees want to use AI responsibly and are looking to their organizations for direction. Establishing such guidelines can also enhance trust (linking to RQ4 concerns about ethical use): when employees know the rules of the road, they are less likely to fear unseen pitfalls, such as inadvertently violating privacy or losing control of data.

In essence, proper training and policies act as **confidence boosters** that complement the existing high self-confidence of employees. Management support, which was reported to be high in principle, needs to translate into these concrete actions. Encouragingly, our respondents seem ready to embrace any opportunities provided – they are effectively saying, “Give us the tools and knowledge, and we will run with it.”

Another important aspect from RQ3 is the emphasis on inclusive change management. This highlights a broader lesson: **technological change is also a human change**. Employees wish to be partners in the AI journey, not passive recipients of new tools. The benefits of involving employees are

well-documented in change management literature – it increases adoption, reduces resistance, and surfaces practical insights from the front lines. Our participants’ suggestions to involve staff in AI initiatives and to demonstrate AI’s benefits align with best practices for introducing new technology in the workplace. Such involvement not only aids skill development (as employees learn by doing) but also builds a culture of innovation where employees feel their voices matter. When people see that their feedback is valued and that training is provided, their buy-in increases, which in turn makes the implementation of AI more effective.

High **self-efficacy (approximately 91%)** and **upskilling intent (84%)** suggest readiness to adapt. This mirrors external workforce surveys that find strong interest in AI training and a perceived gap in formal provision.

In conclusion, RQ3’s findings underscore that **human capital development and organizational readiness** are critical in the AI era. The workforce in this study is not shying away from AI – quite the opposite, they are hungry to learn and generally supported by management encouragement. However, there is a clear call to action for organizations: match the employees’ proactive spirit with robust training programs, clear guidelines, and participatory implementation strategies. By doing so, companies can ensure that their employees not only feel confident in theory but are truly equipped in practice to harness AI’s potential. This will also help address some of the ethical and trust issues (RQ4) by making sure everyone knows how to use AI correctly and ethically.

Ultimately, investing in people is the surest way to translate AI’s promise into productivity gains and innovation on the ground.

4.4 RQ4 – Ethical Concerns and Future Adoption Intentions

4.4.1 Quantitative Results: Caution Meets Optimism

The final research question examines two related areas: (a) individuals’ **ethical or moral concerns** regarding AI in the workplace, and (b) their **intentions or openness** to adopting AI tools in the future.

The survey results suggest a blend of caution and enthusiasm – respondents did have some ethical reservations about AI, yet they remained broadly enthusiastic about AI’s role in their work and were keen to use more AI going forward.

On the ethics side, responses to the statement “*I am concerned that AI might be used in ways that are unfair or unethical at my workplace*” (Q21) were mixed. Around **45%** of participants agreed (to some degree) that they held such concerns, whereas roughly one-third (**33%**) disagreed that they were concerned (the remaining approximately 22% were neutral or undecided). The average rating for this item was about 3.15, near the midpoint of the scale, indicating a fairly divided stance overall. In other words, nearly half of the sample did perceive potential ethical issues with AI at work – a

substantial proportion – even though these same individuals were generally positive about AI’s practical benefits in earlier questions.

These ethical concerns likely stem from uncertainties about **how AI is being used and how its decisions are made**. Indeed, related items in the survey shed light on this. Only about **47%** of respondents agreed that *“AI-related decisions in my workplace are transparent to me”* (Q20), with the rest either not finding those decisions transparent or being unsure about it. Furthermore, trust in AI outputs appeared lukewarm. When asked if they *“trust the results or recommendations that AI systems provide”* (Q18), only **45%** of respondents agreed, while about 22% disagreed and the remainder were neutral. The mean score for trust in AI was 3.26, suggesting that on average respondents were only somewhat trusting of AI – essentially on the fence. Additionally, a significant number of participants admitted to confusion about AI’s decision-making: approximately **35%** in total agreed that *“I find it confusing or difficult to understand how the AI makes decisions”* (Q19), whereas about half of respondents disagreed that they were confused (and the rest neither agreed nor disagreed).

Taken together, these findings point to notable **ethical and transparency concerns**. A large subset of respondents had reservations about the opacity of AI systems and the potential for unfair or improper use of AI, which in turn limited their full trust in the technology.

Despite these concerns, the **adoption intentions** among respondents remained very high. In fact, enthusiasm for AI was one of the defining characteristics of this group. Approximately **75%** of respondents agreed that *“I am enthusiastic about the increasing use of AI in my field of work”* (Q26).

An even higher proportion expressed interest in expanding their own use of AI: about **85%** agreed with *“In the future, I am interested in using even more AI tools or features in my work”* (Q31). The mean scores for these items were correspondingly high (around 4.0 or above). Very few people were opposed – for instance, only 4% stated that they are not interested in using more AI going forward. This indicates that even among those who do have ethical qualms, most still want to continue with, or even increase, their use of AI tools.

In general, participants seemed to distinguish between acknowledging issues and **rejecting the technology**. They remained eager to leverage AI’s benefits, presumably expecting that any challenges can be managed or mitigated with the right approaches. The combination of high perceived usefulness of AI (demonstrated by the strong task-level benefits in RQ1) and only moderate levels of trust suggests that, for these respondents, the clear advantages of AI outweigh their reservations when it comes to deciding whether to use the technology. In other words, while lack of complete trust and ethical concerns might act as a brake, it has not prevented them from being enthusiastic adopters. This mixture of caution and optimism reflects a mature perspective: participants are not blindly accepting of AI – they are aware of real risks – but they are generally choosing to engage with it and are optimistic about its value in their work, provided those risks are addressed.

4.4.2 Qualitative Insights: Transparency, Data Privacy, and Responsible Use

The open-ended comments provide concrete examples of the ethical concerns that respondents had, complementing the quantitative findings. A number of participants voiced worries about **data privacy and security** when using AI. For instance, one respondent wrote that their biggest concern was *“the leak of information can happen”* through AI systems, reflecting a fear that using AI tools (especially third-party or cloud-based AI services) could lead to confidential or sensitive company data being inadvertently exposed. Another participant similarly noted that when utilizing AI at work, people often *“lack control of the amount and type of information we provide to the AI,”* highlighting a sense of unease about what happens to data once it’s fed into an AI system.

These remarks underscore a common ethical concern in the AI domain: employees worry about **how data is handled**, who might have access to it, and whether using AI could compromise proprietary information.

Another major theme in the comments was the **lack of transparency or explainability** of AI decisions. One participant bluntly stated that the *“lack of transparency of AI’s decision process”* was their chief concern regarding workplace AI. This aligns with the earlier survey result where many did not find AI decisions to be transparent. The idea of AI as a “black box” – producing recommendations or decisions without a clear explanation – clearly unsettled some respondents. Relatedly, several individuals were concerned about the **accuracy and potential bias** of AI outputs. For example, one comment expressed skepticism about AI’s reliability, noting that AI *“can make many mistakes and inaccurate statements.”* This points out the risk that if organizations place too much trust in AI without verification, they might act on faulty recommendations or **perpetuate biases** present in the algorithms or training data. While the respondent did not use the word “bias,” the mention of inaccurate outputs raises the implicit concern that AI could mislead or produce unfair outcomes if not properly monitored.

In summary, transparency, explainability, and accuracy (including bias) were prominent issues that our participants associated with the ethical dimension of AI at work.

Some respondents also mentioned broader or more **long-term ethical concerns**, though these were less common. A couple of participants alluded to science-fiction scenarios – one half-jokingly cited *“Skynet”* (the fictional rogue AI from *The Terminator* films) and mused about the possibility of AI becoming so advanced that it no longer *“needs us.”* This reflects an existential risk perspective (albeit expressed humorously), where the concern is AI eventually posing a threat to humanity or to the fundamental need for human workers. Another respondent referenced potential large-scale impacts, worrying about *“the impact on the economy as more jobs are replaced by AI”*.

This comment bridges ethical concerns with job security issues (as in RQ2), pointing to a societal-level risk rather than just personal risk. While these extreme views were not widespread in the

data, their presence shows that a few individuals do consider long-term and big-picture implications of AI, beyond the immediate practical and ethical issues in their own workplace.

In terms of **adoption intentions**, the qualitative feedback generally echoed the survey's high enthusiasm, with an emphasis on adopting AI **responsibly**. Many participants, in their suggestions to organizations (as discussed in RQ3), implicitly addressed how to encourage AI use while managing its risks. Even respondents who voiced ethical concerns tended to frame those concerns as issues to be solved so that AI can be used safely, rather than as reasons to avoid AI entirely. For example, some respondents argued that companies should facilitate experimentation with AI and *"amplify [AI] use cases across business lines,"* but at the same time not be overly restrictive – *"not try to control it too much"* – as long as proper guidelines are in place. This suggests that employees want relatively open access to AI tools in their work, provided the organization offers guidance and safeguards to prevent misuse.

Overall, participants seemed to advocate for **embracing AI and its opportunities, while putting measures in place to ensure that its use is ethical, transparent, and well-regulated**.

The tone of the comments was not anti-AI at all; instead, it was, *"let's use AI, but let's do it right."* As one respondent succinctly advised, the goal should be to *"help employees navigate the change so they can improve their work with [AI]"*.

This captures the dual emphasis on moving forward with AI adoption and doing so in a supportive, well-managed way.

4.4.3 Discussion: Balancing Enthusiasm with Ethics and Trust

The findings for RQ4 reveal a workforce that is **highly open to adoption of AI** even as it harbors certain reservations that must be addressed. This combination of enthusiasm and caution can be seen as a paradox of "cautious optimism." On one side of the coin, the overwhelming majority of participants want more AI in their work and are excited about its potential – a testament to the clear productivity and quality benefits they experience (as documented in RQ1). On the other hand, nearly half have noted ethical concerns, primarily around transparency, fairness, and privacy.

The key insight is that these employees are not letting their concerns stop them; instead, they are likely to **advocate for solutions** to those concerns. This perspective is heartening because it suggests employees believe that *with the right safeguards*, AI's benefits can be enjoyed without succumbing to its risks.

The concerns highlighted align with well-known themes in AI ethics and governance. Lack of transparency (*"black box"* algorithms) and data privacy issues are recognized barriers to fully trusting AI systems. Research in human-AI interaction often stresses that trust in AI hinges on factors like explainability, reliability, and perceived fairness. Our participants' moderate trust levels (only

approximately 45% outright trusting AI outputs) reflect this reality – no matter how helpful AI is, if people don't understand how, it works or fear it could behave unfairly, they will reserve some skepticism. Yet, trust is not a static trait; it can be **built or eroded** based on how AI is implemented.

Transparency measures, user education, and ethical guidelines (all of which our respondents called for) are tools to build trust. In fact, by asking for clarity and rules, employees are effectively saying they want to trust AI more, and they need organizational support to do so. It is also noteworthy that **distrust in AI can hinder its full utilization**, which means addressing these ethical concerns isn't just a moral obligation but a practical one for companies. If left unaddressed, issues like opaque decision-making could lead to employees under-utilizing AI or double-checking it so much that efficiency gains are lost. Thus, balancing enthusiasm with ethics is crucial: sustained AI adoption will require that employees feel the technology is **trustworthy** and aligned with their values.

It is encouraging that despite some “worst-case scenario” mentions (like the Skynet joke or economic displacement worries), the overall sentiment was far from technophobic. Even those who imagined negative outcomes framed them as hypotheticals or distant possibilities. The prevailing attitude was proactive: *use AI, but guard against pitfalls*. This implies that organizations have a willing workforce that just needs the right environment to flourish with AI. Ensuring data security, being transparent about where and how AI is used, and involving employees in creating ethical AI policies are likely ways to reinforce trust. Moreover, these ethical considerations tie back to earlier findings – for instance, an absence of clear guidelines (as lamented in RQ3) could exacerbate ethical worries, whereas strong training on “what not to do” with AI could alleviate fears of data leaks.

Similarly, if ethical concerns are left unaddressed, they could eventually feed into job-level anxieties (RQ2) by creating suspicion around AI initiatives. Conversely, addressing ethics can further boost the positive task-level experiences (RQ1) by ensuring that AI tools are not only useful, but also fair and safe to use.

In summary, RQ4's results show that individuals are generally **very open to AI's increasing role**, provided that its implementation is handled responsibly. They are **enthusiastic adopters** of AI who nonetheless expect their organizations to be stewards of ethical AI practices. The dual message to employers and leaders is: *keep pushing forward with AI* (because employees see its value and want more of it) but *do so with careful attention to ethics and trust*. By actively addressing transparency, privacy, and fairness concerns – through measures like explainable AI, robust data protection, and ethical use policies – organizations can foster an environment where employees' optimism is rewarded, and their caution is respected. In doing so, the full potential of AI can be realized in a way that maintains morale, trust, and alignment with organizational values.

The “confidence–concern” paradox we observe (high self-efficacy yet 45% ethical concerns) is well-explained by the JD-R lens (AI as resource and demand) and by literature on explainability/auditing: trust increases when decisions are explainable, audited, and human-reviewable in high-stakes contexts.

4.5 Overall Summary: A Paradox of Enthusiasm and Caution

Overall, the findings across the four research questions paint a comprehensive picture of how this group of AI-experienced professionals views the technology’s impact on their work.

In many ways, their perspective can be characterized as “**cautious optimism**” – a blend of high enthusiasm for AI’s benefits with a grounded understanding of its challenges.

The results indicate several key patterns:

1. **Task-level success drives acceptance:** Participants report strong productivity and efficiency gains from AI (RQ1), which creates a foundation for their positive attitudes. When AI is seen to *clearly augment daily work* – saving time, reducing errors, and handling drudgery – employees are naturally inclined to welcome it. This highly perceived usefulness has likely been a critical factor in driving the high adoption and interest in using AI more (RQ4). In short, visible success in day-to-day tasks is reinforcing employees’ willingness to embrace AI.
2. **Experience buffers anxiety:** The people in this sample, being familiar with and skilled in using AI, exhibited **lower job insecurity and stress** related to AI (RQ2) than what might be expected from general surveys of the broader workforce. Their direct experience with AI appears to mitigate fear – they view AI as a tool rather than a threat. This suggests that *familiarity and adaptability* are powerful antidotes to the fear of displacement. Those who understand how to work alongside AI tend to see it as a complement to their role, whereas fear thrives more on the unknown.
3. **Support gaps persist amid proactive upskilling:** Employees are **actively upskilling and confident** in their ability to learn AI (RQ3), demonstrating a proactive approach to career development in the AI era. They also perceive management as supportive of AI innovation in principle. However, a notable gap exists in formal support: access to training and clear guidelines is not keeping pace with employees’ enthusiasm. In other words, many workers are essentially training themselves and calling for their organizations to catch up. Bridging this support gap – through robust training programs, clear AI usage policies, and involving employees in AI initiatives – is crucial for sustaining the momentum and ensuring that the workforce is truly AI-ready.

4. **Ethics and trust matter for sustained adoption:** Despite their enthusiasm, employees have **significant ethical concerns and trust issues** regarding AI (RQ4). Areas like transparency of AI decisions, data privacy, and potential bias are on their radar. These concerns act as a moderating force – they don't stop people from using AI, but they do inject caution into how people use it and how much they rely on it. Building and maintaining trust in AI systems will be essential for the long-term success of AI integration. This means organizations should implement AI responsibly, with efforts to make AI processes explainable and fair, and to communicate openly about AI's use. By addressing these ethical dimensions, companies can convert employees' cautious optimism into full confidence in AI-driven processes.

Overall, the study suggests that for this sample of professionals, AI is seen as a valuable tool to be embraced – but with eyes open. Participants are neither blindly optimistic nor doom-and-gloom about AI; rather, they appreciate its advantages while advocating for steps to manage its downsides.

This balanced viewpoint challenges the often-polarized public narratives about AI. It shows that with proper education, experience, and organizational support, employees can hold a nuanced position: excited about what AI can do for them, yet mindful of what needs to be done to implement it right.

The insights from these findings offer important implications for organizations and individuals navigating the evolving landscape of AI in the workplace. In the next chapter, we will build on these results to draw overarching conclusions and offer recommendations. We will discuss how companies can leverage the positive perceptions of AI while addressing the highlighted concerns, and we will suggest avenues for future research to further understand AI's impact on employment. The goal is to ensure that the integration of AI into work continues to enhance productivity and job satisfaction, without compromising ethical standards or employee well-being.

5. Discussion

This discussion chapter brings together the major insights of the study, interpreting the empirical results from Chapter 4 in the light of the theoretical and empirical literature reviewed in Chapter 2. It highlights how the findings address each research question and what they mean for understanding AI adoption and its impact on employees. The chapter synthesizes the key findings by research question, discusses the study's theoretical contributions, offers practical recommendations for organizations, employees, and policymakers, and acknowledges the main limitations and avenues for future research.

Throughout, the chapter integrates evidence from both the quantitative survey results and qualitative responses, capturing the complexity of professionals' perceptions regarding AI's effects on daily tasks, job security, skills, and ethical considerations.

5.1 Key Findings

The study addressed four main research questions (RQs) concerning the perceived impact of AI on employees' work. The key findings for each RQ are summarized below:

- **RQ1 (Productivity and Performance):** *How does AI integration affect employees' productivity and work quality?*
 - The survey results indicate that AI has a strongly positive impact on day-to-day productivity. An overwhelming majority of participants reported efficiency gains from using AI tools, citing significant time savings and higher output quality. For example, 84% of respondents agreed that AI has made them more productive in their jobs, with nearly half strongly agreeing. Only a negligible fraction (around 3%) disagreed, confirming that for most workers, AI serves as an effective efficiency booster in their current roles. Qualitative comments reinforced this, describing how AI automates tedious tasks and frees employees to focus on higher-value work.
- **RQ2 (Job Security and Stress):** *To what extent is AI adoption associated with employees' perceptions of job insecurity and work-related stress?*
 - Most participants (56%) do not feel threatened by AI, viewing it as a complement to their work rather than a replacement. However, a notable minority (28%) do express anxiety that increasing workplace AI could diminish their role or even replace their job, with the remainder being neutral. This indicates a tangible job security concern for a subset of workers, even as the majority see AI as a useful tool. Regarding stress, approximately 15% of employees reported that using AI has made their job more stressful, often due to pressures like learning new systems or monitoring AI outputs (a

form of “technostress”). Conversely, some employees noted relief from AI (around 10% felt it actually reduced their stress by handling menial tasks), and the majority experienced no change or slight improvement in stress levels. In summary, AI’s integration is accompanied by moderate job-insecurity worries and a slight uptick in stress for certain employees, even as most do not feel adversely affected. These nuanced results underscore an ambivalent reality: the same person might enjoy productivity gains from AI yet simultaneously worry about long-term job security.

- **RQ3 (Skills and Training Needs):** *How does AI adoption influence employees’ need to develop new skills, and does employer-provided training support moderate this?*
 - The findings show that AI’s arrival has heightened employees’ awareness of the need to upskill, though many feel confident about adapting. A majority of respondents agreed that implementing AI in their job requires learning new skills or competencies, and workers frequently mentioned the importance of “*continuous learning*” to stay relevant. As one participant noted, “*AI is great, but it means I’ll be a student forever – there’s always a new tool to master*”. At the same time, confidence in using AI is very high – approximately 91% reported feeling confident in their ability to work with the AI tools in their job. This suggests that while employees recognize a skills gap, they are generally optimistic about their capacity to bridge it (indeed, many have already begun self-training). Crucially, the availability of employer training significantly eases the upskilling challenge. Respondents who said their organization provides strong AI training and support tended to report lower anxiety about AI and higher acceptance of it, whereas those without such support felt “on their own” and more apprehensive. Respondents with higher perceived training support tended to be more positive.
- **RQ4 (Ethical and Organizational Concerns):** *What key concerns do employees have about AI’s broader impact (e.g. ethical, social, organizational issues), and do these concerns influence their enthusiasm for AI adoption?*
 - The study found that employees do harbor a range of concerns about AI in the workplace, notably around ethics and job implications. The most frequently cited worry is data privacy and security – staff fear that sensitive personal or company data could be misused or leaked by AI systems. The next major concern is job displacement, echoing RQ2’s job-security issue: even AI supporters sometimes fear that widespread AI use could eventually reduce the need for human labor in certain roles. Other common concerns include algorithmic bias and misuse (the risk of AI making unfair or unethical decisions, or being deployed for excessive surveillance), lack of transparency in AI decision-making (the “black box” problem that undermines trust), and over-

reliance on AI leading to dehumanization of work (erosion of human skills and a more impersonal workplace).

- These issues align with widely discussed challenges in AI ethics and change management. Importantly, the presence of such concerns is linked to employees' openness toward further AI adoption. Participants who voiced multiple strong concerns were generally more cautious or hesitant about expanding AI in their jobs – often adding caveats like “*I support AI if it's properly regulated*”. In contrast, those who reported few or no concerns tended to be the most enthusiastic about adopting more AI, looking forward to new AI innovations at work. Qualitative responses made it clear that unaddressed ethical and job-related concerns can dampen employees' eagerness to embrace AI, whereas building trust and safeguards can foster greater enthusiasm.

5.2 Theoretical Contributions

This research advances the theoretical understanding of workplace technology adoption and its impact on employees through three key contributions:

5.2.1 Extending Technology Acceptance Models

The findings both confirm and extend classical technology-acceptance frameworks in the context of artificial intelligence. While the core constructs of the Technology Acceptance Model (TAM) – **perceived usefulness and perceived ease of use** – remain strong predictors of employees' willingness to adopt AI tools, this study reveals that additional dimensions are essential for explaining acceptance in AI-enabled workplaces. In particular, algorithmic trust, perceived transparency, and technostress (user anxiety arising from continuous adaptation to intelligent systems) emerged as salient factors shaping attitudes and sustained use. These variables influence whether employees regard AI as a reliable collaborator or a source of uncertainty.

Consequently, the results support recent scholarship advocating the expansion of TAM and related frameworks to include AI-specific socio-psychological variables. Integrating trust, transparency, and stress considerations alongside the traditional ease-of-use and usefulness constructs yields a more comprehensive model of intelligent-system adoption – one that reflects the distinctive human-machine dynamics characterizing modern workplaces.

5.2.2 Advancing Automation Anxiety Theory

The results provide empirical evidence for a nuanced understanding of automation-related job anxiety. In line with the Smart Technology, Artificial Intelligence, Robotics, and Algorithms (STARA) framework, the study confirms that awareness of AI's automation potential can adversely influence job attitudes: even employees who actively use AI and recognize its benefits may still harbor concerns

about its broader impact on their roles and future employment. In other words, automation anxiety can coexist with positive usage experiences.

This nuance enriches existing theories of workplace automation anxiety by demonstrating that optimism about AI's day-to-day usefulness does not preclude underlying fears about long-term job security or organizational decision-making. The findings also resonate with job-stress frameworks, particularly the Job Demands–Resources (JD-R) model, which posits that technologies such as AI can act simultaneously as job resources – reducing workload and enhancing performance – and as job demands, creating new pressures or uncertainties.

A substantial minority of employees in the sample exhibited this duality, suggesting that theoretical models must explicitly incorporate contextual and individual factors that determine when AI functions as a resource and when it becomes a demand. This integration advances automation-anxiety theory by positioning AI-related stress not as a uniform outcome but as a contingent experience shaped by personal readiness, organizational support, and ethical climate.

5.2.3 Integrating Sociotechnical Systems Perspectives

The study's insights bridge technology-adoption theory with organizational behavior and ethics perspectives. Participants consistently emphasized that successful AI integration depends not only on individual acceptance of a tool but also on the organizational and social context surrounding its use. Issues of training, communication, fairness, and transparency were raised repeatedly. This emphasis validates the relevance of sociotechnical systems theory for AI implementation, highlighting that deploying AI is as much a human and organizational process as it is a technical one.

The findings reinforce arguments from change-management and ethical-governance literatures, which stress that clear communication, employee involvement, and trust-building measures are essential when introducing advanced technologies. Thus, beyond the classical focus on utility and ease of use, the evidence suggests that comprehensive frameworks for AI adoption must explicitly incorporate human-centric factors such as trust, organizational support, and ethical guidance.

The theoretical contribution here is an integrative perspective: effective workplace AI adoption can only be fully understood by synthesizing insights from technology-acceptance models (TAM/TTF), automation-anxiety theories, and sociotechnical change-management frameworks. Future theorizing should build on this integration to develop more holistic models of workplace AI, capturing both the technical and human dimensions of transformation. Bridging acceptance and worker-centric ethics. By combining task-level performance perceptions with job-level risks, organizational supports, and ethical salience in a single instrument, this study unites the TAM/TTF and JD-R lenses with contemporary human–AI teaming concerns.

The observed paradox – simultaneous enthusiasm and ethical caution among early adopters – extends acceptance models by foregrounding appropriate use and governance as co-determinants of positive orientation, while underscoring the protective role of self-efficacy and managerial resources in mitigating technostress.

5.3 Practical Recommendations

Building on these findings, the study proposes a set of evidence-based recommendations for three key stakeholder groups – organizations and managers, individual employees, and policymakers – aimed at maximizing AI’s benefits while mitigating potential downsides for the workforce.

The recommendations translate the study’s insights into actionable guidance to support responsible, human-centered AI integration across diverse workplace contexts.

5.3.1 For Organizations and Managers

- **Invest in comprehensive training:** Establish continuous and structured AI upskilling programs rather than relying on one-off sessions. This can include formal workshops, online courses, mentoring programs, and protected work time for employees to experiment with and master new AI tools. Continuous learning opportunities signal long-term commitment to workforce development and help close the readiness gap identified in this study
- **Ensure transparent communication:** Actively involve employees in AI-implementation decisions and maintain open dialogue about planned changes. Clearly communicate which AI technologies are being adopted, why they are introduced, and how roles or workflows may evolve. Soliciting feedback and addressing concerns candidly builds trust, reduces misinformation, and fosters a shared sense of ownership over technological transformation.
- **Develop ethical-AI frameworks:** Create clear policies and governance mechanisms for the responsible use of AI. These should encompass data privacy and security, algorithmic fairness, and appropriate levels of human oversight in automated decisions. Publicly commit to ethical-AI principles, conduct regular audits or reviews, and communicate outcomes internally so employees know that safeguards are operational and accountability is shared.
- **Promote human–AI collaboration:** Consistently frame AI as a complement to human work, not a substitute for it. Redesign jobs and workflows so AI handles routine or repetitive tasks while employees focus on creative, analytical, and interpersonal responsibilities. Leadership should reinforce this principle through concrete actions – for example, re-skilling employees for new roles rather than pursuing layoffs tied to automation. Recognize and reward teams that integrate AI effectively, underscoring that human expertise remains indispensable even in highly automated environments.

5.3.2 For Individual Employees

- **Embrace continuous learning:** Professionals should proactively develop AI-related skills and stay informed about emerging tools and practices in their field. Treat AI's rise as an opportunity for personal growth and professional renewal: by regularly updating their skill sets, employees enhance both career resilience and long-term value in an evolving labor market.
- **Engage with organizational change:** When AI initiatives are introduced, employees are encouraged to participate actively – attend the available trainings, experiment with new tools, and adapt workflows accordingly. Providing constructive feedback to managers about what works well and where challenges arise can help shape smoother, more inclusive implementation. Active engagement also helps reduce technostress and allows employees to influence how AI is integrated into their daily work in a positive, collaborative way.
- **Advocate for ethical practices:** Employees should feel empowered to raise questions and concerns about AI use in the workplace, whether related to data privacy, fairness, or the scope of automation. By constructively advocating for transparency and accountability – for example, by asking how algorithms make decisions or by flagging potential biases – employees contribute to a culture of responsible and trustworthy AI adoption. Doing so not only protects individual and collective interests but also ensures that AI implementation aligns with shared organizational values.

5.3.3 For Policymakers

- **Support workforce development:** Governments and public agencies should fund retraining programs, continuing education, and AI-literacy initiatives that strengthen the capabilities of both current workers and jobseekers. These programs should span technical skills (e.g., data analysis, automation tools) and transversal competencies (e.g., critical thinking, ethics, collaboration), ensuring that the labor force is well prepared for AI-driven transformation across industries.
- **Establish ethical and regulatory standards:** Policymakers play a pivotal role in setting clear, enforceable frameworks for responsible AI deployment in workplaces. Standards on data protection, algorithmic transparency, non-discrimination, and worker rights – including requirements for notification and consent when AI monitors performance – can guide organizations toward trustworthy practices. By defining these ground rules, policy can reduce employees' fears of unchecked or unethical AI use while promoting fair and transparent innovation.
- **Promote ongoing research and evidence-based policy:** Public institutions should invest in independent research and pilot programs that track AI's long-term effects on employment, job

quality, and economic inclusion. Systematically studying outcomes such as job creation and displacement, skill evolution, and the effectiveness of adaptation strategies will provide an empirical foundation for future legislation and education initiatives. This evidence-based approach ensures that as AI technologies evolve, governments remain proactive in addressing challenges and disseminating best practices for human-centered AI integration.

5.3.4 Implementation roadmap (time horizons)

Short term (0–3 months): Publish concise acceptable-use guidelines for AI tools and designate an AI product owner. Deliver a baseline micro-training module covering privacy, data security, and appropriate reliance; begin lightweight bias- and quality-spot-checks for any AI-assisted process; and define baseline KPIs (productivity, quality, employee stress, and engagement).

Medium term (3–12 months): Roll out role-specific training and human-in-the-loop SOPs (who reviews what and when). Implement a simple model registry and data-handling rules; schedule quarterly bias-and-drift reviews; pilot two AI use cases with pre/post measures and share lessons learned; and expand change-management communications to include peer champions, internal Q&A clinics, and transparent progress updates.

Long term (12–24 months). Redesign jobs to reallocate routine tasks to AI while enriching human roles. Formalize an AI competency framework tied to learning paths and performance metrics; integrate governance mechanisms (risk assessment, audit trails, and vendor oversight); and measure both workforce outcomes (mobility, skill growth) and organizational outcomes (quality, throughput, customer satisfaction) to guide scaling decisions.

This staged approach complements the stakeholder-specific guidance in Section 5.3 and translates it into a sequenced, capacity-aligned plan that balances organizational ambition with risk tolerance.

5.4 Limitations

No research is without limitations, and it is important to acknowledge the constraints of this study when interpreting the results. The key limitations are outlined below.

- **Self-report bias:** The study relied on participants' self-reported perceptions and attitudes, which may be influenced by social-desirability or recall biases. Respondents might have overstated positive outcomes (e.g., productivity or confidence) or under-reported challenges (e.g., stress or ethical concerns). Because all outcome measures – productivity, stress, and skill development – are based on subjective assessments, the inherent limitations of self-report data must be considered despite the assurances of anonymity and encouragement of honest feedback.

- **Cross-sectional design:** The research captured a single time-point snapshot of employee perceptions, limiting causal inference. Observed associations (for example, between training and lower anxiety) cannot confirm cause-and-effect relationships; unmeasured factors may explain these links. A longitudinal design, surveying the same participants before and after AI implementation, would be required to establish temporal dynamics and stronger causal conclusions – an approach beyond the scope of this dissertation.
- **Sampling and generalizability:** The sample, although adequate for exploratory analysis, was non-random and skewed toward tech-savvy professionals with relatively high AI exposure. Approximately 83 % of respondents worked in technology-related roles, suggesting a bias toward early adopters. Consequently, the findings may not fully represent more AI-cautious or less technologically intensive sectors, such as manual or service-oriented occupations, nor capture differences linked to age, geography, or organizational culture. Results should therefore be interpreted as context-bound rather than universally generalizable.
- **Rapidly evolving context:** The pace of AI advancement means that workplace experiences and attitudes can change quickly. The data reflect conditions in mid-2025; subsequent technological progress, regulatory developments, or societal debates may reshape perceptions. Furthermore, the study focused on current attitudes rather than long-term outcomes such as career mobility or sustained job redesign. As AI adoption accelerates, longitudinal and comparative studies will be necessary to track evolving trends.

Despite these constraints, the research provides valuable empirical insights into how professionals currently perceive AI's influence on their work. Several limitations were mitigated through methodological transparency and the inclusion of open-ended qualitative responses, which helped triangulate and enrich the quantitative results. Taken together, the findings constitute a robust foundation for future research and practical reflection on employee experiences in AI-augmented workplaces.

5.5 Future Research Directions

Building on this work, future research should continue to explore the human implications of AI in employment and address the limitations identified. Several promising avenues are outlined below.

5.5.1 Methodological Advances

- **Longitudinal Studies:** Future investigations should **track changes in employee perceptions and outcomes over time**, observing participants before, during, and after AI adoption in their roles. A longitudinal design – surveying employees prior to an AI tool's introduction and following up at multiple intervals – would reveal how initial excitement or apprehension evolves with experience. Such evidence would help establish the **causal effects** of AI on job

performance, satisfaction, and stress, providing temporal depth that cross-sectional designs cannot capture.

- **Experimental approaches:** Controlled trials or organizational pilot programs could isolate the direct effects of AI tools on work processes and employee outcomes. For example, introducing an AI system to one experimental group while maintaining a comparable control group without AI would enable stronger causal inference by minimizing confounding variables. Experiments could also test specific interventions, such as enhanced training or no-layoff guarantees, to determine how these measures influence employees' acceptance, engagement, and well-being.
- **Mixed-methods integration:** Combining quantitative breadth with qualitative depth would yield a more comprehensive understanding of AI's workplace impact. Surveys can map large-scale trends, while interviews, focus groups, or ethnographic case studies can uncover the underlying mechanisms driving those patterns. Embedding researchers within organizations during AI rollouts, for instance, could provide real-time insights into team dynamics, adaptation strategies, and resistance factors. Integrating these strands would produce a richer, context-sensitive account of not only *what* changes occur with AI, but *how and why* they unfold in practice.

5.5.2 Sample and Context Extensions

- **Diverse populations:** Future research should broaden the participant base to include employees with limited or no direct experience of AI, as well as individuals from a wider variety of industries and occupational roles. Findings from this study suggest that perceptions may differ between active AI users and those newly exposed to, or operating in, sectors less affected by automation. Capturing the views of AI non-users and novices – for instance, through surveys in more traditional industries or among older workers – would yield a more complete picture of workforce readiness and concern. It would also be valuable to examine adoption dynamics beyond the predominantly tech-centric context of this study, such as in healthcare, education, manufacturing, or public-sector employment.
- **Comparative studies:** Comparative analyses can illuminate how **organizational structures and cultures** shape employee perceptions of AI. Future work might contrast **small and medium-sized enterprises with large corporations**, or **fast-moving startups with established firms**, to identify whether differences in resources, hierarchy, or innovation climate influence openness to AI. Similarly, sectoral comparisons – between early adopters (e.g., finance, software) and slower adopters (e.g., public administration, manual trades) – could highlight contextual **enablers and barriers** that affect AI acceptance and its perceived effectiveness.

- **Cross-cultural research:** Expanding inquiry across national and regional contexts will clarify how cultural values, economic structures, and policy environments mediate workers' responses to AI. Employees in countries with robust social safety nets may experience less anxiety about automation than those in more precarious labor markets. Cross-cultural surveys or multi-case studies can reveal variations in trust toward technology, tolerance for uncertainty, and normative attitudes toward automation. Understanding these differences is essential for developing globally relevant, human-centered strategies for AI adoption.

5.5.3 Emerging Questions

- **Long-term career impacts:** Future studies should examine how sustained AI augmentation influences career trajectories and skill development over extended periods. Does intensive use of AI tools accelerate professional growth – by enabling employees to take on more complex tasks and progress to higher-level roles – or does it risk skill atrophy if AI assumes too many routine functions? Longitudinal career-tracking or retrospective analyses could clarify whether AI-exposed professionals advance differently from their peers and what new career pathways emerge in increasingly AI-rich environments.
- **Organizational outcomes:** Further research should assess the effectiveness of different AI-implementation strategies on both organizational performance and employee well-being. Comparative studies might contrast firms that invest heavily in training and maintain no-layoff policies with those that pursue automation primarily for cost reduction. Evaluating indicators such as employee engagement, innovation, productivity, and retention can reveal which approaches yield the most sustainable and mutually beneficial results. Such evidence would contribute to defining best practices for orchestrating AI-driven organizational change.
- **Societal implications:** Broader inquiries are needed to understand AI's cumulative impact at the **societal and economic levels**. Research could investigate whether AI is creating as many jobs as it transforms or eliminates, and which roles are emerging – such as AI maintenance, auditing, or strategic-oversight positions. These questions extend to **ethics and policy**, including how to ensure equitable access to AI's benefits and effective support for workers in transition.

In summary, future research should pursue longitudinal depth, diverse representation, and a strong emphasis on intervention and policy evaluation to extend the baseline established by this study. By doing so, scholars and practitioners will be better equipped to anticipate challenges, leverage opportunities, and ensure that theoretical understanding and practical guidance keep pace with technological change.

5.6 Overall synthesis and final reflections

This research was conducted at a pivotal moment in the evolution of work, as artificial intelligence moves from a novel innovation to an increasingly embedded component of everyday organizational life. The findings depict a workforce that is neither naively optimistic nor irrationally fearful about AI, but rather thoughtfully engaged – recognizing its substantial benefits while remaining mindful of its risks and uncertainties. Above all, one theme emerges consistently across the data and literature: the future of work with AI will depend as much on human decisions and values as on technological capabilities.

The human-centered imperative. Perhaps the study's most important insight is that successful AI integration is, fundamentally, a human-centered challenge rather than a purely technical one. Implementation requires not only advanced algorithms but also intentional management of human factors – training, communication, transparency, and ethical oversight. Participants' experiences underscored that these dimensions are pivotal to acceptance: employees want clarity about how AI is used, confidence that they can master it, and assurance that ethical safeguards are in place. Technological innovation, therefore, must be paired with empathy, dialogue, and support. When organizations treat AI adoption as a human transformation journey rather than a software deployment, they are far more likely to unlock AI's potential for productivity and creativity while maintaining morale and trust.

Beyond the replacement paradigm. The findings also challenge simplistic “human versus machine” narratives. Instead of viewing AI as a substitute for human labor, this research reveals a collaborative paradigm of complementary capabilities. Employees in this study largely perceived AI as a tool that enhances their effectiveness – automating routine or repetitive tasks and freeing them for higher-value, strategic, or creative work. Many participants described tangible improvements in job satisfaction when AI relieved them of tedious duties. This suggests that the relationship between AI and labor need not be zero-sum: when implemented responsibly, AI can elevate the human contribution by shifting work toward tasks requiring judgment, empathy, and innovation. By moving beyond competition and focusing on collaboration between humans and intelligent systems, organizations can create workplaces where technology amplifies human potential rather than diminishing it.

A call for responsible innovation. At this technological inflection point, the study's participants deliver a clear message to leaders and policymakers alike: embrace AI's transformative power, but anchor its deployment in transparency, ethics, and human well-being. The future of work will be determined less by the sophistication of AI tools than by the wisdom and accountability with which they are integrated. This entails setting clear governance frameworks, defining ethical boundaries, and ensuring that technological progress does not outpace the social and moral safeguards protecting what is distinctively human about work. Involving employees in decision-making and providing mechanisms

for oversight and appeal will be essential for sustaining trust. Responsible innovation is not merely a normative ideal; it is a pragmatic necessity for ensuring that AI adoption leads to sustainable, equitable, and trustworthy workplaces.

Looking ahead. As AI becomes further woven into the fabric of work, new challenges and opportunities will arise – ranging from the need to continuously reskill the workforce to addressing algorithmic accountability and ensuring inclusive access to technological benefits. Yet the overarching conclusion of this research remains hopeful. With thoughtful leadership, continuous investment in human capability, and an unwavering commitment to ethical principles, the integration of AI can mark not the displacement of human talent but its renewal. By designing systems that complement rather than compete with people, organizations and societies can realize a model of work where intelligence – both human and artificial – co-evolves toward higher levels of creativity, fairness, and purpose.

In essence, this dissertation affirms that the success of the AI era will not be measured solely by computational speed or productivity gains, but by our collective ability to align technology with human flourishing. The challenge is profound, but so too is the opportunity: to shape an AI-augmented future that enhances the dignity, autonomy, and potential of every worker.

6. References

- Acemoglu, D., & Restrepo, P. (2020). Robots and Jobs: Evidence from US Labor Markets. *Journal of Political Economy*, 128(6), 2188–2244. <https://doi.org/10.1086/705716>
- ADP Research Institute. (2024). *People at work 2024: A global workforce view*. <https://www.adpresearch.com/wp-content/uploads/2024/04/People-at-Work-2024-A-Global-Workforce-View.pdf>
- Autor, D. H. (2015). Why Are There Still So Many Jobs? The History and Future of Workplace Automation. *Journal of Economic Perspectives*, 29(3), 3–30. <https://doi.org/10.1257/jep.29.3.3>
- Bakker, A. B., & Demerouti, E. (2017). *Job Demands-Resources Theory: Taking Stock and Looking Forward*. <https://doi.org/10.1037/ocp0000056>
- Bandura, A. (1977). Self-efficacy: Toward a unifying theory of behavioral change. *Psychological Review*, 84(2), 191–215. <https://doi.org/10.1037/0033-295X.84.2.191>
- Bengesi, S., El-Sayed, H., Sarker, M. K., Houkpati, Y., Irungu, J., & Oladunni, T. (2023). Advancements in Generative AI: A Comprehensive Review of GANs, GPT, Autoencoders, Diffusion Model, and Transformers. *IEEE Access*, 12, 69812–69837. <https://doi.org/10.1109/ACCESS.2024.3397775>
- Brougham, D., & Haar, J. (2018). Smart Technology, Artificial Intelligence, Robotics, and Algorithms (STARA): Employees' perceptions of our future workplace. *Journal of Management & Organization*, 24(2), 239–257. <https://doi.org/10.1017/JMO.2016.55>
- Brynjolfsson, E., Li, D., Raymond, L. R., Acemoglu, D., Autor, D., Axelrod, A., Dillon, E., Enam, Z., Garicano, L., Frankel, A., Manning, S., Mullainathan, S., Pierson, E., Stern, S., Rambachan, A., Reenen, J. Van, Sadun, R., Shaw, K., & Stanton, C. (2023). *Generative AI at Work*. <https://doi.org/10.3386/W31161>
- Brynjolfsson, E., Rock, D., Syverson, C., Acemoglu, D., Benzell, S., Fernald, J., Henderson, R., Goolsbee, A., Rogerson, R., Saunders, A., Summers, L., & Trajtenberg, M. (2018). *The Productivity J-Curve: How Intangibles Complement General Purpose Technologies*. <http://www.nber.org/papers/w25148>
- Buolamwini, J., & Gebru, T. (2018). Gender Shades: Intersectional Accuracy Disparities in Commercial Gender Classification. In *Proceedings of Machine Learning Research* (Vol. 81, pp. 77–91). PMLR. <https://proceedings.mlr.press/v81/buolamwini18a.html>

- Camerini, A. L., Albanese, E., & Marciano, L. (2022). The impact of screen time and green time on mental health in children and adolescents during the COVID-19 pandemic. *Computers in Human Behavior Reports*, 7, 100204. <https://doi.org/10.1016/J.CHBR.2022.100204>
- Dastin, J. (2018). Amazon scraps secret AI recruiting tool that showed bias against women. *Reuters*. <https://www.reuters.com/article/us-amazon-com-jobs-automation-insight-idUSKCN1MK08G/>
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly: Management Information Systems*, 13(3), 319–339. <https://doi.org/10.2307/249008>
- Dell’Acqua, F., McFowland III, E., Mollick, E. R., Lifshitz-Assaf, H., Kellogg, K., Rajendran, S., Kraye, L., Candelon, F., & Lakhani, K. R. (2023). Navigating the Jagged Technological Frontier: Field Experimental Evidence of the Effects of AI on Knowledge Worker Productivity and Quality. *SSRN Electronic Journal*. <https://doi.org/10.2139/SSRN.4573321>
- Deming, D. J. (2017). The Growing Importance of Social Skills in the Labor Market. *The Quarterly Journal of Economics*, 132(4), 1593–1640. <https://doi.org/10.1093/QJE/QJX022>
- Dexe, J., Franke, U., Söderlund, K., van Berkel, N., Jensen, R. H., Lepinkäinen, N., & Vaiste, J. (2022). Explaining automated decision-making: a multinational study of the GDPR right to meaningful information. *Geneva Papers on Risk and Insurance: Issues and Practice*, 47(3), 669–697. <https://doi.org/10.1057/S41288-022-00271-9>
- Dzindolet, M. T., Peterson, S. A., Pomranky, R. A., Pierce, L. G., & Beck, H. P. (2003). The role of trust in automation reliance. *International Journal of Human-Computer Studies*, 58(6), 697–718. [https://doi.org/10.1016/S1071-5819\(03\)00038-7](https://doi.org/10.1016/S1071-5819(03)00038-7)
- Eloundou, T., Manning, S., Mishkin, P., & Rock, D. (2023). *GPTs are GPTs: An Early Look at the Labor Market Impact Potential of Large Language Models*. <http://arxiv.org/abs/2303.10130>
- European Commission. (2021). *Proposal for a Regulation of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) and amending certain Union legislative acts*. https://eur-lex.europa.eu/resource.html?uri=cellar:e0649735-a372-11eb-9585-01aa75ed71a1.0001.02/DOC_1&format=PDF
- Felten, E. W., Raj, M., & Seamans, R. (2023). How will Language Modelers like ChatGPT Affect Occupations and Industries? *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.4375268>

- Frey, C. B., & Osborne, M. A. (2017). The future of employment: How susceptible are jobs to computerisation? *Technological Forecasting and Social Change*, 114, 254–280. <https://doi.org/10.1016/J.TECHFORE.2016.08.019>
- Gmyrek, P., Berg, J., & Bescond, D. (2023). Generative AI and jobs: a global analysis of potential effects on job quantity and quality. *Generative AI and Jobs*. <https://doi.org/10.54394/FHEM8239>
- Goodhue, D. L., & Thompson, R. L. (1995). Task-technology fit and individual performance. *MIS Quarterly: Management Information Systems*, 19(2), 213–233. <https://doi.org/10.2307/249689>
- Hötte, K., Somers, M., & Theodorakopoulos, A. (2023). Technology and jobs: A systematic literature review. *Technological Forecasting and Social Change*, 194. <https://doi.org/10.1016/j.techfore.2023.122750>
- Hu, K. (2023). *ChatGPT sets record for fastest-growing user base – Analyst note*. Reuters. <https://www.reuters.com/technology/chatgpt-sets-record-fastest-growing-user-base-analyst-note-2023-02-01/>
- Kellogg, K. C., Valentine, M. A., & Christin, A. (2020). Algorithms at Work: The New Contested Terrain of Control. *Https://Doi.Org/10.5465/Annals.2018.0174*, 14(1), 366–410. <https://doi.org/10.5465/ANNALS.2018.0174>
- Kim, B. J., & Lee, J. (2024). The mental health implications of artificial intelligence adoption: The crucial role of self-efficacy. *Humanities and Social Sciences Communications*, 11(1), 1–15. <https://doi.org/10.1057/s41599-024-04018-w>
- Korinek, A. (2023). Generative AI for Economic Research: Use Cases and Implications for Economists. *Journal of Economic Literature*, 61(4), 1281–1317. <https://doi.org/10.1257/JEL.20231736>
- Lane, M., Williams, M., & Broecke, S. (2023). *The impact of AI on the workplace: Main findings from the OECD AI surveys of employers and workers*. <https://doi.org/10.1787/ea0a0fe1-en>
- Lee, M. K., Kusbit, D., Metsky, E., & Dabbish, L. (2015). Working with machines: The impact of algorithmic and data-driven management on human workers. *Conference on Human Factors in Computing Systems - Proceedings*, 2015-April, 1603–1612. <https://doi.org/10.1145/2702123.2702548>
- Lepri, B., Oliver, N., Letouze, E. F., Pentland, A. P., & Vinck, P. (2019). Fair, Transparent, and Accountable Algorithmic Decision-making Processes. *Springer Netherlands*. <https://doi.org/10.1007/s13347-017-0279-x>

- Lin, L., & Parker, K. (2025). *U.S. workers are more worried than hopeful about future AI use in the workplace* (Vol. 25). <https://www.pewresearch.org/social-trends/2025/02/25/u-s-workers-are-more-worried-than-hopeful-about-future-ai-use-in-the-workplace/>
- Marguerit, D., Bia, M., Salvo, S. Di, Kramarz, F., Gathmann, C., Grasso, G., Gregory, T., Laudi, M., Machado, J., Petit, F., & Tatsiramos, K. (2025). *Augmenting or Automating Labor? The Effect of AI Development on New Work, Employment, and Wages* *.
- McKinsey & Company. (2023). *The state of AI in 2023: Generative AI's breakout year*. <https://www.mckinsey.com/capabilities/quantumblack/our-insights/the-state-of-ai-in-2023-generative-AIs-breakout-year>
- Miller, T. (2019). Explanation in artificial intelligence: Insights from the social sciences. *Artificial Intelligence*, 267, 1–38. <https://doi.org/10.1016/J.ARTINT.2018.07.007>
- Morgan, D. L. (2007). Paradigms Lost and Pragmatism Regained: Methodological Implications of Combining Qualitative and Quantitative Methods. *Journal of Mixed Methods Research*, 1(1), 48–76. <https://doi.org/10.1177/2345678906292462>;WEBSITE:WEBSITE:SAGE;ISSUE:ISSUE:DOI
- Noy, S., & Zhang, W. (2023). Experimental evidence on the productivity effects of generative artificial intelligence. *Science*, 381(6654), 187–192. <https://doi.org/10.1126/SCIENCE.ADH2586>
- Ragu-Nathan, T. S., Tarafdar, M., Ragu-Nathan, B. S., & Tu, Q. (2008). The consequences of technostress for end users in organizations: Conceptual development and validation. *Information Systems Research*, 19(4), 417–433. <https://doi.org/10.1287/isre.1070.0165>
- Raji, I. D., Smart, A., White, R. N., Mitchell, M., Gebru, T., Hutchinson, B., Smith-Loud, J., Theron, D., & Barnes, P. (2020). Closing the AI Accountability Gap: Defining an End-to-End Framework for Internal Algorithmic Auditing. *FAT* 2020 - Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency*, 33–44. <https://doi.org/10.1145/3351095.3372873>
- Remus, D., & Levy, F. S. (2016). Can Robots Be Lawyers? Computers, Lawyers, and the Practice of Law. *SSRN Electronic Journal*. <https://doi.org/10.2139/SSRN.2701092>
- Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). “Why Should I Trust You?” Explaining the Predictions of Any Classifier. <https://doi.org/10.1145/2939672.2939778>
- Selbst, A. D., Anthony, D., Bambauer, J., Bloch-Wehba, H., Boyd, W., Carlson, A., Cohen, J., Gilman, M., Brennan-Marquez, K., Geczy, I., Glater, J., Guggenberger, N., Hirsch, D., Kaminski, M., Kim, P., Solow-Niederman, A., Loo, R. Van, Verstein, A., Volokh, E., ... Watkins, E. (2021). AN

INSTITUTIONAL VIEW OF ALGORITHMIC IMPACT ASSESSMENTS. *Harvard Journal of Law & Technology*, 35(1).

Shneiderman, B. (2020). Human-Centered Artificial Intelligence: Reliable, Safe & Trustworthy. *International Journal of Human-Computer Interaction*, 36(6), 495–504. <https://doi.org/10.1080/10447318.2020.1741118>

Slack (corporate author). (2024). *The Fall 2024 Workforce Index shows executives and employees investing in AI, but uncertainty holding back adoption*. Slack. <https://slack.com/blog/news/the-fall-2024-workforce-index-shows-executives-and-employees-investing-in-ai-but-uncertainty-holding-back-adoption>

U.S. Equal Employment Opportunity Commission. (n.d.). Retrieved October 27, 2025, from <https://www.eeoc.gov/>

Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS Quarterly: Management Information Systems*, 27(3), 425–478. <https://doi.org/10.2307/30036540>

Vu, A., & Oppenlaender, J. (2025). *Prompt Engineer: Analyzing Skill Requirements in the AI Job Market*. <https://arxiv.org/pdf/2506.00058>

World Economic Forum. (2023). *The Future of Jobs Report 2023*. https://www3.weforum.org/docs/WEF_Future_of_Jobs_2023.pdf

Ziegler, A., Kalliamvakou, E., Li, X. A., Rice, A., Rifkin, D., Simister, S., Sittampalam, G., & Aftandilian, E. (2022). *Productivity assessment of neural code completion*. 21–29. <https://doi.org/10.1145/3520312.3534864>

7. Appendix A: Survey Instrument – Questionnaire

Survey on The Impact of Artificial Intelligence on Employment: Individuals' Perceptions.

This appendix contains the full questionnaire used in this research, aiming to capture individual perceptions regarding artificial intelligence (AI) and its impact on productivity, skills, job security, and ethics in the workplace.

Participant Information and Consent

Purpose:

Investigate employee perceptions about AI's impact on productivity, job security, and overall work experience as part of an ISTE C Master's thesis.

Procedures:

Participants completed one anonymous online questionnaire (approximately 10 minutes) about their background, exposure to AI, and perceptions of AI at work.

Voluntary and Confidential:

Participation was voluntary and confidential, with no personal identifiers collected. Data is securely stored and analyzed in aggregate.

Use of Data:

Results are used in this thesis and may be published academically. All data will be securely stored for five years (until December 31, 2029) and then permanently deleted.

Section A: Demographic and Background Information

Q1. Age Group: (18–24; 25–34; 35–44; 45–54; 55 or above)

Q2. Gender: (Man; Woman; Non-binary or gender diverse; Prefer to self-describe; Prefer not to say)

Q3. Highest Education Level: (High school or equivalent; Some college/university no degree; Bachelor's degree; Master's degree; Doctorate/PhD; Other)

Q4. Industry Sector: (Information Technology/Software; Financial Services; Healthcare & Life Sciences; Manufacturing & Industrial; Consumer & Retail; Professional Services; Public Sector/Education/Non-profit; Energy & Utilities; Transportation & Logistics; Media & Telecommunications; Other)

Q5. Job Function: (Executive/Senior Leadership; Engineering/Software Development; Data/AI/Analytics; Product/Project Management; Sales/Business Development; Marketing/Growth;

Customer Success/Support; Operations/Supply Chain; Finance/Accounting; HR/Talent Management; Legal/Compliance/Risk; R&D; Other)

Q6. Years of Professional Experience: (<1 year; 1–5 years; 6–10 years; 11–20 years; >20 years)

Q7. Familiarity with AI Concepts: (Scale: 1 = Not at all familiar, 5 = Extremely familiar)

Section B: Exposure to AI in the Workplace

Q8. Current Use of AI Tools: (Yes; No; Not sure / Not applicable)

Q9. Types of AI Tools Used: (Generative AI; AI-based data analysis/forecasting; Intelligent assistants/chatbots; Automation/RPA tools; Machine-learning development tools; Computer vision systems; NLP tools; Decision-support AI systems; Other)

Q10. Frequency of AI Use: (Daily; A few times/week; A few times/month; Rarely)

Q11. Training on AI Tools: (Yes-extensive; Yes-some; No formal training, learned on own; No training required)

Section C: Perceptions of AI's Impact on Work

Participants indicated their level of agreement on each statement using the scale: **Strongly Disagree, Disagree, Neither Agree nor Disagree, Agree, Strongly Agree, I don't know/Not sure.**

AI and Work Tasks:

Q12. Using AI tools has made me more productive.

Q13. AI has helped reduce errors or mistakes in my work.

Q14. AI tools have taken over the most repetitive/mundane parts of my job.

Q15. AI tools have improved the quality of my work outputs.

Q16. Using AI saves me time on tasks.

Working with AI Systems:

Q17. AI systems fit well with tasks I need to accomplish (relevant and useful).

Q18. I trust the results/recommendations from AI systems.

Q19. I find it confusing/difficult to understand how AI makes decisions.

Q20. Decisions made by AI systems are transparent (I usually know how or why recommendations are made).

Q21. I am concerned AI might be used unfairly or unethically at work.

AI's Impact on My Job:

Q22. I feel anxious that AI could replace my job/responsibilities.

Q23. AI has made my job more stressful.

Q24. AI is creating opportunities for more interesting/higher-value tasks.

Q25. Overall, I believe AI's impact on my job has been positive.

Q26. I am enthusiastic about increasing AI use in my field.

Adapting to AI:

Q27. I feel confident in my ability to work with AI technologies.

Q28. My organization supports employees learning and using AI.

Q29. I have sufficient training/resources available to learn AI tools.

Q30. I actively seek new skills to remain relevant as AI becomes common.

Q31. I am interested in using even more AI tools/features in the future.

Section D: Open-Ended Questions

Participants provided written responses to the following:

Q32. Benefits of AI: What do you see as the biggest benefit or opportunity from AI adoption in your job or industry?

Q33. Concerns about AI: What is your biggest concern or fear regarding AI in the workplace?

Q34. Suggestions for Implementation: Any recommendations for your organization (or generally) on handling AI introduction?

Q35. Additional Comments: Any further thoughts or experiences related to AI's impact on your work or career?

End of Questionnaire

Thank you for your participation. Your insights significantly contribute to understanding employee perceptions of AI's role and influence in contemporary workplaces.

For questions or a summary of results, please contact: daniellourenco.oliveira@my.istec.pt