

# MGI

Master Degree Program in  
**Information Management**

**The Effect of Generative AI on the Productivity of Consultants**

Marijke Brommann

Master Thesis

presented as a partial requirement for obtaining the Master Degree in Information Management

**NOVA Information Management School**  
**Instituto Superior de Estatística e Gestão de Informação**

Universidade Nova de Lisboa

**NOVA Information Management School**  
**Instituto Superior de Estatística e Gestão de Informação**  
Universidade Nova de Lisboa

**The Effect of Generative AI on the Productivity of Consultants**

by  
Marijke Brommann

Master Thesis presented as a partial requirement for obtaining the Master Degree in Information Management, with a specialization in Information Systems and Technologies Management

**Supervised by**

Maria Manuela Simões Aparício da Costa, PhD

February, 2024

## **STATEMENT OF INTEGRITY**

I hereby declare having conducted this academic work with integrity. I confirm that I have not used plagiarism or any form of undue use of information or falsification of results along the process leading to its elaboration. I further declare that I have fully acknowledged the Rules of Conduct and Code of Honor from the NOVA Information Management School.

*[Cologne, 22<sup>nd</sup> of February 2024]*

## ABSTRACT

Powered by technological advancements and evolving client expectations, the consulting industry is currently undergoing a transformative digital transformation. This study focuses on generative artificial intelligence (GEN-AI), a technology that, despite its decades-long existence, has recently seen significant performance improvements, expanding its capabilities and business use cases. While other studies already investigate the separate productivity effects in the consulting industry and of GEN-AI, a research gap persists in the integration of both aspects. Hence, this research aims to identify what impact the integration of GEN-AI has on the individual productivity in the consulting industry. To address this research question, the study identifies GEN-AI determinants influencing the productivity of consultants, proposes a theoretical model, and validates it empirically. The quantitative research approach utilizes a survey, distributed to several consulting enterprises and the partial least squares-structural equation modeling technique to analyze the data. The results not only confirm that GEN-AI significantly improves the productivity of consultants but also provide insights into the adaptability of the Delone and McLean (D&M) model of information system success. Other findings highlight age-specific challenges in GEN-AI adoption and a positive feedback loop between the usage of and satisfaction with the system. The study concludes by suggesting practical measures for organizations adopting the technology, emphasizing tailored training initiatives, or securing the service as well as information quality.

## KEYWORDS

Consulting Industry; Generative Artificial Intelligence; Productivity Effects; Structural Equation Modeling; Ethical Aspects

### Sustainable Development Goals (SDG):



## TABLE OF CONTENTS

|                                                                  |    |
|------------------------------------------------------------------|----|
| 1. Introduction .....                                            | 1  |
| 2. Theoretical Background on Generative AI in Consulting.....    | 3  |
| 2.1. Digital Transformation in Consulting.....                   | 3  |
| 2.1.1. Employee Training.....                                    | 3  |
| 2.1.2. Ethical Behavior.....                                     | 4  |
| 2.2. Productivity in Consulting .....                            | 6  |
| 3. Research Model Proposal .....                                 | 9  |
| 4. Empirical Study and Results .....                             | 14 |
| 4.1. Sampling Design Process.....                                | 14 |
| 4.2. Measurement Model Results.....                              | 15 |
| 4.3. Structural Model Results.....                               | 17 |
| 5. Discussion.....                                               | 21 |
| 6. Conclusions and Future Works.....                             | 24 |
| 6.1. Conclusions .....                                           | 24 |
| 6.2. Implications and Future Works .....                         | 24 |
| 6.2.1. Theoretical Implications.....                             | 24 |
| 6.2.2. Practical Implications .....                              | 25 |
| 6.2.3. Limitations and Future Works.....                         | 26 |
| Bibliographical References.....                                  | 27 |
| Appendix A. System Success Dimensions in Empirical Studies ..... | 34 |
| Appendix B. Measurement Model.....                               | 35 |
| Appendix C. Ethics Committee Approval .....                      | 36 |
| Appendix D. Extract Survey Design .....                          | 37 |
| Appendix E. Measurement Item Cross Loadings .....                | 40 |

## **LIST OF FIGURES**

|                                                             |    |
|-------------------------------------------------------------|----|
| Figure 1 - Generative AI Adoption Impact Model .....        | 13 |
| Figure 2 - Results Generative AI Adoption Impact Model..... | 19 |

## LIST OF TABLES

|                                                               |    |
|---------------------------------------------------------------|----|
| Table 1 - Recent Related Studies on the Usage of GEN-AI ..... | 8  |
| Table 2 - Construct Operationalization .....                  | 9  |
| Table 3 - Reliability and Validity Results .....              | 16 |
| Table 4 - Fornell-Larcker Results .....                       | 16 |
| Table 5 - HTMT Results .....                                  | 17 |
| Table 6 - VIF Results .....                                   | 17 |
| Table 7 - Research Hypotheses Results.....                    | 18 |

## **LIST OF ABBREVIATIONS AND ACRONYMS**

|                |                                                   |
|----------------|---------------------------------------------------|
| <b>A</b>       | Age                                               |
| <b>AI</b>      | Artificial Intelligence                           |
| <b>AVE</b>     | Average Variance Extracted                        |
| <b>CA</b>      | Cronbach's Alpha                                  |
| <b>D&amp;M</b> | Delone and Mclean                                 |
| <b>E</b>       | Ethics of AI                                      |
| <b>GEN-AI</b>  | Generative Artificial Intelligence                |
| <b>HTMT</b>    | Heterotrait-Monotrait                             |
| <b>IQ</b>      | Information Quality                               |
| <b>LLM</b>     | Large Language Model                              |
| <b>NB</b>      | Net Benefits                                      |
| <b>PLS</b>     | Partial Least Squares                             |
| <b>PLS-SEM</b> | Partial Least Squares-Structure Equation Modeling |
| <b>SE</b>      | Self-Efficacy                                     |
| <b>SEM</b>     | Structural Equation Modeling                      |
| <b>SEQ</b>     | Service Quality                                   |
| <b>SU</b>      | System Use                                        |
| <b>SYQ</b>     | System Quality                                    |
| <b>T</b>       | Training                                          |
| <b>US</b>      | User Satisfaction                                 |
| <b>VIF</b>     | Variance Inflation Factor                         |

# 1. INTRODUCTION

The consulting industry is currently undergoing a significant digital transformation, driven by technological advancements, and changing client expectations. To remain competitive, the efficient use of data and technology is essential for delivering tailored services to clients. However, this transformation also presents challenges, including the need to adapt to new tools or to address ethical concerns (Gînguță et al., 2023). This study takes these factors along with recent developments in generative artificial intelligence (GEN-AI) into consideration. More precisely, it examines how the uprising disruptive technology affects the work of consultants.

GEN-AI has been around for several decades. Yet, its performance and relative accessibility have recently improved drastically which increases the opportunities the technology offers (Koh & Doroudi, 2023). The bandwidth of tasks GEN-AI can support nowadays is extensive and ranges for instance from automating repetitive tasks to improving software quality or automating workflows. However, the technology is neither deterministic nor fully explainable in most cases which creates an immense risk, for example in the context of cybersecurity (Ebert & Louridas, 2023).

While previous studies have separately examined the productivity effects of consulting (Bologa & Lupu, 2014; Nachum, 1999) and those associated with GEN-AI (Damioli et al., 2021; Noy & Zhang, 2023), there is a lack of research that combines both areas. Furthermore, a research gap remains concerning whether the explainability of artificial intelligence (AI) positively or negatively impacts the competitiveness and productivity of organizations (Dwivedi et al., 2023) as well as the broader personal impact of recent AI advancements (Collins et al., 2021). Additionally, there is a scarcity of research exploring ethical concerns (Gînguță et al., 2023; Siau & Wang, 2020) and the influence of training and mindset on employees (Dwivedi et al., 2023), despite the established relevance of the latter in technology adoption (Hanaysha, 2016).

Based on the previously defined gaps the research question is specified:

"What impact does the integration of GEN-AI have on the individual productivity in the consulting industry?"

The following research objectives are determined to answer the research question:

1. Identify the GEN-AI determinants that influence the productivity of consultants.
2. Propose a theoretical model to explain the impact on the productivity of consultants in the context of GEN-AI.
3. Validate the model empirically.

This research adopts a quantitative approach, utilizing a conceptual model and advanced statistical techniques. The study employs partial least squares (PLS) as a data analysis method which falls under the structural equation modeling (SEM) approach (Hair et al., 2011; Urbach & Ahlemann, 2010). Using a single cross-sectional design, the research provides an overview of participants' opinions and experiences. Data collection is conducted through a survey method, involving the creation and distribution of a questionnaire to multiple consulting enterprises.

This study contributes to the comprehension of how GEN-AI shapes the individual productivity of consultants. It achieves this by employing and expanding upon the Delone and McLean (D&M) model for information systems, with a specific emphasis on the impact of employee training and ethics within the context of GEN-AI. This explores a so far non-existent perspective on the adoption and utilization of the technology. In particular, the study serves as a valuable guide to support consulting enterprises in leveraging GEN-AI in their daily business and supporting their employees in the adoption process while considering the advantages of the disruptive technology as well as the potential risks.

Laying the foundation for a comprehensive exploration of GEN-AI productivity effects in the consulting industry, this thesis begins with an introduction that sets the stage for a detailed examination. Following the introduction, Chapter 2 presents the theoretical background, providing an overview of existing research regarding the digital transformation in the consulting industry. This chapter specifically emphasizes the topics of training, ethics as well as productivity measurement in consulting, along with the contextual relevance of GEN-AI. Moving forward, Chapter 3 unfolds the research model, detailing its constructs and hypotheses. Chapter 4 explains the underlying methodology, offering further insights into the data generation and analysis process and presenting the study's measurement and structural model results. The ensuing Chapter 5 discusses the results to address the research question. Finally, Chapter 6 concludes the study by considering implications and limitations.

## **2. THEORETICAL BACKGROUND ON GENERATIVE AI IN CONSULTING**

### **2.1. DIGITAL TRANSFORMATION IN CONSULTING**

Consulting is a division of the service industry. The research area of information systems for instance defines it as advising firms to digitalize their business processes (Bode et al., 2022). The consulting industry is prone to disruption because of its competitive nature. Consequently, the topic is widely spoken about in recent studies. Consulting companies must act fast when new technologies arise and grow in importance to adapt to changing client and market needs (Crişan & Stanca, 2021; Oesterle et al., 2016). The aim is to continue to be a driver of innovation in the client's firm (Tavoletti et al., 2021). Oesterle et al. (2016) further emphasize that the expertise of consultants on the technological front significantly influences the quality of services in the consulting industry. It is consequently crucial to think ahead, keep the technical knowledge updated, and educate the consultants accordingly. Hence, navigating the dynamics of digital transformation is central to the continuous success and relevance of the consulting industry. In this context, the term digital transformation signifies leveraging digitalization to adapt business processes and enhance overall effectiveness (Bode et al., 2022).

GEN-AI tools such as ChatGPT or GitHub CoPilot are amongst the most disruptive technologies of recent years (Dwivedi et al., 2023). They can therefore change the role of consultants. In this specific case, the role is changed by introducing a collaboration between humans and machines. These uprising GEN-AI systems are machine learning technologies that can help users with a variety of tasks, for example in the creation of text or visual output, leveraging large amounts of training data (Noy & Zhang, 2023). The application generates an answer specifically to that user's question instead of providing a prefabricating one like search engines (Ebert & Louridas, 2023). More precisely, Ebert & Louridas (2023) explain that GEN-AI is based on Large Language Models (LLM) which are large neural network models constantly trained on an extensive amount of data. It is already proven that consultants regard GEN-AI tools as valuable but not as a universal solution. Therefore, extensive analysis is required to determine the suitability of GEN-AI for certain needs and tasks in the industry (Bayati, 2019) as well as to quantify the general and personal impact of these recent advancements (Collins et al., 2021).

#### **2.1.1. Employee Training**

The challenge of introducing new technologies in a company is that employees need to be willing as well as capable of using them to keep up with the rapid technological changes (Yaghmaie & Jayasuriya, 2004). This is specifically true regarding GEN-AI due to its highly disruptive nature (Nah et al., 2023) which is the reason why the topic requires further research according to Dwivedi et al. (2023).

Firstly, a positive pre-adoptive attitude of employees is fundamental for the successful implementation of GEN-AI (Chiu et al., 2021). Additionally, the attitude moderates the

perceived information quality and satisfaction with the system, making it a factor of success (Prentice et al., 2020). Yet, if the pre-adoptive attitude is positive and AI is appreciated as a supporting tool, the capability to properly use the technology can still create a boundary (Fügener et al., 2022). Employees naturally try to investigate how to adapt their professional role identity to a new technology. In the worst case, they fail to integrate the technology and consequently feel threatened by it which also creates a risk of fluctuation (Strich et al., 2021). Despite the assumption that users appreciate a high task efficiency in GEN-AI tools, Baek & Kim (2023) found that an efficiency exceeding expectations increases the perceived skepticism and creepiness towards the tool. This, in turn, can harm the likelihood of usage. Since the complexity of a specific technology negatively impacts its adoption, supporting employees is crucial (wael AL-khatib, 2023). The company therefore needs to step in, offer training courses to educate employees about the technology, and perhaps point out alternative job tasks that cannot be replaced by that specific technology (Chiu et al., 2021; Strich et al., 2021). An AI training course improves people's understanding and attitudes toward AI (Budhwar et al., 2023), simultaneously influences the awareness of ethical factors (Shih et al., 2021), and positively impacts the commitment to the organization (Hanaysha, 2016b). Training can also decrease anxiety when dealing with new technologies (Yaghmaie & Jayasuriya, 2004). Furthermore, in the specific case of AI, the user's proficiency with the tool positively correlates with the generation of higher-quality information (Dwivedi et al., 2023). This expertise also enhances the user's ability to differentiate between tasks suitable for human intervention and those better suited for AI, thereby contributing to the overall success of the tool's utilization among employees (Fügener et al., 2022). Early adopters among the employees can be consulted to set up these training courses since they are aware of the technology's benefits and generally have a more positive attitude towards it (Fosso Wamba et al., 2023). Based on the principle of human-AI collaboration the focus in the upskilling training should lie on five main skills: data analysis, digital, complex cognitive, decision-making, and continuous learning skills (Jaiswal et al., 2022). Finally, and specifically relevant to this study is Hanaysha's (2016) finding, affirming that employee training significantly enhances productivity. However, the research was exclusively limited to the educational context within the region of Malaysia. The author therefore points out the generalizability as a limitation for future research. In conclusion, the goal of training in the context of GEN-AI is to ensure that disruptive technologies are rather a competitive advantage than a source of anxiety (Gînguță et al., 2023).

### **2.1.2. Ethical Behavior**

A significant factor in fostering trust between consultant and client is profound ethical behavior (Mingaleva, 2013). When it comes to the usage of GEN-AI, there is still a lack of overarching regulations that are essential to prevent the technology from creating serious harm (Sætra, 2023). Currently, the global understanding includes five ethical AI principles which are transparency, justice/fairness, non-maleficence, responsibility, and privacy. The interpretation, way of implementation, or level of importance still differs immensely though

(Jobin et al., 2019). AI ethics do not have a strong influence on decision-making in technological professions yet, especially if economic interests are contradictory. Simultaneously, the topic of using and developing AI technologies also does not raise enough attention on an enterprise level (Hagendorff, 2020; Siau & Wang, 2020) even though some firms are starting to set up policies involving the technology (Zhou et al., 2020). The lack of these initiatives creates a serious problem since missing out on addressing these ethical or mindset challenges might even lead to a failing integration (Gînguță et al., 2023). Underlining the significance of responsible AI development, Siau & Wang (2020) and Wach et al. (2023) emphasize the immediate need for actions to secure the future of AI technologies. Responsible AI includes the ethical, social, and environmental effects of the technology (Gyllensvärd & Hiljegren, 2023). Mechanisms such as audit trails, interpretability or algorithmic design choices can amongst others play a role in ensuring the latter (Fenwick & Molnar, 2022). Hence, ethical principles should be incorporated into every part of the AI lifecycle (Zhou et al., 2020). Baabdullah (2024) leveraged the information quality model to define several recommendations for decision-makers to ensure a safe and effective use of GEN-AI. He specifically points out the importance of reviewing the content to assure believability and accuracy as well as defining certain standards for feasibility and usability such as timeliness. The output should also be consistent and easy to interpret. Closely related to Baabdullah's findings is the topic of AI explainability and transparency. Users should have a right to understand how exactly the specific AI model works to better categorize the quality or legal implications of a certain output (Clinciu & Hastie, 2019; Dwivedi et al., 2023). Transparency also significantly increases trust in the technology (Yu & Li, 2022) while concerns about data privacy rather lead to a feeling of discomfort (Baek & Kim, 2023). One solution to these challenges might be to make GEN-AI more human-centered. In that case, the algorithm does not act like a human but the previously discussed attributes such as transparency, ethics but also empathy would be incorporated into the development of the technology as human traits (Nah et al., 2023). According to Sætra (2023), this also includes promoting societal values such as freedom or democracy. Yu & Li (2022) also detected that decision-making transparency makes AI more human-like and therefore increases the overall trust in the technology. This leads to increased effectiveness and decreased discomfort in the usage of the technology. However, Baek & Kim (2023) found a contradictory outcome in their study where a more human-like appearance of the GEN-AI tool increases the perceived creepiness and discomfort and harms the likeliness of usage. Evidently, the results provided by GEN-AI are based on probabilities and not necessarily authorized sources which makes a human review process indispensable at the current point in time and will most likely continuously grow in relevance (Ebert & Louridas, 2023). A research gap remains regarding whether the explainability of AI positively or negatively impacts the productivity of organizations (Dwivedi et al., 2023). Furthermore, factors to control ethical risks need to be determined (Gînguță et al., 2023) and the topic needs to be a central aspect of research at the current point in time (Siau & Wang, 2020).

## 2.2. PRODUCTIVITY IN CONSULTING

With increasing spending in the technology sector, managers are highly interested in measuring the impact of their investments. Torkzadeh & Doll (1999) define four dimensions to measure the impact of technology, offering a framework to assess the influence of information systems. These dimensions include individual impact, organizational impact, task impact, and impact on people-task-technology fit. For this study, task as well as individual productivity are specifically interesting. While productivity is generally described as the relationship between the output and input of resources or factors (Björkman, 1992), task productivity assesses the influence of technology on specific tasks and activities within the organization and falls under the task impact dimension. Productivity is also measured with the dimension of individual impact which refers to the overall effectiveness or job performance of the individual (Torkzadeh & Doll, 1999). Another model that offers valuable insights into how technology investments impact individual productivity as well as overall organizational performance is the D&M model which was first developed in 1992 and updated in 2003. The model provides a framework of six interdependent variables including system quality, information quality, service quality, system use, user satisfaction, and net benefits to evaluate the success of information systems.

From an employee's perspective productivity is strictly tied to motivation and the diverse and individual factors that influence the latter such as interesting work or opportunities for advancement (Sabir, 2017). The amount of existing research on employee productivity is still limited. Yet, it is highly relevant since it is tightly connected to the overall performance of the company (Hanaysha, 2016b). Nachum (1999) found that in professional service firms like the consulting industry, productivity is generally – compared to manufacturing firms – difficult to define since the existing measurements consider inputs and outputs as tangible standardized quantities. These mostly do not exist in the service industry. Hence, while managers in that field are aware of the need to measure their productivity, they are missing suitable instruments. The scale introduced in Nachum's study suggests considering more industry-specific measurement items such as the stock of knowledge including the required training as well as considering the perspective of the consulting firm and the client equally. Despite the measurement scale, other studies have found additional productivity-enhancing dimensions. For instance, knowledge-sharing frameworks, employee empowerment, and teamwork have a significant impact on the productivity of consultants (Bologa & Lupu, 2014; Hanaysha, 2016). Competitive dynamics – previously mentioned as a characteristic of the consulting industry - can additionally uplift productivity levels since low competitiveness permits inefficient behavior (Baily & Solow, 2001).

There already is scientific evidence that GEN-AI can generally enhance productivity (Cockburn et al., 2019). Sometimes productivity growth and technical progress are considered to be alike by researchers, but some also define it as an interdependent relationship. According to them, productivity is the *“net change in output due to change in efficiency and technical change”*

(Grosskopf, 1993, p.160). Nevertheless, like the challenges of measuring productivity in the consulting industry, the productivity paradox describes the difficulties of connecting technological advancements to an increase in performance. While factors like mismeasurements or premature measurement explain the paradox to some extent, a gap remains (Brynjolfsson, 1993). This also holds for GEN-AI since productivity growth generally stagnated in the market about five years ago when the technology started to gain more importance but was not mature enough yet (Gries & Naudé, 2018). Meanwhile, analysis indicates that small and medium enterprises, along with companies in the service industry with AI patent applications, have experienced a boost in labor productivity. This underscores the significance of promptly adapting organizations to new technologies (Damioli et al., 2021). The same effect can be observed for startups. GEN-AI improves efficiency and can even lead to a decrease in the likelihood of failure (Gyllensvärd & Hiljegren, 2023). Larger companies are still missing out on this productivity growth since their processes usually take longer to adapt (Damioli et al., 2021). To ensure productivity effects on a task level the organization should encourage human-GEN-AI interactions by paying special attention to the right mapping of employee capabilities to the specific GEN-AI tool to prevent error multiplication (Walkowiak, 2023). Moreover, looking at the importance of generating customer and market insights as well as creating tailored client offerings as tasks in the consulting industry, GEN-AI is significantly more effective than previous digital technologies. These effects should be leveraged and there is evidence they can also be used to improve the sales lead generation process and extend the customer base more efficiently (Kshetri et al., 2023). However, to leverage the technology as efficiently as possible and achieve the best possible increase in productivity, a company-wide GEN-AI strategy is essential (Dell'Acqua et al., 2023). Furthermore, productivity growth can also be observed on an individual level. Noy & Zhang (2023) detected in their experiment that a productivity increase is given in terms of quality when using GEN-AI. Furthermore, task productivity can be increased significantly by reducing the amount of time spent. Workers with lower abilities hereby benefit from using GEN-AI more extensively. By now GEN-AI is also capable of supporting complex and knowledge-intensive tasks and can create significant productivity growth. Since the technology is constantly developing and advancing it is crucial to closely monitor the boundaries of the technology. For tasks that are too complex for the current state of GEN-AI, the effect can reverse, and quality and therefore efficiency and productivity decrease, especially if the results are not critically reviewed by an employee (Dell'Acqua et al., 2023). Even though a growing amount of literature exists regarding the productivity effects in the consulting industry as well as the productivity impact of GEN-AI, research integrating both areas is still insufficient.

Table 1 - Recent Related Studies on the Usage of GEN-AI

| Author                  | Journal                                                | Study Objective                                                                                                                                         | Theory/ Model                                                   |
|-------------------------|--------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------|
| (Prentice et al., 2020) | International Journal of Hospitality Management        | Comprehend the connections among customer perceptions of AI service quality, customer satisfaction, and customer engagement.                            | Affordance and rational choice theory                           |
| (Damioli et al., 2021)  | Eurasian Business Review                               | Examine how AI patenting activities affect the labor productivity of companies.                                                                         | General method of moments model                                 |
| (Jaiswal et al., 2022)  | The International Journal of Human Resource Management | Explore essential skills needed for employee upskilling in the context of AI adoption.                                                                  | Dynamic skill, neo-human capital, and AI job replacement theory |
| (Gînguță et al., 2023)  | Electronics                                            | Examine the ethical implications and challenges of AI in the business consulting industry.                                                              | Structural equation model                                       |
| (Baek & Kim, 2023)      | Telematics and Informatics                             | Investigate the influence of user motivations on the perceived creepiness, trust, and intention to use AI chatbot technology in ChatGPT interactions.   | Use and gratifications theory                                   |
| (Noy & Zhang, 2023)     | Science                                                | Explore the impact of ChatGPT on the productivity of mid-level professional writing tasks.                                                              | Experiment                                                      |
| (Baabdullah, 2024)      | Technological Forecasting & Social Change              | Examine and analyze the factors influencing the content quality of generative conversational AI agents and their effects on decision-making efficiency. | Information quality model                                       |

To conclude this chapter and lay the theoretical foundation for this study, a curated compilation of recent and relevant papers on GEN-AI usage research has been selected and presented in *Table 1*. Each paper is drawn from a different publication, providing a comprehensive view of the subject. It becomes apparent that recent studies encompass a wide range of objectives and methodologies. Some delve into the topic of quality and adoption (Baabdullah, 2024; Hyun Baek & Kim, 2023; Jaiswal et al., 2022; Prentice et al., 2020), while others explore the impact on labor productivity (Damioli et al., 2021; Noy & Zhang, 2023) or identify ethical challenges (Gînguță et al., 2023). However, these studies employ diverse theories and models, and a notable gap exists in the absence of comprehensive structural models measuring benefits, highlighting the need for a deeper understanding of success factors in the field of GEN-AI.

### 3. RESEARCH MODEL PROPOSAL

This study adopts the revised D&M model of information system success from 2003 as its foundational framework within the distinctive context of GEN-AI. The primary objective of the model is to offer insights into the quality dimensions of satisfaction while concurrently assessing the success of information systems (DeLone & McLean, 1992). The selection of the D&M model is based on its well-established reputation, supported by numerous empirical studies, showcasing its robustness in evaluating information system success across various scenarios (*Appendix A*). Its proven effectiveness and comprehensive coverage of important dimensions make it a suitable choice for this study. To enhance the applicability of the model to the unique characteristics of GEN-AI, additional constructs were incorporated to develop the final conceptual model for this study, guided by insights from the literature review. This strategic adaptation aligns with the recommendation by Petter et al. (2013), emphasizing the importance of extending established models to improve contextual relevance. This adaption of the D&M model aligns with observed practices in other studies, documented in *Appendix A* (Aparicio et al., 2019; Elsdaig, 2019; Ojo, 2017; Yu & Qian, 2018).

The main constructs of the model, namely information quality, system quality, and service quality, affect user satisfaction and system use. User satisfaction and system mutually influence each other as well as the construct of net benefits. Delone & McLean (2003) specifically point out that net benefits need to be clearly defined for each study individually. To understand the effects that GEN-AI has on the work of consultants, net benefits are specified on an individual level as the productivity of a consultant in this study. Additionally, the model was extended with the constructs training, adapted from Hanaysha (2016a), self-efficacy, adapted from Venkatesh & Bala (2008), and finally, ethics of AI was added as an influential construct based on Siau & Wang (2020). At the time this study was conducted, to the best of knowledge no study had been carried out applying the D&M model to the context of GEN-AI. *Table 2* shows the specific constructs including the acronym, definition as well as specific reference.

Table 2 - Construct Operationalization

| Construct     | Acronym | Definition                                                                                                                                                                                                           | Reference                      |
|---------------|---------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------|
| Training      | T       | <i>"Systematic process that aims to help employees enhance their knowledge and skills and develop positive behavior through learning experience that is expected to help employees achieve greater performance."</i> | (Hanaysha 2016a, p. 299)       |
| Self-Efficacy | SE      | <i>"Computer self-efficacy refers to individuals' control beliefs regarding his or her personal ability to use a system."</i>                                                                                        | (Venkatesh & Bala, 2008, p. 9) |

| Construct           | Acronym | Definition                                                                                                                                       | Reference                     |
|---------------------|---------|--------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------|
| Ethics of AI        | E       | <i>"The ethics of AI is part of the ethics of advanced technology that focuses on robots and other artificially intelligent agents."</i>         | (Siau & Wang, 2020, p. 76)    |
| Service Quality     | SEQ     | <i>"Quality of the support that system users receive from the IS department and support personnel."</i>                                          | (Petter et al., 2008, p. 239) |
| Information Quality | IQ      | <i>"Desirable characteristics of the system output."</i>                                                                                         | (Petter et al., 2008, p. 239) |
| System Quality      | SYQ     | <i>"The desirable characteristics of an information system."</i>                                                                                 | (Petter et al., 2008, p. 238) |
| User Satisfaction   | US      | <i>"Users' level of overall satisfaction with their interaction with a system."</i>                                                              | (Petter et al., 2008, p. 239) |
| System Use          | SU      | <i>"Degree and manner in which staff and customers utilize the capabilities of a system."</i>                                                    | (Petter et al., 2008, p. 239) |
| Net Benefits        | NB      | <i>"The extent to which information systems are contributing to the success of individuals, groups, organizations, industries, and nations."</i> | (Petter et al., 2008, p. 239) |

An individual's proficiency in computer skills generally exerts a positively significant impact on the intention to use a system (Yu et al., 2009). Moreover, several studies mention the importance of training in the context of GEN-AI to ensure that employees are skilled and comfortable using it (Alexandre & Blanckaert, 2020; Budhwar et al., 2023; Shih et al., 2021; Strich et al., 2021). The more familiar a potential user is with the technology, the better the task quality that can be generated by using GEN-AI and therefore the more efficient the usage of the specific tool (Dwivedi et al., 2023). Noy & Zhang (2023) indicate that exposure to a GEN-AI tool increases self-efficacy towards automation technologies which indicates a relationship between the two and leads to the first hypothesis:

**H1:** Training positively and significantly influences self-efficacy.

The increasing relevance of ethics in AI education has been emphasized in recent studies, as the technology poses potential harm in case of misuse. Dedicated training courses can help individuals reflect on the connection between ethics and AI, enabling them to act responsibly when using AI tools (Borenstein & Howard, 2021; Garrett et al., 2020). Moreover, the field of AI ethics shares a close relationship with business ethics to the extent that the content mediation and development of both fields can mutually benefit (Schultz & Seele, 2023). Based on the results of Shih et al. (2021) that an AI training course not only improves the general understanding but also influences the awareness of ethical factors, the second hypothesis was added:

**H2:** Training positively and significantly influences Ethics of AI.

Computer self-efficacy was derived from self-efficacy which refers to an individual's judgement of how to perform a certain task. Computer self-efficacy specifically refers to the ability to use a computer. Usually, the longer the time of usage and therefore experience with

a system, the higher the level of computer self-efficacy (Hung & Liang, 2001). Self-efficacy is directly related to the usage of a newly introduced system (Howard, 2014; Noy & Zhang, 2023; Yu & Qian, 2018), specifically to the perceived ease of use (Venkatesh & Bala, 2008). In this context, self-efficacy has a more pronounced influence on system usage compared to behavioral intention (Yi & Hwang, 2003). Therefore, the hypothesis determines:

**H3:** Self-efficacy positively and significantly influences use.

AI ethics generally refers to ethical issues regarding AI which require defining ethical principles, rules, guidelines, policies, and regulations for the technology to ensure ethical behavior (Siau & Wang, 2020). According to the literature, ethical principles should be incorporated into every part of the AI lifecycle (Zhou et al., 2020) with a strong focus on the explainability and transparency of a model to let the user better categorize a certain output (Clinciu & Hastie, 2019; Dwivedi et al., 2023). Baabdullah (2024) also points out the relevance of setting up an ethical reference framework in the content of GEN-AI to secure the generated information.

AI ethics not only directly influence the use of an AI tool (Harrer, 2023) but increased transparency even leads to higher effectiveness and likeliness in using GEN-AI tools (Riedl, 2019; Yu & Li, 2022) which leads to framing the following hypothesis:

**H4:** Ethics of AI significantly influence use.

Service quality measures the help that a tool's support department provides for the user regarding attributes like responsiveness or assurance (Urbach & Müller, 2012). So far, service quality has not been a common research topic concerning GEN-AI. Usually, literature talks about the human-AI-relationship and how one supports the other (Fügener et al., 2022). However, Prentice et al. (2020) identified a connection between perceptions of AI service quality and satisfaction and engagement. In the context of GEN-AI, the dynamic is particularly interesting as users actively contribute to system training through their interactions, creating a close association with the service quality construct (Haleem et al., 2022). The defined hypotheses are:

**H5a:** Service Quality positively and significantly influences use.

**H5b:** Service Quality positively and significantly influences user satisfaction.

Information quality incorporates attributes such as accuracy, completeness, relevance, and understandability of the system (Delone & Mclean, 2003). The dimension has a proven strong impact on user satisfaction while other studies have found contradictory results regarding the impact on the dimension system use (Petter et al., 2008). Nevertheless, it is recognized that the quality of the training data and hence the LLM significantly affects the quality of AI-generated answers for a specific tool, subsequently influencing the likelihood of usage

(Dwivedi et al., 2023). Conversely, users might experience a feeling of discomfort if the provided information exceeds their expectations (Baek & Kim, 2023). Information quality control is consequently crucial (Wach et al., 2023). The construct is specifically important when looking at the consulting industry since information is considered one of the main resources (Oesterle et al., 2016). Therefore, the following two hypotheses are framed:

**H6a:** Information Quality positively and significantly influences use.

**H6b:** Information Quality positively and significantly influences user satisfaction.

System quality measures the essential attributes of a system, encompassing features or ease of use (Urbach & Müller, 2012). While numerous studies have explored the influence of system quality on user satisfaction (Petter et al., 2008) and defined it as the most significant dimension in explaining the independent variable (Urbach et al., 2010), the findings regarding its impact on system use are inconclusive. Some studies establish a significant impact, while others do not (Petter et al., 2008). Yet, Prentice et al. (2020) underscore the importance of a user-friendly design for AI tools, as quality influences both, customer satisfaction and engagement. Moreover, ease of use is explicitly highlighted as a success factor for tools like ChatGPT, exemplifying GEN-AI (Dwivedi et al., 2023). Consequently, the hypotheses are as follows:

**H7a:** System Quality positively and significantly influences use.

**H7b:** System Quality positively and significantly influences user satisfaction.

System use, often described as the system's usage pattern, encompassing aspects like frequency or duration of use relates to user satisfaction, characterized by the perceived level of favorability toward the whole system. Previous research in the field of technology adoption determines that user satisfaction plays a significant role in driving technology use, with satisfied users displaying a higher likelihood of sustained use (Urbach & Müller, 2012). This establishes one plausible interrelation between the two dimensions. Other studies suggest an interdependent relationship between the two variables (Delone & Mclean, 2003; Urbach et al., 2010). In the specific context of GEN-AI, heightened expertise correlates with improved content generation, potentially elevating the overall user experience with the tool (Dwivedi et al., 2023). This scenario encourages a positive interaction between system use and user satisfaction. The hypotheses marking the bidirectional relationship are the following:

**H8:** User Satisfaction positively and significantly influences use.

**H9:** Use positively and significantly influences user satisfaction.

Additionally, user satisfaction and system use – according to Delone & McLean (2003) – both have a positive and significant impact on the independent variable net benefits. Urbach et al. (2010) confirm these hypotheses. Nevertheless, other authors suggest otherwise and could

not detect a significant impact on net benefits (Petter et al., 2008). The usage is specifically important when it comes to GEN-AI since the interaction shapes the overall system and degree of possible personalization (Guerreiro & Loureiro, 2023). This is one example of a positive outcome resulting from the relationship. Noy & Zang (2023) also found an increase in productivity by the usage of GEN-AI. Consequently, the hypotheses are:

**H10:** Use positively and significantly influences net benefits.

**H11:** User Satisfaction positively and significantly influences net benefits.

Older users generally are less likely to enjoy new technologies. Though this user group is more invested in technologies such as smartphones than ever before, it is still naturally more difficult for older users to start using new technologies compared to other age groups. This also applies in the specific case of GEN-AI. The reason is mainly cognitive. A lower processing speed or attention span coupled with anxiety decreases the chances of making proper use of a new technology (Chattaraman et al., 2019). The final hypothesis of the model therefore states:

**H12:** Age negatively and significantly influences user satisfaction.

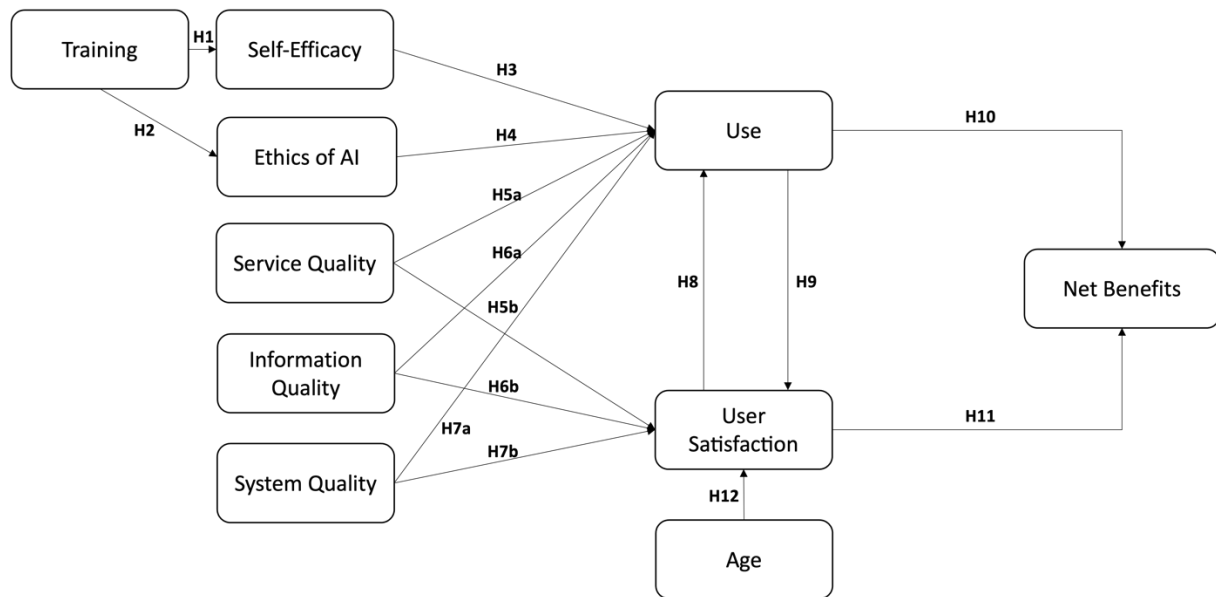


Figure 1 - Generative AI Adoption Impact Model

Figure 1 presents the proposed conceptual model of how training and self-efficacy as well as ethics and the variables defined by Delone & McLean (2003) influence the net benefits of a consultant on an individual level resulting from the adoption of a GEN-AI tool in their professional life.

## 4. EMPIRICAL STUDY AND RESULTS

### 4.1. SAMPLING DESIGN PROCESS

A questionnaire was developed to quantitatively validate the model. The measurement items were adopted from existing literature and are therefore statistically validated. Apart from that, the selection process focused on finding existing measurement items that would match the characteristics and tone of this study as well as possible. This led to a variety of study sources for the different measurement items. The questionnaire consists of 32 questions which are divided into 9 measurement items. The translated measurement items as well as the respective sources can be found in *Appendix B* of this study. General demographical questions and questions referring to the frequency of usage as well as an introduction text complete the questionnaire.

Before the distribution, the questionnaire was approved by the Ethics Committee of NOVA IMS. The approval is documented in *Appendix C*. Subsequently, a pilot study was conducted with twenty participants to validate the correctness and understandability of the questionnaire. Extracts of the final questionnaire design can be found in *Appendix D*. After successful validation in the pilot stage, the study was randomly distributed to several consulting companies, adopting a single cross-sectional design, providing a snapshot of participants' opinions and experiences (Allen, 2017). A broad variety of distribution channels was used to reach a statistically representative number of responses. Many participants were recruited via LinkedIn groups or direct messages. Furthermore, the study was published in internal chats of several consulting companies. On top of that, professors of NOVA IMS with connections to the consulting industry shared the survey in their network.

The survey was closed after collecting 201 valid responses within four weeks. Valid responses were complete and confirmed the participants' prior use of GEN-AI tools in their professional lives. Among respondents, 69% identified as male, 30% as female, and 1% chose not to disclose their gender. The average age of the study group, as stated by the participant in years, was 32, with a total range spanning from 18 to 64 years. A predominant 79.6% of respondents currently resided in Germany but the study group consisted of 20 residence nationalities in total. Notably, 89% listed ChatGPT as one of the GEN-AI technologies they utilized. Other represented tools were for instance Dall-E2, GitHub CoPilot, Bard, and AlphaCode. The participants had differing levels of experience, daily usage times, and application fields related to GEN-AI.

To evaluate the data, the partial least squares-structure equation modeling (PLS-SEM) method was leveraged. The PLS-SEM method is a popular approach when analyzing information systems. It is a statistical technique to test and estimate the causal relationships between several independent and dependent variables in structural models (Hair et al., 2011; Urbach & Ahlemann, 2010). Additionally, PLS-SEM is more flexible when it comes to model complexity

and sample size and allows the empirical validation of a non-normally distributed sample as required in this case (Hair et al., 2011). The skewness checks validated that the sample was skewed either left or right and therefore no normal distribution was given.

#### **4.2. MEASUREMENT MODEL RESULTS**

In the following paragraphs, the results of the first phase of the PLS-SEM method - the measurement analysis - including the reliability as well as the convergent and discriminant validity analysis are described in detail. The measurement analysis verifies if the single measurement items are measuring their specific construct and only that construct.

SmartPLS was used as a tool for the analysis. As a basis for the analysis, the conceptual model of this study was built within the tool.

Cronbach's alpha (CA) was used to validate the measurement items by assessing the internal consistency reliability of the test scale. According to the literature, CA must be greater than 0.7 since a value below that threshold suggests an inadequate internal consistency among the items and does therefore potentially compromise the reliability of the scale (Hair et al., 2011). Six measurement items did not fulfill this criterion and needed to be eliminated (E2, SE3, SE4, IQ4, SU2 and SU3). This means that not all study participants understood these measurement items in the same way which led to the low value. Consequently, because of the required deletion of SU2 and SU3, the variable system use became a univariate variable, only measured by a single item. After the deletion of the six items, all other items had a CA greater than 0.7 which meant that reliability was given.

Following the elimination of items, the analysis extended to examining cross-loadings to ensure that each indicator primarily measured its intended construct. The results, detailed in *Appendix E*, confirmed the appropriateness of the model. Additionally, the composite reliability was assessed, and values exceeding 0.7, as recommended by the literature, were observed. The average variance extracted (AVE) was assessed to further examine the convergent validity of the test scale. According to the literature, the value needs to be greater than 0.5 (Hair et al., 2021) which was given for all items. CA, AVE as well as the composite reliability can be observed in *Table 3* below.

Table 3 - Reliability and Validity Results

|                          | CA    | Composite reliability_a | Composite reliability_c | AVE   |
|--------------------------|-------|-------------------------|-------------------------|-------|
| Ethics_AI (E)            | 0.749 | 0.776                   | 0.887                   | 0.797 |
| Information Quality (IQ) | 0.813 | 0.814                   | 0.889                   | 0.728 |
| Net Benefits (NB)        | 0.871 | 0.88                    | 0.906                   | 0.661 |
| Self_Efficacy (SE)       | 0.747 | 0.781                   | 0.886                   | 0.796 |
| Service Quality (SEQ)    | 0.804 | 0.808                   | 0.872                   | 0.632 |
| System Quality (SYQ)     | 0.788 | 0.79                    | 0.877                   | 0.703 |
| Training (T)             | 0.849 | 0.866                   | 0.908                   | 0.766 |
| User Satisfaction (US)   | 0.86  | 0.863                   | 0.915                   | 0.781 |
| System Use (SU)          | 1     | 1                       | 1                       | 1     |
| Age (A)                  | 1     | 1                       | 1                       | 1     |

Subsequently, the discriminant validity was evaluated using the Fornell-Larcker criterion. Discriminant validity, in essence, explains how distinct a construct is from others within the model. The Fornell-Larcker criterion accomplishes this by scrutinizing the square root of the AVE for each construct and consequently comparing it within that construct and with other constructs in the model (Hair et al., 2011). This model under consideration successfully satisfied this criterion, as indicated in *Table 4*.

Table 4 - Fornell-Larcker Results

|     | A      | E            | IQ           | NB           | SE           | SEQ          | SYQ          | T            | SU       | US           |
|-----|--------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|----------|--------------|
| A   | 1      |              |              |              |              |              |              |              |          |              |
| E   | -0.072 | <b>0.893</b> |              |              |              |              |              |              |          |              |
| IQ  | -0.102 | 0.436        | <b>0.853</b> |              |              |              |              |              |          |              |
| NB  | -0.149 | 0.372        | 0.645        | <b>0.813</b> |              |              |              |              |          |              |
| SE  | -0.148 | 0.308        | 0.489        | 0.451        | <b>0.892</b> |              |              |              |          |              |
| SEQ | -0.075 | 0.247        | 0.541        | 0.486        | 0.439        | <b>0.795</b> |              |              |          |              |
| SYQ | -0.198 | 0.432        | 0.693        | 0.605        | 0.473        | 0.484        | <b>0.839</b> |              |          |              |
| T   | -0.051 | 0.16         | 0.35         | 0.35         | 0.44         | 0.451        | 0.222        | <b>0.875</b> |          |              |
| SU  | -0.091 | 0.213        | 0.428        | 0.608        | 0.421        | 0.413        | 0.449        | 0.344        | <b>1</b> |              |
| US  | -0.24  | 0.505        | 0.685        | 0.667        | 0.459        | 0.478        | 0.65         | 0.274        | 0.49     | <b>0.884</b> |

The assessment of the discriminant validity was further conducted using the heterotrait-monotrait (HTMT) ratio of correlations criterion. This criterion, established in the literature with a recommended threshold of 0.9, examines the extent to which constructs differ from each other in a model (Hair et al., 2021). Considering this threshold, all HTMT ratio values in the present study were found to be below 0.9, confirming the discriminant validity of the conceptual model. Detailed results are provided in *Table 5*.

Table 5 - HTMT Results

|     | A     | E     | IQ    | NB    | SE    | SEQ   | SYQ   | T     | SU    | US |
|-----|-------|-------|-------|-------|-------|-------|-------|-------|-------|----|
| A   |       |       |       |       |       |       |       |       |       |    |
| E   | 0.085 |       |       |       |       |       |       |       |       |    |
| IQ  | 0.111 | 0.555 |       |       |       |       |       |       |       |    |
| NB  | 0.16  | 0.451 | 0.762 |       |       |       |       |       |       |    |
| SE  | 0.169 | 0.426 | 0.63  | 0.562 |       |       |       |       |       |    |
| SEQ | 0.085 | 0.311 | 0.668 | 0.584 | 0.559 |       |       |       |       |    |
| SYQ | 0.224 | 0.554 | 0.865 | 0.727 | 0.618 | 0.606 |       |       |       |    |
| T   | 0.06  | 0.207 | 0.425 | 0.416 | 0.538 | 0.539 | 0.267 |       |       |    |
| SU  | 0.091 | 0.243 | 0.474 | 0.651 | 0.479 | 0.461 | 0.505 | 0.375 |       |    |
| US  | 0.259 | 0.633 | 0.816 | 0.76  | 0.575 | 0.572 | 0.789 | 0.321 | 0.526 |    |

Finally, the analysis incorporated the variance inflation factor (VIF), a measure designed to evaluate the extent of multicollinearity or degree of correlation among predictor variables (Hair et al., 2021). The recommended threshold to avoid multicollinearity issues is a VIF of less than five for all variables. In this study, all variables met this criterion which is documented in *Table 6*.

Table 6 - VIF Results

|     | Variables                                | VIF   |
|-----|------------------------------------------|-------|
| H1  | Training -> Self_Efficacy                | 1     |
| H2  | Trainings -> Ethics_AI                   | 1     |
| H3  | Self_Efficacy -> Use                     | 1.45  |
| H4  | Ethics_AI -> Use                         | 1.296 |
| H5a | Service Quality -> Use                   | 1.53  |
| H5b | Service Quality -> User Satisfaction     | 1.527 |
| H6a | Information Quality -> Use               | 2.308 |
| H6b | Information Quality -> User Satisfaction | 2.184 |
| H7a | System Quality -> Use                    | 2.127 |
| H7b | System Quality -> User Satisfaction      | 2.134 |
| H8  | User Satisfaction -> Use                 | 2.347 |
| H9  | Use -> User Satisfaction                 | 1.352 |
| H10 | Use -> Net Benefits                      | 1.316 |
| H11 | User Satisfaction -> Net Benefits        | 1.316 |
| H12 | Age -> User Satisfaction                 | 1.044 |

In conclusion, the conducted analysis demonstrates that the model is reliable and valid. The positive outcomes in terms of reliability and validity, coupled with the successful mitigation of multicollinearity issues confirm the robustness of the underlying model.

### 4.3. STRUCTURAL MODEL RESULTS

After determining the validity of the conceptual model, the focus now shifts to the analysis of the structural model as the second phase of the PLS-SEM method. This phase aims to evaluate the significance of the hypotheses and understand the contribution of each indicator. The p-

values are computed using bootstrapping with a sub-sample size of five thousand. The use of bootstrapping, a resampling technique generating multiple samples, enhances the robustness of the statistical analysis by providing a reliable assessment of significance (Hair et al., 2011).

The model encompasses a reciprocal relationship between the variables of system use and user satisfaction. To facilitate the analysis, the model was divided into model A and model B. Model A explores the influence of user satisfaction on system use while model B assesses the impact of system use on user satisfaction. The results of both models are indicated in *Table 7*. A thorough examination of the results in both tables reveals discrepancies between the two models, further explained in the subsequent paragraphs. The bootstrapping criteria outlined in both tables serve as benchmarks for the evaluation of a potential significance.

Table 7 - Research Hypotheses Results

| Hypotheses | Variables | $\beta$ -value |         | p-value |         | Significance | Conclusion    |
|------------|-----------|----------------|---------|---------|---------|--------------|---------------|
|            |           | Model A        | Model B | Model A | Model B |              |               |
| H1         | T -> SE   | 0.437          | 0.437   | 0       | 0       | ***          | supported     |
| H2         | T -> E    | 0.163          | 0.163   | 0.03    | 0.03    | *            | supported     |
| H3         | SE -> SU  | 0.184          | 0.205   | 0.01    | 0.004   | **           | supported     |
| H4         | E -> SU   | -0.087         | -0.028  | 0.21    | 0.693   | NS           | not supported |
| H5a        | SEQ-> SU  | 0.147          | 0.174   | 0.045   | 0.017   | *            | supported     |
| H5b        | SEQ -> US | 0.105          | 0.067   | 0.069   | 0.269   | NS           | not supported |
| H6a        | IQ -> SU  | 0.003          | 0.099   | 0.98    | 0.336   | NS           | not supported |
| H6b        | IQ -> US  | 0.418          | 0.393   | 0       | 0       | ***          | supported     |
| H7a        | SYQ -> SU | 0.14           | 0.211   | 0.149   | 0.033   | NS / *       | not supported |
| H7b        | SYQ -> US | 0.283          | 0.241   | 0       | 0.002   | *** / **     | supported     |
| H8         | US -> SU  | 0.287          | -       | 0.002   | -       | **           | supported     |
| H9         | SU -> US  | -              | 0.174   | -       | 0.002   | **           | supported     |
| H10        | SU -> NB  | 0.369          | 0.37    | 0       | 0       | ***          | supported     |
| H11        | US-> NB   | 0.486          | 0.486   | 0       | 0       | ***          | supported     |
| H12        | A -> US   | -0.133         | -0.131  | 0.015   | 0.013   | *            | supported     |

**Notes:** NS = not significant; \* significant at  $p < 0.05$ ; \*\* significant at  $p < 0.01$ ; \*\*\* significant at  $p < 0.001$  (Cohen, 1992)

The structural model utilizes  $R^2$ , a metric that quantifies the degree of explanation and therefore the proportion of variance in a dependent variable explained by an independent variable (Hair et al., 2011). The following paragraphs explain the results of the two models in detail, starting with the degree of explanation and ending with the significance of the different hypotheses.

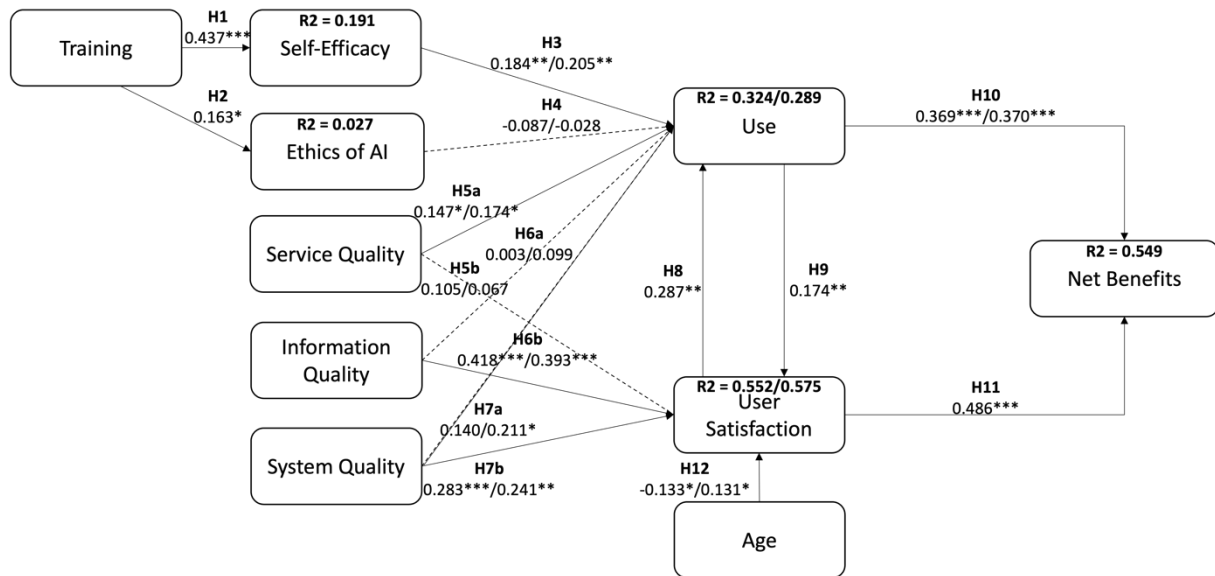


Figure 2 - Results Generative AI Adoption Impact Model

Figure 2 displays the hypotheses test results for model A and model B. In both models, the variable net benefits expresses a degree of explanation of 54.9%. Notably, the beta value for user satisfaction ( $\beta = 0.486$ ) exceeds that of system use ( $\beta = 0.37$ ), suggesting that net benefits are better explained by user satisfaction than by system use. The degree of explanation for system use and user satisfaction differs between the two models. In model A, system use is explained by 32.4% with self-efficacy ( $\beta = 0.184$ ) and service quality ( $\beta = 0.147$ ), where self-efficacy has a stronger impact. In model B, system use is explained by 28.9% with self-efficacy ( $\beta = 0.205$ ), service quality ( $\beta = 0.174$ ) as well as system quality ( $\beta = 0.211$ ) with the latter having the strongest influence. Remarkably, although the beta values of the independent variables are higher in model B, the total degree of explanation is lower. User satisfaction is impacted by information quality, system quality, and age in both models, with slight differences. In model A, the degree of explanation is 55.2%, and in model B it is 57.5%. In both models, information quality has the strongest impact on user satisfaction. For the variables, self-efficacy ( $R^2 = 0.191$ ) and ethics of AI ( $R^2 = 0.027$ ), which are both explained by training, the degree of explanation remains consistent between both models.

Subsequently, building on the insights gained from the degree of explanation, the significance of each hypothesis is determined based on the calculated p-values, with a threshold of 0.05 (Cohen, 1992). In model A, several hypotheses demonstrate statistical significance (H1, H2, H3, H5a, H6b, H7b, H8, H10, H11, and H12).

Thus, training significantly impacts self-efficacy ( $\beta = 0.437$ ,  $p < 0.001$ ) and ethics of AI ( $\beta = 0.163$ ,  $p < 0.05$ ). Furthermore, system use is positively and significantly influenced by self-efficacy ( $\beta = 0.184$ ,  $p < 0.01$ ), service quality ( $\beta = 0.147$ ,  $p < 0.05$ ), and user satisfaction ( $\beta = 0.287$ ,  $p < 0.01$ ). User satisfaction, in turn, is positively influenced by information quality ( $\beta = 0.418$ ,  $p < 0.001$ ) and system quality ( $\beta = 0.283$ ,  $p < 0.001$ ), while being negatively impacted by age ( $\beta = -0.133$ ,  $p < 0.05$ ). Net benefits are positively and significantly impacted by user satisfaction ( $\beta = 0.486$ ,  $p < 0.001$ ) and system use ( $\beta = 0.369$ ,  $p < 0.001$ ).

It is equally crucial to acknowledge non-significant findings which contribute to the overall understanding of the model. In model A, H4, H5b, H6a, and H7a did not have a significant impact on their independent variables. Consequently, ethics of AI ( $\beta = -0.087$ ,  $p > 0.05$ ), system quality ( $\beta = 0.140$ ,  $p > 0.05$ ) and information quality ( $\beta = 0.003$ ,  $p > 0.05$ ) do not significantly impact system use. Additionally, service quality ( $\beta = 0.105$ ,  $p > 0.05$ ) does not significantly impact user satisfaction.

The findings prove noteworthy relationships within this model. Model B, while being largely consistent with model A and only containing differences in certain values which can be observed in detail in *Figure 2*, contains an additional significant hypothesis. Hypotheses 7a, indicating a significant impact of system quality on system use when system use impacts user satisfaction, is distinct ( $\beta = 0.283$ ,  $p < 0.001$ ). However, if user satisfaction impacts use, system quality does not have a significant impact on system use. Moreover, it is important to mention that hypothesis H9, the bidirectional counterpart to H8 in model A is significant. Hence, user satisfaction is significantly and positively influenced by system use ( $\beta = 0.174$ ,  $p < 0.01$ ).

In summary, the analysis revealed significant influences while non-significant findings offered insights into the limits of certain variables. The consistency between the models strengthens the findings, and the unique significance of system quality on system use in model B adds an interesting dimension for further discussion in the next chapter.

## 5. DISCUSSION

This chapter engages in a comprehensive discussion of the previously presented results, drawing comparisons with existing literature, and highlighting both, points of alignment and contradiction. Notably, Yu & Qian (2018) expanded the D&M model of information system success, incorporating training and self-efficacy as measurement items, revealing their significant influence on system use. The current study confirms the significance of these factors within the GEN-AI context. Additionally, the outcome lends general support to the relevance of training in both, research, and practice for a successful implementation of this technology, as suggested by prior studies (Budhwar et al., 2023; Shih et al., 2021). Furthermore, the noteworthy impact of the variable age on user satisfaction implies that older users may experience less favorable interactions with GEN-AI. This aligns with the broader observation that older adults tend to lag in the adoption of AI (Chattaraman et al., 2019).

Moreover, the finding that ethics have a negative yet nonsignificant impact on system use aligns with the literature, indicating that employees may not currently prioritize ethical considerations when dealing with GEN-AI (Siau & Wang, 2020). Surprisingly, this trend is observed even though ethical behavior is essential for a good relationship with the client in the consulting industry (Mingaleva, 2013). The disconnect between the importance pointed out in research (Baabdullah, 2024; Wach et al., 2023; Zhou et al., 2020) and its missing practical application highlights a gap that needs to be addressed. Moreover, the necessary elimination of the item E2 measuring transparency due to its low CA value, despite its high relevance in research (Clinciu & Hastie, 2019; Dwivedi et al., 2023) points to the fact that the topic of AI explainability is not yet fully understood by people and hence leads to confusion and different interpretations.

The application of the D&M model to the context of GEN-AI provides interesting results. Service quality, according to the results of this study, has a significant impact on system use but not on user satisfaction. This implies that services such as a qualified and available user support or an easy API integration make the adoption of the system more likely. This aligns with the assumption of Delone & McLean (2003) that service quality is a relevant item to consider. However, the finding challenges the direct linkage between service quality and user satisfaction by Prentice et al. (2020) and suggests that while service quality might encourage initial adoption, it is not as influential in improving the overall satisfaction that a user has with the tool. Since service quality has been measured differently in different studies, a significance cannot always be detected (Petter et al., 2008; Urbach et al., 2010) but certain studies confirm the findings of this research (Nunes & Javier, 2014).

The results of this study indicate that information quality does not significantly influence use but user satisfaction. Factors related to this variable such as relevance or timeliness therefore do not encourage the user to initially engage with the system. Nevertheless, the quality of

information does contribute to the overall satisfaction with the system. This correlation aligns with the findings of another study in the context of employee portals (Urbach et al., 2010). Supporting this statement, Petter et al. (2008) claim that the impact of information quality on user satisfaction is strongly validated, while the one with system use is more fragile and therefore not consistently significant. Since users can generally interpret and enhance the information quality only after adopting the system, this makes sense (Dwivedi et al., 2023). However, Wu & Wang (2006) found a significant impact on both, system use and user satisfaction. The results of this study are also inconsistent with the findings of Dwivedi et al. (2023) who reported a positive influence of information quality on the likelihood of GEN-AI system usage. Moreover, given the relevance of information quality in the consulting industry (Oesterle et al., 2016) one would anticipate a significant impact on both independent variables. Surprisingly, this expectation is not met in this study, marking an interesting inconsistency with broader literature.

The evaluation of system quality, measuring the characteristics of the system, is a specifically interesting case due to the observed variations between model A and model B. If system use impacts user satisfaction, system quality has a significant impact on both, system use and user satisfaction. However, if user satisfaction impacts system use, system quality only impacts user satisfaction. In the latter scenario, system quality's impact on system use is indirect, meaning that only a satisfied user is more likely to continue using the system. Other studies emphasize the complex dimensions of system quality by providing differing results. Urbach et al. (2010) consider it the most significant dimension explaining user satisfaction which does not align with the results of this study in which information quality has a stronger impact. Additionally, there are studies in which system quality does not have a significant impact at all (Aparicio et al., 2019) or aligns with model A (Urbach et al., 2010). Especially in the context of GEN-AI, studies confirm a positive effect of ease of use on satisfaction, engagement, and overall success (Dwivedi et al., 2023; Prentice et al., 2020), partly agreeing with this study.

The interconnected relationship between system use, characterized by the frequency and duration of interactions and user satisfaction is evident, as both variables positively and significantly impact each other. This positive feedback loop, identified by Delone & McLean in their 2003 model and supported by Urbach et al. (2010), contributes to the overall success of the system. However, it is worth noting that not all studies observe this interconnection (Aparicio et al., 2019). In the context of GEN-AI, where system interaction enhances personalization and therefore likely satisfaction with the system due to a higher relevance, this interconnected relationship becomes even more apparent (Guerreiro & Loureiro, 2023).

The results most closely related to the research question of this study originate from the net benefits variable, significantly influenced by both, system use and user satisfaction. This strong impact aligns with findings from several other studies (Aparicio et al., 2019; Urbach et al., 2010), including a significant impact on task and overall job performance (Petter et al., 2008)

and therefore individual productivity which is the definition of net benefits in this study. Factors such as engaging work or opportunities for advancement enhance individual motivation and subsequently productivity. The utilization of a disruptive technology like GEN-AI, coupled with the curiosity of employees, suggests one aspect of how system use influences net benefits (Sabir, 2017). These findings also align with recent scientific experiments, exploring the productivity effects of GEN-AI in terms of quality and efficiency (Noy & Zhang, 2023). However, if the task itself is too advanced for the current state of GEN-AI, the productivity effect can be reversed which conflicts with the results of this study (Dell'Acqua et al., 2023).

In conclusion, the application of the D&M model to the GEN-AI context highlights insightful findings. The identified gaps and variations contribute to a proper understanding of the technology's dynamics and provide potential future enhancements to the theoretical framework further explained in the following chapter.

## **6. CONCLUSIONS AND FUTURE WORKS**

### **6.1. CONCLUSIONS**

The main purpose of this study was to determine if GEN-AI enhances the productivity of consultants. To achieve this objective, the determinants influencing productivity in the consulting industry were identified. Consequently, a comprehensive model was developed and empirically tested with a survey approach and the PLS-SEM statistical analysis. The results confirm that GEN-AI has indeed a significant and positive impact on the productivity of consultants underscoring the technology's potential as a transformative and supportive tool in the consulting industry.

The outcome of this research not only enhances the theoretical understanding of information system success within the GEN-AI context but also offers practical insights for organizations aiming to boost the productivity of consultants through strategic AI integration. As organizations navigate the adoption of advanced technologies, these findings can serve as a valuable guide, facilitating decision-making and providing a foundation for future research in the dynamic landscape of AI applications in the industry.

### **6.2. IMPLICATIONS AND FUTURE WORKS**

The proposed implications offer a valuable understanding for researchers as well as practitioners involved in the study and implementation of GEN-AI technologies to refine theoretical models or maximize the benefit of GEN-AI for an organization.

#### **6.2.1. Theoretical Implications**

The extension of the D&M model for information system success with the variables training, self-efficacy, and ethics as well as its adaptation to the context of GEN-AI proves the adaptability and versatility of the established model to new disruptive technologies. The identification of age as a significant factor impacting user satisfaction theoretically highlights age-specific challenges in the adoption of GEN-AI. This underscores the importance of incorporating demographic variables in comprehending system success. Theoretical models might need to be refined in certain circumstances based on this understanding.

The observed variations in the impact of system quality in different scenarios contribute to the theoretical understanding of the rather complex dimension. It emphasizes that the influence on user satisfaction and system use is not consistent and may be context-dependent. The theoretical implications suggest that a generalized approach to understanding the dimension in theoretical models may not capture its diverse effects accurately. Therefore, a well-balanced and context-sensitive approach is required to assess the multifaced impact of system quality within the GEN-AI landscape.

The study's findings suggest that in the context of GEN-AI, information quality demonstrates a significant impact, influencing user satisfaction rather than initial system use. This challenges prevailing assumptions about the immediate role of information quality in system adoption. It implies that users may assess and appreciate information quality post-adoption, emphasizing its long-term contribution to user satisfaction. Moreover, service quality's theoretical implications highlight its distinctive impact on encouraging system use rather than directly influencing user satisfaction. The study underscores the need for theoretical models to differentiate between factors driving initial adoption and those shaping prolonged user satisfaction.

The validated positive feedback loop between system use and user satisfaction, as revealed in this study, adds to a deeper understanding of the dynamics between the two. This offers theoretical insights into the contributions of these dimensions to the overall success of the system, thereby supporting and expanding existing theories.

### **6.2.2. Practical Implications**

The practical implications derived from the study underscore the complexity of successfully integrating GEN-AI into the organization. To maximize the benefits and ensure sustainable productivity growth, organizations should consider a comprehensive approach that goes beyond the initial adoption of the technology.

This involves the prioritization of tailored and in-depth training initiatives to enhance both, technical aspects, and a profound understanding of ethical issues. The substantial influence of age as a variable on user satisfaction indicates that older users may have a less positive experience with GEN-AI, emphasizing the need for tailored training sessions for different age groups or even tailored user interfaces or support mechanisms for older users. The detected lack of knowledge of AI explainability also requires further training in this specific context. This involves the development of training programs either in-house or with an agency. Resources need to focus on the specific needs of employees interacting with GEN-AI in that organization. Moreover, practical measures should be taken to bridge the ethics awareness gap. Besides training, this includes incorporating ethical considerations throughout the entire lifecycle of GEN-AI technologies, from deployment to daily usage, and decision-making processes. Practical guidelines or frameworks can support embedding ethics into organizational operational practices.

Highlighting the importance of service quality during the initial adoption stages of the technology suggests a focus on this dimension to promote engagement with GEN-AI. This includes investing in automated support services, real-time assistance, or user-friendly manuals to guarantee sustained engagement with the technology.

Since the proper quality of information is a success criterion in the consulting industry, measures should be taken to ensure the latter provided by GEN-AI tools. Standards, protocols, or codes of conduct should be established to guarantee information quality, ultimately enhancing user satisfaction not only on the consultant but also on the client side and positioning GEN-AI as a reliable and trustworthy source in the industry.

However, to achieve the best possible increase in productivity, a company-wide GEN-AI strategy is essential including the previously mentioned measures. Hereby, an ongoing monitoring of GEN-AI capabilities, coupled with a constant review process and continuous competence development prevents a decrease in quality and productivity, ensuring the sustainable success of GEN-AI within the organization.

### **6.2.3. Limitations and Future Works**

This study is subject to certain limitations. The continuous and rapid development of GEN-AI results in a growing number of tools. Hence, employees use several applications for a broad range of tasks. This creates a potential bias in the survey responses due to the diversity of tools and tasks being evaluated which can influence the survey's applicability.

Furthermore, the demographic sample of the study has a high concentration of male participants as well as respondents from Germany and is therefore skewed. This may limit the generalizability of the findings to a broader population or a different cultural context. The study's focus on the consulting industry with its specific characteristics further narrows its generalizability. Hence, the results cannot be applied to every other professional context. Future research can enhance the validity of the results by diving deeper into the consulting industry itself or comparing them to different professional settings or demographics as comparative studies.

Additionally, the study only reflects opinions at a certain moment in time. Longitudinal studies could offer further insights into the evolution of perceptions over time. Likewise, this study specified net benefits as the individual productivity of a consultant. In future research, the dimension could be interpreted on an organizational level or even a societal one.

Determining that ethical considerations have a negative yet non-significant impact on system use partially addresses one of the previously identified research gaps. However, the lack of significance suggests a potential gap in the theoretical understanding and practical prioritization of ethical considerations. Additional research is needed to deepen the comprehension of the ethical perception within the GEN-AI context and bridge the gap between theoretical concepts and practical implementation.

Lastly, the data collection method using a questionnaire can potentially introduce a response bias by relying on subjective responses. Subsequent research could explore alternative validation approaches, utilizing different methodologies to enhance the robustness of the findings.

## BIBLIOGRAPHICAL REFERENCES

- Alexandre, C., & Blanckaert, L. (2020). *The Influence of Artificial Intelligence on The Consulting Industry* [Master Thesis]. Université catholique de Louvain.
- Allen, M. (2017). Cross-Sectional Design. In *The SAGE Encyclopedia of Communication Research Methods*. SAGE Publications, Inc. <https://doi.org/10.4135/9781483381411.n118>
- Aparicio, M., Oliveira, T., Bacao, F., & Painho, M. (2019). Gamification: A key determinant of massive open online course (MOOC) success. *Information and Management*, 56(1), 39–54. <https://doi.org/10.1016/j.im.2018.06.003>
- Baabdullah, A. M. (2024). Generative conversational AI agent for managerial practices: The role of IQ dimensions, novelty seeking and ethical concerns. *Technological Forecasting and Social Change*, 198, 122951. <https://doi.org/10.1016/j.techfore.2023.122951>
- Baily, M. N., & Solow, R. M. (2001). International Productivity Comparisons Built from the Firm Level. *Journal of Economic Perspectives*, 15(3), 151–172. <https://doi.org/10.1257/jep.15.3.151>
- Bayati, M. (2019). *AI in Consulting* [Bachelor Thesis]. Haaga-Helia University of Applied Sciences.
- Björkman, M. (1992). What is Productivity? *IFAC Proceedings Volumes*, 25(8), 203–210. [https://doi.org/10.1016/S1474-6670\(17\)54065-3](https://doi.org/10.1016/S1474-6670(17)54065-3)
- Bode, M., Daneva, M., & van Sinderen, M. J. (2022). Characterising the digital transformation of IT consulting services—Results from a systematic mapping study. *IET Software*, 16(5), 455–477. <https://doi.org/10.1049/sfw2.12068>
- Bologa, R., & Lupu, A. R. (2014). Organizational learning networks that can increase the productivity of IT consulting companies. A case study for ERP consultants. *Expert Systems with Applications*, 41(1), 126–136. <https://doi.org/10.1016/j.eswa.2013.07.016>
- Borenstein, J., & Howard, A. (2021). Emerging challenges in AI and the need for AI ethics education. *AI and Ethics*, 1(1), 61–65. <https://doi.org/10.1007/s43681-020-00002-7>
- Brynjolfsson, E. (1993). The productivity paradox of information technology. *Communications of the ACM*, 36(12), 66–77. <https://doi.org/10.1145/163298.163309>
- Budhwar, P., Chowdhury, S., Wood, G., Aguinis, H., Bamber, G. J., Beltran, J. R., Boselie, P., Lee Cooke, F., Decker, S., DeNisi, A., Dey, P. K., Guest, D., Knoblich, A. J., Malik, A., Paauwe, J., Papagiannidis, S., Patel, C., Pereira, V., Ren, S., ... Varma, A. (2023). Human resource management in the age of generative artificial intelligence: Perspectives and research directions on ChatGPT. *Human Resource Management Journal*, 33(3), 606–659. <https://doi.org/10.1111/1748-8583.12524>
- Chattaraman, V., Kwon, W.-S., Gilbert, J. E., & Ross, K. (2019). Should AI-Based, conversational digital assistants employ social- or task-oriented interaction style? A task-competency and reciprocity perspective for older adults. *Computers in Human Behavior*, 90, 315–330. <https://doi.org/10.1016/j.chb.2018.08.048>

- Chiu, Y.-T., Zhu, Y.-Q., & Corbett, J. (2021). In the hearts and minds of employees: A model of pre-adoptive appraisal toward artificial intelligence in organizations. *International Journal of Information Management*, 60, 102379. <https://doi.org/10.1016/j.ijinfomgt.2021.102379>
- Cliniciu, M.-A., & Hastie, H. (2019). A Survey of Explainable AI Terminology. *Proceedings of the 1st Workshop on Interactive Natural Language Technology for Explainable Artificial Intelligence (NL4XAI 2019)*, 8–13. <https://doi.org/10.18653/v1/W19-8403>
- Cockburn, I. M., Henderson, R., & Stern, S. (2019). The Impact of Artificial Intelligence on Innovation An Exploratory Analysis. In *The Economics of Artificial Intelligence* (pp. 115–146). University of Chicago Press. <https://doi.org/10.7208/chicago/9780226613475.001.0001>
- Cohen, J. (1992). Statistical Power Analysis. *Current Directions in Psychological Science*, 1(3), 98–101. <https://doi.org/10.1111/1467-8721.ep10768783>
- Collins, C., Dennehy, D., Conboy, K., & Mikalef, P. (2021). Artificial intelligence in information systems research: A systematic literature review and research agenda. *International Journal of Information Management*, 60, 102383. <https://doi.org/10.1016/j.ijinfomgt.2021.102383>
- Crișan, E. L., & Stanca, L. (2021). The digital transformation of management consulting companies: a qualitative comparative analysis of Romanian industry. *Information Systems and E-Business Management*, 19(4), 1143–1173. <https://doi.org/10.1007/s10257-021-00536-1>
- Damioli, G., Van Roy, V., & Vertesy, D. (2021). The impact of artificial intelligence on labor productivity. *Eurasian Business Review*, 11(1), 1–25. <https://doi.org/10.1007/s40821-020-00172-8>
- Dell’Acqua, F., McFowland, 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>
- DeLone, W. H., & McLean, E. R. (1992). Information Systems Success: The Quest for the Dependent Variable. *Information Systems Research*, 3(1), 60–95. <https://doi.org/10.1287/isre.3.1.60>
- Delone, W. H., & Mclean, E. R. (2003). The DeLone and McLean Model of Information Systems Success: A Ten-Year Update. In *Information Systems Research, Journal of Management Information Systems* (Vol. 19, Issue 4). <https://doi.org/10.1080/07421222.2003.11045748>
- Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., Baabdullah, A. M., Koohang, A., Raghavan, V., Ahuja, M., Albanna, H., Albashrawi, M. A., Al-Busaidi, A. S., Balakrishnan, J., Barlette, Y., Basu, S., Bose, I., Brooks, L., Buhalis, D., ... Wright, R. (2023). Opinion Paper: “So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research,

- practice and policy. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- Ebert, C., & Louridas, P. (2023). Generative AI for Software Practitioners. *IEEE Software*, 40(4), 30–38. <https://doi.org/10.1109/MS.2023.3265877>
- Elsdaig, M. (2019). Evaluation of Healthcare Information System Using Delone and McLean Quality Model, Case study KSA. *International Journal of Advanced Trends in Computer Science and Engineering*, 8(1.4), 522–527. <https://doi.org/10.30534/ijatcse/2019/8181.42019>
- Fenwick, A., & Molnar, G. (2022). The importance of humanizing AI: using a behavioral lens to bridge the gaps between humans and machines. *Discover Artificial Intelligence*, 2(1), 14. <https://doi.org/10.1007/s44163-022-00030-8>
- Fosso Wamba, S., Queiroz, M. M., Chiappetta Jabbour, C. J., & Shi, C. (Victor). (2023). Are both generative AI and ChatGPT game changers for 21st-Century operations and supply chain excellence? *International Journal of Production Economics*, 265, 109015. <https://doi.org/10.1016/j.ijpe.2023.109015>
- Fügener, A., Grahl, J., Gupta, A., & Ketter, W. (2022). Cognitive Challenges in Human–Artificial Intelligence Collaboration: Investigating the Path Toward Productive Delegation. *Information Systems Research*, 33(2), 678–696. <https://doi.org/10.1287/isre.2021.1079>
- Garrett, N., Beard, N., & Fiesler, C. (2020). More Than “If Time Allows.” *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*, 272–278. <https://doi.org/10.1145/3375627.3375868>
- Gînguță, A., Ștefea, P., Noja, G. G., & Munteanu, V. P. (2023). Ethical Impacts, Risks and Challenges of Artificial Intelligence Technologies in Business Consulting: A New Modelling Approach Based on Structural Equations. *Electronics*, 12(6), 1462. <https://doi.org/10.3390/electronics12061462>
- Gries, T., & Naudé, W. (2018). Artificial Intelligence, Jobs, Inequality and Productivity: Does Aggregate Demand Matter? *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.3301777>
- Grosskopf, S. (1993). Efficiency and Productivity. In H. O. Fried, C. A. K. Lovell, & S. S. Schmidt (Eds.), *The Measurement of Productive Efficiency and Productivity Change* (pp. 160–196). Oxford University Press. <https://doi.org/10.1093/acprof:oso/9780195183528.001.0001>
- Guerreiro, J., & Loureiro, S. M. C. (2023). I am attracted to my Cool Smart Assistant! Analyzing Attachment-Aversion in AI-Human Relationships. *Journal of Business Research*, 161. <https://doi.org/10.1016/j.jbusres.2023.113863>
- Gyllensvärd, F., & Hiljegren, O. (2023). *Generative AI as a Tool for Swedish Startups - A qualitative study on how Generative AI can affect Swedish startups* [Master Thesis]. Göteborg Universitet.
- Hagendorff, T. (2020). The Ethics of AI Ethics: An Evaluation of Guidelines. *Minds and Machines*, 30(1), 99–120. <https://doi.org/10.1007/s11023-020-09517-8>

- Hair, J. F., Hult, G. T. M., Ringle, C. M., Sarstedt, M., Danks, N. P., & Ray, S. (2021). *Partial Least Squares Structural Equation Modeling (PLS-SEM) Using R*. Springer International Publishing. <https://doi.org/10.1007/978-3-030-80519-7>
- Hair, J. F., Ringle, C. M., & Sarstedt, M. (2011). PLS-SEM: Indeed a Silver Bullet. *Journal of Marketing Theory and Practice*, 19(2), 139–152. <https://doi.org/10.2753/MTP1069-6679190202>
- Haleem, A., Javaid, M., & Singh, R. P. (2022). An era of ChatGPT as a significant futuristic support tool: A study on features, abilities, and challenges. *BenchCouncil Transactions on Benchmarks, Standards and Evaluations*, 2(4), 100089. <https://doi.org/10.1016/j.tbench.2023.100089>
- Hanaysha, J. (2016a). Examining the Effects of Employee Empowerment, Teamwork, and Employee Training on Organizational Commitment. *Procedia - Social and Behavioral Sciences*, 229, 298–306. <https://doi.org/10.1016/j.sbspro.2016.07.140>
- Hanaysha, J. (2016b). Testing the Effects of Employee Empowerment, Teamwork, and Employee Training on Employee Productivity in Higher Education Sector. *International Journal of Learning and Development*, 6(1), 164. <https://doi.org/10.5296/ijld.v6i1.9200>
- Harrer, S. (2023). Attention is not all you need: the complicated case of ethically using large language models in healthcare and medicine. *EBioMedicine*, 90, 104512. <https://doi.org/10.1016/j.ebiom.2023.104512>
- Howard, M. C. (2014). Creation of a Computer Self-Efficacy Measure: Analysis of Internal Consistency, Psychometric Properties, and Validity. *Cyberpsychology, Behavior, and Social Networking*, 17(10), 677–681. <https://doi.org/10.1089/cyber.2014.0255>
- Hung, S., & Liang, T. (2001). Effect of computer self-efficacy on the use of executive support systems. *Industrial Management & Data Systems*, 101(5), 227–237. <https://doi.org/10.1108/02635570110394626>
- Hyun Baek, T., & Kim, M. (2023). Is ChatGPT scary good? How user motivations affect creepiness and trust in generative artificial intelligence. *Telematics and Informatics*, 83, 102030. <https://doi.org/10.1016/j.tele.2023.102030>
- Jaiswal, A., Arun, C. J., & Varma, A. (2022). Rebooting employees: upskilling for artificial intelligence in multinational corporations. *The International Journal of Human Resource Management*, 33(6), 1179–1208. <https://doi.org/10.1080/09585192.2021.1891114>
- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1(9), 389–399. <https://doi.org/10.1038/s42256-019-0088-2>
- Koh, E., & Doroudi, S. (2023). Learning, teaching, and assessment with generative artificial intelligence: towards a plateau of productivity. *Learning: Research and Practice*, 9(2), 109–116. <https://doi.org/10.1080/23735082.2023.2264086>
- Kshetri, N., Dwivedi, Y. K., Davenport, T. H., & Panteli, N. (2023). Generative artificial intelligence in marketing: Applications, opportunities, challenges, and research agenda. *International Journal of Information Management*, 102716. <https://doi.org/10.1016/j.ijinfomgt.2023.102716>

- Mingaleva, Z. (2013). Ethical Principles in Consulting. *Procedia - Social and Behavioral Sciences*, 84, 1740–1744. <https://doi.org/10.1016/j.sbspro.2013.07.024>
- Mohammadi, H. (2015). Investigating users' perspectives on e-learning: An integration of TAM and IS success model. *Computers in Human Behavior*, 45, 359–374. <https://doi.org/10.1016/j.chb.2014.07.044>
- Nah, F. F.-H., Zheng, R., Cai, J., Siau, K., & Chen, L. (2023). Generative AI and ChatGPT: Applications, challenges, and AI-human collaboration. *Journal of Information Technology Case and Application Research*, 25(3), 277–304. <https://doi.org/10.1080/15228053.2023.2233814>
- 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>
- Nunes, G. S., & Javier, M. G. F. (2014). Hospital Information System Satisfaction in Brazil: Background and Moderating Effects. *International Journal of Research Foundation of Hospital and Healthcare Administration*, 2(1), 1–9. <https://doi.org/10.5005/jp-journals-10035-1007>
- Oesterle, S., Buchwald, A., & Urbach, N. (2016). Understanding the Co-Creation of Value Emerging from the Collaboration between IT Consulting Firms and their Customers. *Thirty Seventh International Conference on Information Systems*.
- Ojo, A. I. (2017). Validation of the DeLone and McLean Information Systems Success Model. *Healthcare Informatics Research*, 23(1), 60. <https://doi.org/10.4258/hir.2017.23.1.60>
- Petter, S., DeLone, W., & McLean, E. (2008). Measuring information systems success: Models, dimensions, measures, and interrelationships. *European Journal of Information Systems*, 17(3), 236–263. <https://doi.org/10.1057/ejis.2008.15>
- Prentice, C., Weaven, S., & Wong, I. A. (2020). Linking AI quality performance and customer engagement: The moderating effect of AI preference. *International Journal of Hospitality Management*, 90, 102629. <https://doi.org/10.1016/j.ijhm.2020.102629>
- Riedl, M. O. (2019). Human-centered artificial intelligence and machine learning. *Human Behavior and Emerging Technologies*, 1(1), 33–36. <https://doi.org/10.1002/hbe2.117>
- Sabir, A. (2017). Motivation: Outstanding Way to Promote Productivity in Employees. *American Journal of Management Science and Engineering*, 2(3), 35. <https://doi.org/10.11648/j.ajmse.20170203.11>
- Sætra, H. S. (2023). Generative AI: Here to stay, but for good? *Technology in Society*, 75. <https://doi.org/10.1016/j.techsoc.2023.102372>
- Schultz, M. D., & Seele, P. (2023). Towards AI ethics' institutionalization: knowledge bridges from business ethics to advance organizational AI ethics. *AI and Ethics*, 3(1), 99–111. <https://doi.org/10.1007/s43681-022-00150-y>
- Shih, P.-K., Lin, C.-H., Wu, L. Y., & Yu, C.-C. (2021). Learning Ethics in AI—Teaching Non-Engineering Undergraduates through Situated Learning. *Sustainability*, 13(7), 3718. <https://doi.org/10.3390/su13073718>

- Siau, K., & Wang, W. (2020). Artificial Intelligence (AI) Ethics. *Journal of Database Management*, 31(2), 74–87. <https://doi.org/10.4018/JDM.2020040105>
- Strich, F., Mayer, A.-S., & Fiedler, M. (2021). What Do I Do in a World of Artificial Intelligence? Investigating the Impact of Substitutive Decision-Making AI Systems on Employees' Professional Role Identity. *Journal of the Association for Information Systems*, 22(2), 304–324. <https://doi.org/10.17705/1jais.00663>
- Tavoletti, E., Kazemargi, N., Cerruti, C., Grieco, C., & Appolloni, A. (2021). Business model innovation and digital transformation in global management consulting firms. *European Journal of Innovation Management*, 25(6), 612–636. <https://doi.org/10.1108/EJIM-11-2020-0443>
- Torkzadeh, G., & Doll, W. J. (1999). The development of a tool for measuring the perceived impact of information technology on work. *Omega*, 27(3), 327–339. [https://doi.org/10.1016/S0305-0483\(98\)00049-8](https://doi.org/10.1016/S0305-0483(98)00049-8)
- Urbach, N., & Ahlemann, F. (2010). Structural Equation Modeling in Information Systems Research Using Partial Least Squares. *Journal of Information Technology Theory and Application (JITTA)*, 11(2), 5–40.
- Urbach, N., & Müller, B. (2012). The Updated DeLone and McLean Model of Information Systems Success. In *Information Systems Theory* (pp. 1–18). [https://doi.org/10.1007/978-1-4419-6108-2\\_1](https://doi.org/10.1007/978-1-4419-6108-2_1)
- Urbach, N., Smolnik, S., & Riempp, G. (2010). An empirical investigation of employee portal success. *The Journal of Strategic Information Systems*, 19(3), 184–206. <https://doi.org/10.1016/j.jsis.2010.06.002>
- Venkatesh, V., & Bala, H. (2008). Technology Acceptance Model 3 and a Research Agenda on Interventions. *Decision Sciences*, 39(2), 273–315. <https://doi.org/10.1111/j.1540-5915.2008.00192.x>
- Wach, K., Duong, C. D., Ejdy, J., Kazlauskaitė, R., Korzynski, P., Mazurek, G., Paliszkievicz, J., & Ziemba, E. (2023). The dark side of generative artificial intelligence: A critical analysis of controversies and risks of ChatGPT. *Entrepreneurial Business and Economics Review*, 11(2), 7–30. <https://doi.org/10.15678/eber.2023.110201>
- wael AL-khatib, A. (2023). Drivers of generative artificial intelligence to fostering exploitative and exploratory innovation: A TOE framework. *Technology in Society*, 75, 102403. <https://doi.org/10.1016/j.techsoc.2023.102403>
- Walkowiak, E. (2023). Task-interdependencies between Generative AI and Workers. *Economics Letters*, 231. <https://doi.org/10.1016/j.econlet.2023.111315>
- Wang, Y.-S., Wang, H.-Y., & Shee, D. Y. (2007). Measuring e-learning systems success in an organizational context: Scale development and validation. *Computers in Human Behavior*, 23(4), 1792–1808. <https://doi.org/10.1016/j.chb.2005.10.006>
- Wu, J.-H., & Wang, Y.-M. (2006). Measuring KMS success: A respecification of the DeLone and McLean's model. *Information & Management*, 43(6), 728–739. <https://doi.org/10.1016/j.im.2006.05.002>

- Yaghmaie, F., & Jayasuriya, R. (2004). The roles of “subjective computer training” and management support in the use of computers in community health centres. *Journal of Innovation in Health Informatics*, 12(3), 163–170. <https://doi.org/10.14236/jhi.v12i3.122>
- Yi, M. Y., & Hwang, Y. (2003). Predicting the use of web-based information systems: self-efficacy, enjoyment, learning goal orientation, and the technology acceptance model. *International Journal of Human-Computer Studies*, 59(4), 431–449. [https://doi.org/10.1016/S1071-5819\(03\)00114-9](https://doi.org/10.1016/S1071-5819(03)00114-9)
- Yu, L., & Li, Y. (2022). Artificial Intelligence Decision-Making Transparency and Employees’ Trust: The Parallel Multiple Mediating Effect of Effectiveness and Discomfort. *Behavioral Sciences*, 12(5), 127. <https://doi.org/10.3390/bs12050127>
- Yu, P., Li, H., & Gagnon, M.-P. (2009). Health IT acceptance factors in long-term care facilities: A cross-sectional survey. *International Journal of Medical Informatics*, 78(4), 219–229. <https://doi.org/10.1016/j.ijmedinf.2008.07.006>
- Yu, P., & Qian, S. (2018). Developing a theoretical model and questionnaire survey instrument to measure the success of electronic health records in residential aged care. *PLOS ONE*, 13(1), e0190749. <https://doi.org/10.1371/journal.pone.0190749>
- Zhou, J., Chen, F., Berry, A., Reed, M., Zhang, S., & Savage, S. (2020). A Survey on Ethical Principles of AI and Implementations. *2020 IEEE Symposium Series on Computational Intelligence (SSCI)*, 3010–3017. <https://doi.org/10.1109/SSCI47803.2020.9308437>

## APPENDIX A. SYSTEM SUCCESS DIMENSIONS IN EMPIRICAL STUDIES

The table presents several empirical studies verifying the D&M model in a variety of use cases and contexts.

| Reference               | Topic                         | SEQ | IQ | SYQ | SU | US | NB | Other Constructs                                                                                         |
|-------------------------|-------------------------------|-----|----|-----|----|----|----|----------------------------------------------------------------------------------------------------------|
| (Wu & Wang, 2006)       | Knowledge management system   |     | x  | x   | x  | x  | x  | -                                                                                                        |
| (Wang et al., 2007)     | E-learning system             | x   | x  | x   | x  | x  | x  | -                                                                                                        |
| (Urbach et al., 2010)   | Employee portal               | x   | x  | x   | x  | x  | x  | Process quality, collaboration quality, knowledge intensity, process standardization, management support |
| (Mohammadi, 2015)       | E-learning                    | x   | x  | x   | x  | x  |    | Educational quality, perceived ease of use, perceived usefulness, intention to use                       |
| (Ojo, 2017)             | Hospital information system   | x   | x  | x   | x  | x  | x  | -                                                                                                        |
| (Yu & Qian, 2018)       | Electronic health records     |     | x  | x   | x  | x  | x  | Training, self-efficacy                                                                                  |
| (Elsdaig, 2019)         | Healthcare information system | x   | x  | x   | x  | x  | x  | -                                                                                                        |
| (Aparicio et al., 2019) | Gamification                  | x   | x  | x   | x  | x  | x  | Gamification, enjoyment, challenge                                                                       |

**Notes:** SEQ – Service Quality, IQ – Information Quality, SYQ – System Quality, SU – System Use, US – User Satisfaction, NB – Net Benefits.

## APPENDIX B. MEASUREMENT MODEL

Using a seven-point Likert scale with 1 translating to “strongly disagree” and 7 translating to “strongly agree”, the variables are measured by asking consultants to rate their perception of GEN-AI.

| Construct           | Measurement Items                                                                                                                                                                                                                                                                                                                                                                                                                                          | Reference                      |
|---------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------|
| Training            | T1. <i>“My employer provides generative AI learning/training opportunities to meet the changing needs of the workplace (e.g., certifications).”</i><br>T2. <i>“Generative AI training and development are encouraged and rewarded in my company.”</i><br>T3. <i>“Overall, I am satisfied with the amount of generative AI training I receive on the job.”</i>                                                                                              | (Hanaysha, 2016b, p. 178)      |
| Self-Efficacy       | SE1. <i>“I can always manage to solve difficult generative AI problems if I try hard enough.”</i><br>SE2. <i>“I can remain calm when facing generative AI difficulties because I can rely on my abilities.”</i><br>SE3. <i>“Failing to do something on the computer makes me try harder.”</i><br>SE4. <i>“There are a few things that I cannot do with generative AI.”</i>                                                                                 | (Howard, 2014, p. 681)         |
| Ethics of AI        | E1. <i>“I think generative AI should ensure people’s security and privacy in terms of data.”</i><br>E2. <i>“I think using generative AI requires an understanding of the purpose for which it was developed and designed.”</i><br>E3. <i>“I think the development of generative AI should be in line with ethical values.”</i>                                                                                                                             | (Shih et al., 2021)            |
| Service Quality     | SEQ1. <i>“The generative AI application provides a proper level of online assistance and explanation (e.g., chatbot).”</i><br>SEQ2. <i>“The application staff provides high availability for consultation.”</i><br>SEQ3. <i>“The application staff responds in a cooperative manner to suggestions for future enhancements of the application.”</i><br>SEQ4. <i>“The application provides satisfactory support to users using the generative AI tool.”</i> | (Wang et al., 2007, p. 1805)   |
| Information Quality | IQ1. <i>“The generative AI tool provides information that is relevant to my job.”</i><br>IQ2. <i>“The generative AI tool provides sufficient information.”</i><br>IQ3. <i>“The generative AI tool provides information that is easy to understand.”</i><br>IQ4. <i>“The generative AI tool provides up-to-date information.”</i>                                                                                                                           | (Wang et al., 2007, p. 1804)   |
| System Quality      | SYQ1. <i>“Generative AI is easy to use.”</i><br>SYQ2. <i>“I find it easy to get generative AI to do what I want it to do.”</i><br>SYQ3. <i>“It is easy for me to become skillful at using generative AI.”</i>                                                                                                                                                                                                                                              | (Nunes & Javier, 2014, p. 9)   |
| User Satisfaction   | US1. <i>“The generative AI tool is efficient.”</i><br>US2. <i>“The generative AI tool is effective.”</i><br>US3. <i>“Overall, I am satisfied with the generative AI tool.”</i>                                                                                                                                                                                                                                                                             | (Aparicio et al., 2019, p. 50) |
| System Use          | SU1. <i>“The frequency of use with generative AI is high.”</i><br>SU2. <i>“The generative AI usage is voluntary.”</i><br>SU3. <i>“I depend upon generative AI.”</i>                                                                                                                                                                                                                                                                                        | (Wang et al., 2007, p. 1805)   |
| Net Benefits        | NB1. <i>“Generative AI helps me to improve my job performance.”</i><br>NB2. <i>“Generative AI helps me to think through problems.”</i>                                                                                                                                                                                                                                                                                                                     | (Wang et al., 2007, p. 1805)   |
|                     | NB3. <i>“Generative AI helps me to acquire new knowledge and innovative ideas.”</i><br>NB4. <i>“Generative AI enables me to accomplish tasks more efficiently.”</i><br>NB5. <i>“Generative AI improves the quality of my work life.”</i>                                                                                                                                                                                                                   | (Wu & Wang, 2006, p. 738)      |

## APPENDIX C. ETHICS COMMITTEE APPROVAL



This is to certify that

Project No.: **INFSYS2023-10-179195**

Project Title: **The effect of generative AI on the productivity of consultants**

Principal Researcher: **Marijke Brommann**

according to the regulations of the Ethics Committee of NOVA IMS and MagIC Research Center this project was considered to meet the requirements of the NOVA IMS Internal Review Board, being considered **APPROVED** on 10/17/2023.

It is the Principal Researcher's responsibility to ensure that all researchers and stakeholders associated with this project are aware of the conditions of approval and which documents have been approved.

The Principal Researcher is required to notify the Ethics Committee, via amendment or progress report, of

- Any significant change to the project and the reason for that change;
- Any unforeseen events or unexpected developments that merit notification;
- The inability of the Principal Researcher to continue in that role or any other change in research personnel involved in the project.

Lisbon, 10/17/2023

NOVA IMS Ethics Committee  
ethicscommittee@novaims.unl.pt

## APPENDIX D. EXTRACT SURVEY DESIGN

### Training

\*What is your level of agreement regarding the following statements? [1 - Strongly Disagree; 2 - Disagree, 3 - Somewhat Disagree, 4 - Agree Nor Disagree, 5 - Somewhat Agree, 6 - Agree, 7 - Strongly Agree]

|                                                                                                                                        | 1                     | 2                     | 3                     | 4                     | 5                     | 6                     | 7                     |
|----------------------------------------------------------------------------------------------------------------------------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|
| My employer provides generative AI learning/training opportunities to meet the changing needs of the workplace (e.g., certifications). | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| Generative AI training and development are encouraged and rewarded in my company.                                                      | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| Overall, I am satisfied with the amount of generative AI training I receive on the job.                                                | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |

### Self-Efficacy

|                                                                                              | 1                     | 2                     | 3                     | 4                     | 5                     | 6                     | 7                     |
|----------------------------------------------------------------------------------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|
| I can always manage to solve difficult generative AI problems if I try hard enough.          | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| I can remain calm when facing generative AI difficulties because I can rely on my abilities. | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| Failing to do something on the computer makes me try harder.                                 | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| There are a few things that I cannot do with generative AI.                                  | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |

### Ethics of AI

|                                                                                                               | 1                     | 2                     | 3                     | 4                     | 5                     | 6                     | 7                     |
|---------------------------------------------------------------------------------------------------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|
| I think generative AI should ensure people's security and privacy in terms of data.                           | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| I think using generative AI requires an understanding of the purpose for which it was developed and designed. | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| I think the development of generative AI should be in line with ethical values.                               | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |

**Service Quality**

\*What is your level of agreement regarding the following statements? [1 - Strongly Disagree; 2 - Disagree, 3 - Somewhat Disagree, 4 - Agree Nor Disagree, 5 - Somewhat Agree, 6 - Agree, 7 - Strongly Agree]

|                                                                                                                       | 1                     | 2                     | 3                     | 4                     | 5                     | 6                     | 7                     |
|-----------------------------------------------------------------------------------------------------------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|
| The generative AI application provides a proper level of online assistance and explanation (e.g., chatbot).           | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| The application staff provides high availability for consultation.                                                    | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| The application staff responds in a cooperative manner to your suggestion for future enhancements of the application. | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| The application provides satisfactory support to users using the generative AI tool.                                  | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |

**Information Quality**

|                                                                         | 1                     | 2                     | 3                     | 4                     | 5                     | 6                     | 7                     |
|-------------------------------------------------------------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|
| The generative AI tool provides information that is relevant to my job. | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| The generative AI tool provides sufficient information.                 | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| The generative AI tool provides information that is easy to understand. | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| The generative AI tool provides up-to-date information.                 | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |

**System Quality**

|                                                                 | 1                     | 2                     | 3                     | 4                     | 5                     | 6                     | 7                     |
|-----------------------------------------------------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|
| Generative AI is easy to use.                                   | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| I find it easy to get generative AI to do what I want it to do. | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| It is easy for me to become skillful at using generative AI.    | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |

**User Satisfaction**

\*What is your level of agreement regarding the following statements? [1 - Strongly Disagree; 2 - Disagree, 3 - Somewhat Disagree, 4 - Agree Nor Disagree, 5 - Somewhat Agree, 6 - Agree, 7 - Strongly Agree]

|                                                      | 1                     | 2                     | 3                     | 4                     | 5                     | 6                     | 7                     |
|------------------------------------------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|
| The generative AI tool is efficient.                 | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| The generative AI tool is effective.                 | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| Overall, I am satisfied with the generative AI tool. | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |

|                                                                       |                       |                       |                       |                       |                       |                       |                       |
|-----------------------------------------------------------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|
| System Use                                                            | 1                     | 2                     | 3                     | 4                     | 5                     | 6                     | 7                     |
| The frequency of use with generative AI is high.                      | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| The generative AI usage is voluntary.                                 | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| I depend upon generative AI.                                          | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| Net Benefits                                                          | 1                     | 2                     | 3                     | 4                     | 5                     | 6                     | 7                     |
| Generative AI helps me to improve my job performance.                 | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| Generative AI helps me to think through problems.                     | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| Generative AI helps me to acquire new knowledge and innovative ideas. | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| Generative AI enables me to accomplish tasks more efficiently.        | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| Generative AI improves the quality of my work life.                   | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |

## APPENDIX E. MEASUREMENT ITEM CROSS LOADINGS

|             | A        | E            | IQ           | NB           | SE           | SEQ          | SYQ          | T            | SU       | US           |
|-------------|----------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|----------|--------------|
| <b>A</b>    | <b>1</b> | -0.072       | -0.102       | -0.149       | -0.148       | -0.075       | -0.198       | -0.051       | -0.091   | -0.24        |
| <b>E1</b>   | -0.076   | <b>0.867</b> | 0.356        | 0.309        | 0.289        | 0.183        | 0.33         | 0.132        | 0.167    | 0.468        |
| <b>E3</b>   | -0.055   | <b>0.918</b> | 0.418        | 0.352        | 0.265        | 0.251        | 0.431        | 0.151        | 0.209    | 0.44         |
| <b>IQ1</b>  | -0.161   | 0.365        | <b>0.874</b> | 0.592        | 0.457        | 0.427        | 0.58         | 0.328        | 0.375    | 0.609        |
| <b>IQ2</b>  | -0.053   | 0.313        | <b>0.859</b> | 0.522        | 0.396        | 0.535        | 0.589        | 0.299        | 0.365    | 0.574        |
| <b>IQ3</b>  | -0.043   | 0.439        | <b>0.825</b> | 0.534        | 0.395        | 0.425        | 0.606        | 0.267        | 0.354    | 0.568        |
| <b>NB1</b>  | -0.122   | 0.359        | 0.598        | <b>0.848</b> | 0.368        | 0.367        | 0.511        | 0.356        | 0.524    | 0.571        |
| <b>NB2</b>  | -0.031   | 0.301        | 0.524        | <b>0.717</b> | 0.326        | 0.392        | 0.469        | 0.253        | 0.458    | 0.499        |
| <b>NB3</b>  | -0.177   | 0.207        | 0.485        | <b>0.79</b>  | 0.407        | 0.417        | 0.503        | 0.264        | 0.49     | 0.448        |
| <b>NB4</b>  | -0.127   | 0.365        | 0.577        | <b>0.87</b>  | 0.357        | 0.383        | 0.54         | 0.246        | 0.523    | 0.678        |
| <b>NB5</b>  | -0.15    | 0.256        | 0.419        | <b>0.831</b> | 0.386        | 0.43         | 0.425        | 0.308        | 0.47     | 0.481        |
| <b>SE1</b>  | -0.147   | 0.225        | 0.424        | 0.416        | <b>0.92</b>  | 0.447        | 0.422        | 0.428        | 0.427    | 0.4          |
| <b>SE2</b>  | -0.114   | 0.34         | 0.454        | 0.389        | <b>0.863</b> | 0.323        | 0.426        | 0.35         | 0.312    | 0.425        |
| <b>SEQ1</b> | -0.083   | 0.206        | 0.42         | 0.377        | 0.357        | <b>0.721</b> | 0.401        | 0.347        | 0.284    | 0.392        |
| <b>SEQ2</b> | -0.021   | 0.143        | 0.399        | 0.349        | 0.372        | <b>0.817</b> | 0.33         | 0.341        | 0.355    | 0.31         |
| <b>SEQ3</b> | -0.016   | 0.226        | 0.498        | 0.434        | 0.326        | <b>0.844</b> | 0.423        | 0.412        | 0.348    | 0.428        |
| <b>SEQ4</b> | -0.122   | 0.203        | 0.395        | 0.376        | 0.345        | <b>0.792</b> | 0.379        | 0.326        | 0.324    | 0.381        |
| <b>SU1</b>  | -0.091   | 0.213        | 0.428        | 0.608        | 0.421        | 0.413        | 0.449        | 0.344        | <b>1</b> | 0.49         |
| <b>SYQ1</b> | -0.163   | 0.387        | 0.599        | 0.509        | 0.371        | 0.393        | <b>0.826</b> | 0.167        | 0.312    | 0.583        |
| <b>SYQ2</b> | -0.188   | 0.335        | 0.505        | 0.46         | 0.387        | 0.378        | <b>0.817</b> | 0.19         | 0.378    | 0.518        |
| <b>SYQ3</b> | -0.149   | 0.364        | 0.635        | 0.549        | 0.431        | 0.445        | <b>0.872</b> | 0.201        | 0.436    | 0.536        |
| <b>T1</b>   | -0.073   | 0.14         | 0.301        | 0.349        | 0.353        | 0.429        | 0.161        | <b>0.885</b> | 0.352    | 0.274        |
| <b>T2</b>   | 0.009    | 0.212        | 0.355        | 0.341        | 0.338        | 0.334        | 0.192        | <b>0.866</b> | 0.28     | 0.241        |
| <b>T3</b>   | -0.062   | 0.085        | 0.273        | 0.247        | 0.446        | 0.413        | 0.222        | <b>0.874</b> | 0.276    | 0.212        |
| <b>US1</b>  | -0.205   | 0.497        | 0.619        | 0.63         | 0.447        | 0.44         | 0.554        | 0.291        | 0.481    | <b>0.896</b> |
| <b>US2</b>  | -0.221   | 0.448        | 0.555        | 0.554        | 0.404        | 0.413        | 0.549        | 0.194        | 0.37     | <b>0.878</b> |
| <b>US3</b>  | -0.211   | 0.394        | 0.636        | 0.582        | 0.365        | 0.414        | 0.619        | 0.235        | 0.443    | <b>0.877</b> |

