



Acceptance of Artificial Intelligence in the Workplace:

The impact of self-affirmation and the mediating role of emotions, trust, and insecurity at work

Maria Inês Galrote da Mata Nunes

Dissertation written under the supervision of Professor

Cristina Mendonça

Dissertation submitted in partial fulfilment of requirements for the MSc in Business, at the Universidade Católica Portuguesa, 5 January 2026.

Abstract

Title: Acceptance of Artificial Intelligence in the Workplace: The impact of self-affirmation and the mediating role of emotions, trust, and insecurity at work

Author: Maria Inês Nunes

The rapid growth of Artificial Intelligence (AI) in the workplace has raised questions about the factors that influence its acceptance by workers. This dissertation investigated whether self-affirmation, a psychological intervention aimed at reinforcing personal worth, could increase acceptance of AI at work, as well as the mediating role of positive emotions, perceived job insecurity and trust in AI. The study used a between-subjects experimental design, with a convenience sample consisting of 200 participants with some type of professional experience. Participants were randomly assigned to a self-affirmation or control condition and exposed to a hypothetical scenario involving the implementation of an AI system. The variables were measured using different scales and analyzed using PROCESS (model 4), allowing direct and indirect effects to be tested. The results showed that self-affirmation significantly increased positive emotions, but did not influence job insecurity, trust in AI or acceptance of the technology. Among the mediators, only trust significantly influenced AI acceptance, while positive emotions and job insecurity did not show significant effects. No indirect effects of self-affirmation through mediators were identified. These findings suggest that while self-affirmation has a limited impact, emotional responses and trust play a central role in how workers view AI. Future research should confirm the stability of these effects and verify whether they hold up in real-world organizational contexts.

Keywords: Self-affirmation; Artificial Intelligence; Emotions; Job insecurity; Trust in AI; Technology adoption; Acceptance of AI at work

Sumário

Título: Aceitação da Inteligência Artificial no Local de Trabalho: O impacto da autoafirmação e o papel mediador das emoções, da confiança e da insegurança no trabalho

Autor: Maria Inês Nunes

O rápido crescimento da Inteligência Artificial (IA) no local de trabalho levantou questões sobre os fatores que influenciam a sua aceitação por parte dos trabalhadores. Esta dissertação investigou se a autoafirmação, uma intervenção psicológica destinada a reforçar o valor pessoal, poderia aumentar a aceitação da IA no trabalho, bem como o papel mediador das emoções positivas, da insegurança no trabalho percebida e da confiança na IA. O estudo utilizou um desenho experimental entre sujeitos, com uma amostra de conveniência constituída por 200 participantes com algum tipo de experiência profissional. Os participantes foram aleatoriamente alocados a uma condição de autoafirmação ou de controlo e expostos a um cenário hipotético que envolvia a implementação de um sistema de IA. As variáveis foram medidas através de diferentes escalas e analisadas através do PROCESS (modelo 4), permitindo testar os efeitos diretos e indiretos. Os resultados mostraram que a autoafirmação aumentou significativamente as emoções positivas, mas não influenciou a insegurança no trabalho, a confiança na IA ou a aceitação da tecnologia. Entre os mediadores, apenas a confiança influenciou significativamente a aceitação da IA, enquanto as emoções positivas e a insegurança no trabalho não apresentaram efeitos significativos. Não foram identificados efeitos indiretos da autoafirmação através dos mediadores. Estes resultados sugerem que, embora a autoafirmação tenha um impacto limitado, as respostas emocionais e a confiança desempenham um papel central na forma como os trabalhadores veem a IA. Pesquisas futuras deverão confirmar a estabilidade destes efeitos e verificar se se mantêm em contextos organizacionais reais.

Palavras-chave: Autoafirmação; Inteligência Artificial; Emoções; Insegurança Laboral; Confiança na IA; Adoção de Tecnologia; Aceitação da IA no Trabalho

List of Figures

Figure 1 - Conceptual Model	13
--	----

List of Tables

Table 1 - Scale's reliability test results	21
Table 2 - Means and Standard Deviations (SD) of all the variables used	21
Table 3 - Sample demographic characteristics	40
Table 4- Mean age and SD recorded	40
Table 5 - Sample sector distribution	40
Table 6 - Sample nationality distribution	41
Table 7 - Bivariate correlations	48

Glossary

α	Cronbach's index of reliability
β	Standardized coefficient
b	Unstandardized coefficient
&	And
AI	Artificial Intelligence
ANOVA	Analysis of Variance
CI	Confidence Interval
F	F distribution, fishers F ratio
H	Hypothesis
M	Mean
N	Total number of cases
p	p-value
r	Pearson correlation coefficient
R^2	Multiple correlation squared
RQ	Research Question
SD	Standard Deviation
SE	Standard Error
t	T-test
TAM	Technology Acceptance Model
UTAUT	Unified Theory of Acceptance and Use of Technology

Table of Contents

Abstract	i
Sumário.....	ii
List of Figures	iii
List of Tables.....	iii
Glossary	iv
1. Introduction	1
1.1. Relevance of the topic & problem statement.....	2
1.2. Academic and practical relevance	2
1.3. Structure.....	3
2. Literature review	4
2.1. Artificial Intelligence.....	4
2.1.1. AI in organizations and implications for workers	5
2.2. Acceptance of AI.....	6
2.2.1. Theories of technology acceptance.....	6
2.2.2. Acceptance of AI at work.....	7
2.3. Psychological factors influencing the acceptance of AI.....	8
2.3.1. Emotions.....	8
2.3.2. Perception of job insecurity	9
2.3.3. Trust in AI	10
2.4. Self-affirmation theory	10
2.5. Self-affirmation as a mechanism to promote acceptance of AI.....	12
2.6. Conceptual model	13
3. Methodology.....	14
3.1. Research design	14
3.2. Participants	14
3.3. Procedure.....	15
3.4. Variable measurement	16

3.4.1.	Independent variables	16
3.4.2.	Dependent variables	16
3.4.3.	Mediating variables	17
3.4.4.	Control variables.....	19
4.	Results.....	20
4.1.	Scale reliability	20
4.2.	Data preparation	21
4.3.	Descriptive statistics	21
4.4.	Bivariate correlations.....	22
4.4.1.	Relationships between primary variables and control/demographic variables	22
4.5.	Hypothesis testing.....	22
4.5.1.	Effect of experimental condition on mediating variables.....	23
4.5.2.	Effects of mediators on AI acceptance	23
4.5.3.	Direct and Total Effects of Self-Affirmation on AI Acceptance.....	24
4.5.4.	Mediation analyses	24
5.	Discussion	26
5.1.	Summary of results.....	26
5.2.	Connection to the existing literature.....	27
5.3.	Implications	28
5.4.	Limitations and future studies	29
5.5.	Conclusion	31
6.	References.....	32
7.	Appendix.....	40

1. Introduction

Recent technological advances, more specifically Artificial Intelligence (AI), are one of the main reasons for the significant transformation of various industry sectors. These changes are already transforming the present and having a major impact on the functioning of various organizations, leading to the creation of new business and work models (Philbeck & Davis, 2018).

Any change, regardless of the topic, brings with it many challenges and significant impacts. The adoption of AI in the workplace is often associated with fears of replacement, job insecurity, and negative impacts on the psychological well-being of workers (Frey & Osborne, 2017; Makarius et al., 2020). A study conducted by the Organization for Economic Cooperation and Development in 2019 predicted that 14% of jobs existing that year would disappear in the next 15 to 20 years and that 32% would change radically as a result of automation (OECD, 2019). Other reports point out that technologies such as automated learning, AI, and robotics already have a major impact on how people live and work and will therefore continue to influence the nature of work as well as the workplace itself (Manyika et al., 2017; World Economic Forum, 2020). Thus, although its use can generate efficiency gains, increase productivity, and contribute to reducing skills gaps (Maslej et al., 2025), AI is also perceived as a threat to fairness and job insecurity (Vander et al., 2013; De Witte et al., 2015). The acceptance and effectiveness of this technology depend not only on technical factors, but also on how workers perceive it, trust it, and are willing to implement it in their daily lives (Glikson & Woolley, 2020; Makarius et al., 2020). Thus, understanding workers' emotional and perceptual reactions to AI has become critical to ensuring an implementation that truly adds value to the company (Makarius et al., 2020).

Given this context, it is necessary to understand the processes that can reduce workers' defensiveness and increase acceptance of AI. Self-affirmation theory (Steele, 1998) is a promising tool. According to this theory, reflecting on fundamental personal values reinforces self-integrity, which leads to less defensiveness and greater openness to change and threatening scenarios (Sherman & Cohen, 2006; McQueen & Klein, 2006). Thus, the objective of this thesis is to assess whether self-affirmation can effectively intervene in the acceptance of AI in the workplace.

1.1.Relevance of the topic & problem statement

On the one hand, technological acceptance is a topic that has been studied repeatedly. One of the classic models is the Technology Acceptance Model (TAM), which presents perceived usefulness and ease of use as determining factors for the successful acceptance of new technologies (Davis, 1989). However, as studies have progressed, it has been found that AI acceptance also depends on emotional factors such as trust, technological anxiety, and perceived threat (Glikson & Woolley, 2020).

On the other hand, although not yet used in the context of AI acceptance, self-affirmation has been used in multiple other contexts, such as health (Sweeney & Moyer, 2015) and education (Escobar-Soler et al., 2024), showing that self-affirmation can reduce anxiety, increase engagement, and promote behavioral change. For example, individuals who performed self-affirmation tasks were more willing to adopt healthier behaviors, such as quitting smoking or improving their diet, after being exposed to threatening health messages (Sweeney & Moyer, 2015). This data shows how self-affirmation helps individuals cope with threats and defensive reactions, suggesting its potential as an effective psychological intervention to increase openness to AI in the workplace.

Given this context, this study aims to assess whether a self-affirmation intervention (Steele, 1998) can increase the acceptance of AI in the workplace by positively influencing three central psychological mechanisms: positive emotions, perception of job insecurity, and trust in AI. As a result, the following research questions were developed, which will be answered throughout this thesis:

RQ1: How do psychological factors, namely positive emotions, perceived job insecurity, and trust in AI, influence workers' AI acceptance?

RQ2: Does self-affirmation impact AI acceptance?

RQ3: Do emotions, perception of job insecurity, and trust help explain the impact of self-affirmation on AI acceptance?

1.2.Academic and practical relevance

From an academic point of view, this research contributes to filling a gap in the literature, since, although self-affirmation theory has proven effective in reducing defensive attitudes and promoting openness to change in other contexts (Steele, 1988; Sherman & Cohen, 2006), its

application to the context of AI acceptance at work is still an under-explored area. By testing this theory in the area of technology acceptance, this study aims to understand the impact of self-affirmation on how workers react to the implementation of AI systems. Its influence on psychological mechanisms, namely emotions, perceptions of job insecurity, and trust in AI, which can influence the willingness to accept these technologies (Fredrickson, 2001; Schepman & Rodway, 2020; Brougham & Haar, 2018; Makarius et al., 2020; De Visser et al., 2018; Glikson & Woolley, 2020), will be analyzed.

From a practical standpoint, this study may offer relevant guidance for managers and human resources professionals in people management and digital transformation processes, as self-affirmation can be a simple, low-cost intervention that can be applied in internal communication strategies to reduce defensive attitudes, increase confidence, and promote employee well-being (McQueen & Klein, 2006; Sweeney & Moyer, 2015).

1.3.Structure

This dissertation is organized into five chapters. The first, and current, chapter is the introduction, which presents the topic, contextualizes the research problem, presents the objectives and research questions, and justifies the academic and practical relevance of the study. The second chapter develops the literature review, addressing the main concepts and theories that underpin the conceptual model. The following topics are discussed: AI and its acceptance in the workplace, the psychological variables that influence its acceptance, namely, emotions, perception of job insecurity, and trust in AI, as well as the role of self-affirmation as a potential intervention to promote greater openness to AI adoption. This second chapter also presents the research hypotheses. The third chapter describes the study methodology, detailing the experimental design, the variables analyzed, and the data collection procedures. The fourth chapter presents and analyzes the results obtained, highlighting the relationships between the variables and testing the proposed hypotheses. Finally, the fifth chapter integrates the discussion and conclusions of the study, presenting the main theoretical and practical implications, as well as limitations and suggestions for future research.

2. Literature review

2.1. Artificial Intelligence

Artificial Intelligence can be defined as the ability of technological systems to perform tasks that would initially require human capabilities and intelligence, such as learning, reasoning, problem solving, and decision making (Russell & Norvig, 1995). In practical terms, AI refers to "a system's ability to interpret external data correctly, to learn from such data, and use those learnings to achieve specific goals and tasks through flexible adaptation" (Kaplan & Haenlein, 2019, p. 17).

Historically, it is difficult to date the beginning of AI research, but it is considered that the term AI was first officially mentioned in 1956 by John McCarthy, a mathematics professor at Dartmouth College at the time, during the Dartmouth Summer Research Project on Artificial Intelligence (Haenlein & Kaplan, 2019). In the conference proposal, the study of AI was justified as follows: "every aspect of learning or any other feature of intelligence can in principle be so precisely described that a machine can be made to simulate it" (as cited in Moor, 2006, p. 87). The first decades of study were accompanied by high expectations and grand plans, but at the same time by major technological limitations due to the limited processing capabilities of computers, the so-called AI Winters, leading to cycles of enthusiasm and retraction- (Haenlein & Kaplan, 2019).

However, since its emergence, AI has undergone significant evolution, moving from systems with fairly defined and fixed rules to systems with machine learning and deep learning algorithms, which allow the system to learn and improve its performance (Haenlein & Kaplan, 2019; Russell & Norvig, 1995).

This technology quickly became useful and necessary for various sectors, such as healthcare, finance, logistics, education, and marketing (Dwivedi et al., 2021). Currently, AI is seen as a strategic technology that generates competitiveness, optimizes processes, increases efficiency, and creates new forms of collaboration between humans and machines (Makarius et al., 2020). This rapid growth has led more traditional companies to adapt, keep up with these advances, and adopt AI technologies in their operations in order to reap their benefits and not fall behind in the evolution of the sector (Bean, 2019).

Despite significant advances in AI systems, they remain dependent on several factors (Davenport & Ronanki, 2018; Makarius et al., 2020): 1) they depend on technical factors: often, their integration into pre-existing systems is not possible or is not easy and may take some time; 2) they depend on organizational factors, such as the adaptation of the company's organizational culture (Makarius et al., 2020). Finally, they depend on human factors, such as acceptance and trust (Davenport & Ronanki, 2018). Thus, understanding these factors is fundamental to analyzing the role of AI in organizations and its implications for workers.

2.1.1. AI in organizations and implications for workers

The implementation of AI in organizations has been considered one of the main reasons for digital transformation and job restructuring (Davenport & Ronanki, 2018; Makarius et al., 2020). These new tools enable the automation of repetitive tasks, the processing of large volumes of data (i.e., big data), and support for decision-making, thereby contributing to gains in efficiency, productivity, and innovation (Davenport & Ronanki, 2018; Makarius et al., 2020). Many companies use these systems to optimize processes, improve their services (customer and internal customer), and increase their competitiveness, while benefiting from systems that are capable of learning and adapting to different organizational contexts (Brynjolfsson & McAfee, 2017; Haenlein & Kaplan, 2019).

There are many challenges that come with this technological revolution. One of the major challenges is how this implementation affects worker well-being, often causing anxiety and resistance (Brougham & Haar, 2018; Frey & Osborne, 2017). These feelings are often associated with fear of replacement, loss of autonomy, and the perception of devaluation of skills and values (Brougham & Haar, 2018; Frey & Osborne, 2017). The introduction of intelligent systems has a major impact on organizations, particularly on power, trust, and control relationships (Glikson & Woolley, 2020).

AI is redefining the tasks and role of humans at work, eliminating and simplifying repetitive work, thus giving workers space and time for more cognitive, creative, and collaborative tasks (Wilson & Daugherty, 2018). Although these systems are creating new opportunities for development and learning, they require some adaptation, such as process restructuring, but it is also necessary to acquire new digital skills and adapt to enable better collaboration between

humans and machines (Brougham & Haar, 2018; Davenport, 2018). While workers recognize the benefits of these systems, they have fears associated with job insecurity, loss of autonomy, the need for retraining, and the impact on their employment, especially among workers who perceive AI as a direct substitute for their jobs (Brougham & Haar, 2017; Makarius et al., 2020).

The implementation of AI systems in organizations goes far beyond technological or strategic issues; it is a process that depends on social factors, especially on the acceptance and involvement of employees (Davenport & Ronanki, 2018). While job insecurity, unfairness, and loss of control are associated with more defensive attitudes and resistance to change, when employees perceive AI as a support rather than a replacement, they tend to show more confidence, openness, and willingness to accept new technologies, thus adding value to the company (Jussupow et al., 2021; Makarius et al., 2020; Shin, 2020).

Taking all these aspects into account, organizational communication must be clear and employee involvement must be present from the outset of projects in order to ensure effective implementation (Brougham & Haar, 2020).

2.2. Acceptance of AI

2.2.1. Theories of technology acceptance

Technology acceptance has been intensively studied in recent decades, particularly with regard to organizational behavior (Venkatesh et al., 2003). Thus, there are several theories that attempt to justify individuals' decisions to accept or reject new technologies. One of the classic and most influential models is the Technology Acceptance Model (TAM), developed by Davis (1989), which presents perceived usefulness and perceived ease of use as the main motivators for accepting a new technology. This model argues that the more users believe that the technology will be useful and easy to use, the greater their intention to use it will be (Davis, 1989).

A few years later, Venkatesh et al. (2003) developed the Unified Theory of Acceptance and Use of Technology (UTAUT), a model with a more in-depth explanation that uses the main principles of TAM and contributions from eight other models. UTAUT includes four key determinants of the intention to use new technologies: performance expectancy, effort

expectancy, social influence, and facilitating conditions (Venkatesh et al., 2003). This model is thus a little more complete and attempts to explain much of the behavior related to technology acceptance (Venkatesh et al., 2003). It demonstrates that acceptance depends not only on technological factors, but also on cognitive and social factors (Venkatesh et al., 2003).

There are also other theories such as Diffusion of Innovations Theory, which analyzes innovation characteristics such as relative advantage, compatibility, and complexity, and relates them to acceptance over time (Rogers, 2014).

All these theories continue to be applied in organizational contexts in an attempt to understand why some individuals accept, and others reject, the introduction of new technologies. Although these models are still relevant today, AI has characteristics that distinguish it from other types of technology, such as autonomy, unpredictability, and decision-making capacity, which require the models to be adapted (Glikson & Woolley, 2020).

2.2.2. Acceptance of AI at work

Given the differences between AI and more traditional technologies, namely AI's greater level of autonomy, clarity and less human control (Burrell, 2016; Lee & See, 2004; Siau & Wang, 2018), it is expected that the processes involved in its acceptance may also be different. In an organizational context, it will be an even greater challenge, taking into account that adapting to a new technology will likely be accompanied by structural changes in the nature of work and the professional identity of workers (Brougham & Harr, 2018; Makarius et al., 2020).

Several researchers have argued that beneficial and effective acceptance of AI depends not only on rational factors and technical characteristics, such as cost-benefit assessment, but also on emotional factors, such as enthusiasm, curiosity, fear, or threat (Glikson & Woolley, 2020; Schepman & Rodway, 2020). All these reactions influence how individuals perceive technology and, as a result, how they accept and integrate it into their work routines (Schepman & Rodway, 2020). The acceptance of AI is also influenced by the relationship between humans and machines. Unlike other technologies, AI will not only depend on the cognitive dimensions explored in TAM (Davis, 1989) and UTAUT (Venkatesh et al., 2003), but also on affective responses related to fairness, control, and personal value (De Visser et al., 2018; Makarius et al., 2020).

When the introduction of AI is perceived as a threat, whether to job security, autonomy, or professional status, workers tend to show greater resistance, mistrust, and defensive attitudes (Brougham & Haar, 2018; Vander Elst et al., 2013). In contrast, when it is seen as a support, it can promote positive emotions, motivation, and confidence in workers, thus facilitating its acceptance (Makarius et al., 2020).

Therefore, to understand the AI acceptance in an organizational context, it is necessary to focus on emotional, perceptual, and relational dimensions, which assess emotions associated with technology, perceptions of job insecurity, and trust in AI systems (Glikson & Woolley, 2020).

This thesis does precisely that, advocating a psychological and organizational perspective on AI acceptance, focusing on positive emotions, perceived job insecurity, and trust in AI. These variables will be explored in the following subchapters.

2.3. Psychological factors influencing the acceptance of AI

As seen in the previous subchapter, the acceptance of AI in the workplace does not depend solely on technical factors, but also on psychological factors that may be related to how workers perceive and feel about the introduction of this technology (Glikson & Woolley, 2020; Makarius et al., 2020). This thesis highlights three factors that are decisive in the acceptance of AI: positive emotions (Schepman & Rodway, 2020), the perception of job insecurity (De Witte, 2000), and trust in AI (Hoff & Bashir, 2015).

2.3.1. Emotions

Emotions play an important role in how each individual evaluates, interprets, and reacts to the introduction of new technologies. Given this scenario, it is important to mention Fredrickson's (2004) "broaden-and-build" theory of positive emotions. This theory argues that more positive emotional states increase attention, creativity, and adaptation, thus contributing to greater openness to change (Fredrickson, 2004).

In the organizational context, positive emotions have been directly associated with greater motivation and involvement in what we call digital transformation (Seo et al., 2004). Negative

emotions, such as anxiety and fear, have been associated with resistance and defensive attitudes toward the same topic (Saks et al., 2007).

Therefore, the introduction of AI in the workplace can generate any of these emotional effects. Recent studies indicate that workers who experience positive emotions, such as enthusiasm and pride, tend to perceive AI as an opportunity for professional growth and development, thus showing a greater willingness to accept it (Schepman & Rodway, 2020).

So, my first hypothesis is:

***H1:** Higher levels of positive emotions are associated with greater acceptance of AI in the workplace.*

2.3.2. Perception of job insecurity

The perception of job insecurity refers to the feeling that workers have regarding the continuity of their employment and its future conditions (De Witte, 2015). As many workers perceive AI as a threat to their job security and the importance of their roles, the introduction of automated and AI systems increases this feeling (Brougham & Haar, 2018; Frey & Osborne, 2017).

One of the aspects that most influences job insecurity is the fear of replacement. This is characterized by the fear that AI will be so well developed that it will perform human tasks more efficiently and thus reduce the need for human labor (Makarius et al., 2020).

Job insecurity is directly related to higher levels of stress and greater resistance to change, which can hinder and make it more difficult to accept new technologies (Cheng & Chan, 2008). Similarly, greater perceptions of job insecurity have led workers to have more negative attitudes toward innovation (Vander Elst et al., 2013).

However, it all depends on how organizational communication is structured. If the introduction of AI and its objectives are communicated clearly, transparently, and always as a means of supporting human tasks, the feeling of threat tends to diminish, and acceptance tends to increase (Makarius et al., 2020).

Therefore, my second hypothesis is:

***H2:** Higher levels of perceived job insecurity are associated with lower acceptance of AI in the workplace.*

2.3.3. Trust in AI

Trust is one of the variables that most influences behavior related to technology acceptance (Mayer et al., 1995; McKnight et al., 2011). In the case of AI technologies, trust refers to believing that these technologies are fair, safe, and benevolent in their decisions (Glikson & Woolley, 2020).

Trust plays a very important psychological role in reducing uncertainty, changing how comfortable (or uncomfortable) workers feel about relying on automated systems (Hoff & Bashir, 2015). When AI is perceived as clear and predictable, it increases feelings of security and, consequently, acceptance of AI (De Visser et al., 2018).

On the contrary, a lack of or low confidence in this technology is associated with higher levels of resistance and rejection, even when its results are as good as or better than those of humans (Mayer et al., 1995). It is therefore necessary to promote an organizational environment of trust in intelligent technology systems in order to ensure their successful integration (Glikson & Woolley, 2020).

Therefore, my third hypothesis is:

***H3:** Higher levels of trust in AI are associated with greater acceptance of the AI in the workplace.*

2.4. Self-affirmation theory

Steele (1988) developed the theory of self-affirmation, proposing that individuals have a biological need to maintain a positive and integrated self-image, seeing themselves as moral, competent, and consistent people. When this perception is threatened, for example through criticism or mistakes, defensive responses are triggered in order to reestablish personal value, even if this often involves ignoring relevant information (Sherman & Cohen, 2006).

Self-affirmation is a psychological process in which individuals reflect on personally important values such as honesty, friendship, or competence (Sherman & Cohen, 2006; Steele, 1988).

This reflection reinforces one's sense of personal worth and supports self-regulation (Sherman & Hartson, 2011). According to self-affirmation theory, restoring this sense of integrity helps individuals deal with threatening situations in a less defensive and more adaptive way (Sherman & Cohen, 2006).

Research also propose that the process of self-affirmation can act through other complementary mechanisms. On the one hand, it can promote emotional regulation, providing tools for individuals to recover a positive emotional state and reduce stress (Creswell et al., 2005; Koole et al., 1999). On the other hand, it can allow individuals to assess threats differently, often making them less personal and more solution-oriented (Sherman & Hartson, 2011). It also contributes to the process of perceiving autonomy and control, which are important factors in self-integrity (Sherman & Cohen, 2006; Steele, 1988).

Self-affirmation theory has been tested in various contexts, demonstrating its effectiveness in reducing defensiveness, improving well-being, and promoting a more open outlook on change (Sherman & Cohen, 2006). In recent decades, over a hundred experimental studies have been conducted in the areas of health, education, and social behavior (Cohen & Sherman, 2014).

In the health field, a meta-analysis by Sweeney et al. (2015) used 16 studies that demonstrated that self-affirmation increases the positive acceptance of messages about health risks. The authors concluded that individuals who practice self-affirmation tend to react less defensively to threatening information and demonstrate a greater intention to adopt healthy behaviors, such as quitting smoking and improving their diet. This information indicates that self-affirmation reduces defensive information processing and increases receptiveness to threatening messages.

In the context of education, a meta-analysis conducted by Escobar-Soler et al. (2023) synthesized results from over 100 independent studies that confirmed that self-affirmation is effective in reducing stress and increasing academic performance, especially in groups more exposed to threats to identity, such as minorities or students in demanding courses.

In addition to the evidence presented in educational and health contexts, there is also research that demonstrates the role of self-affirmation in work environments. Studies show that self-affirmation can reduce defensiveness, increase openness to feedback and promote more ethical

behaviors in the organizational context (Sherman & Cohen, 2006; Legault et al., 2012). These results suggest that, by reinforcing self-integrity, self-affirmation helps workers to deal less defensively with potentially challenging information, and may facilitate healthier adaptation to challenges, feedback and changes in the work environment.

So, given that self-affirmation works in education and healthcare, there may be other areas, such as AI acceptance, where it can be effective.

2.5. Self-affirmation as a mechanism to promote acceptance of AI

As discussed in previous subchapters, the implementation of AI in the workplace represents not only a technological change, but also a threat to the role of workers (Frey & Osborne, 2017; Brougham & Haar, 2018). The negative perceptions associated with this implementation can generate defensive actions on the part of workers, leading them to resist change (Makarius et al., 2020; Brougham & Haar, 2018).

The theory of self-affirmation argues that when the self is threatened, people feel the need to seek to restore their personal integrity in some way, reinforcing positive aspects of their identity (Steele, 1988). This reflection helps to reduce defensiveness and perceive the threat in a less personal way (Sherman & Cohen, 2006). Based on this reasoning, in the workplace context, self-affirmation may also help workers cope with negative emotions and the fear of replacement associated with AI implementation. Thus:

H4: Self-affirmation will improve AI acceptance.

The literature review contained in this chapter suggests that implementation may influence behavior related to technology acceptance through three main psychological mechanisms:

1. By increasing positive emoticons, as self-affirmation creates a beneficial emotional state that facilitates the process of accepting new information and reduces resistance to change (Creswell et al., 2005; Fredrickson, 2001).
2. By restoring feelings of value and control, as self-affirmation reduces perceptions of threat and insecurity (Brockner & Wiesenfeld, 1996; Wiesenfeld et al., 1999). Consequently, this mechanism can help alleviate employees' fears that AI may make their roles less necessary.
3. By increasing trust, as self-affirmation makes people more open and willing to evaluate changes more objectively, reducing defensiveness and increasing receptiveness to new

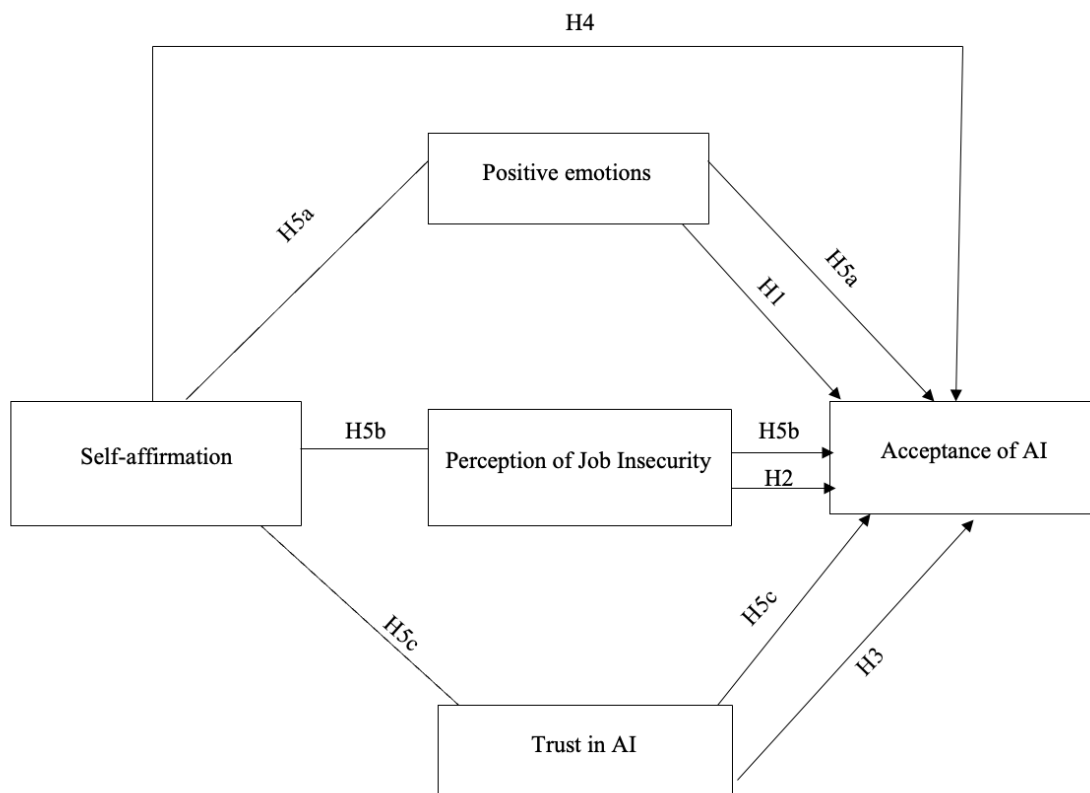
information (Sherman & Cohen, 2006) and greater objectivity in evaluation may, in turn, support the development of trust in AI (Glikson & Woolley, 2020).

These processes may allow workers to perceive AI as a means of assistance rather than a substitute, encouraging its acceptance (Brougham & Haar, 2018; Glikson & Woolley, 2020). Thus, I propose a final hypothesis:

H5: *Self-affirmation will impact AI acceptance through (H5a) more positive emotions, (H5b) lower perceived job insecurity, and (H5c) greater trust in AI.*

2.6. Conceptual model

Figure 1 - Conceptual Model



3. Methodology

3.1. Research design

The main objective of this study is to test the impact of self-affirmation on the acceptance of AI in the workplace, investigating not only its direct effect but also its indirect effects through the mediation of positive emotions, perception of job insecurity, and trust in AI. The aim is to understand the extent to which reflection on important personal values can influence how workers accept and react to the introduction of AI systems. Thus, an experimental approach was implemented with a single between-subjects factor with two levels (condition: self-affirmation vs. control). Thus, in this study, each participant is exposed to only one of the conditions in order to avoid possible knowledge transfer from one condition to another or bias in the conditions. Data collection was carried out using an online questionnaire created through Qualtrics, an online survey platform. Through the platform, it was possible to randomly assign participants to each of the conditions in order to reduce possible biases associated with demographic characteristics such as age, gender, and education level.

3.2. Participants

The study questionnaire was distributed through various online channels, namely social networks, personal contacts, and questionnaire-sharing platforms. Participation was voluntary, anonymous, and without financial compensation.

A total of 264 questionnaires were initiated, of which 238 had all responses completed. As can be seen in the data, the study dropout rate was not very high (only 10%), but this dropout occurred mainly at the beginning of the questionnaire (in the first 20%), which corresponds to the written and reflective question. After an initial analysis of the defined validation criteria, excluding responses from people who had never had professional experience, since the sample under study consists of people with work experience, 38 responses were excluded, thus 200 responses were considered valid for the study. Of the relevant sample, 67.5% identified as women, 31% as men, and the remainder were distributed among other gender categories or preferred not to answer. The average age of the participants was 28 years. Regarding the level of education, most had a master's degree (40.2%) or a bachelor's degree (39.2%). The participants in this study were individuals of 36 different nationalities, with India (18.5%), Portugal (18%), and the United States (12.7%) being the most represented.

More demographic information can be found in Appendix A.

3.3.Procedure

Participants began by reading an introductory page that briefly presented the study's objective, an estimate of the time required (five to eight minutes), and a statement of anonymity and voluntary participation. On this page, participants were also informed that, by proceeding to the next page, they would be agreeing to participate in this research.

Next, participants answered an exclusion question, which was intended to confirm whether the participant had current or previous work experience, one of the necessary conditions for the responses to be considered valid, since the study assesses the introduction of AI in the workplace.

Participants deemed eligible were randomly assigned to one of two conditions, self-affirmation or control, using the Qualtrics automated randomization system. After being assigned a condition, participants moved on to the first block, which consisted of a writing task that was different for each condition (self-affirmation or description of the commute to work). After the writing task, participants briefly assessed their current emotional state, indicating the intensity with which they were feeling various positive and negative emotions. Next, participants were presented with a hypothetical scenario of implementing an AI system in their workplaces. In this scenario, the worker receives an email from the management team announcing that a new AI system will soon be implemented to assist the worker in various tasks, while acknowledging that it may alter how some tasks are performed and, in some cases, render certain tasks irrelevant or redundant.

In the following section, participants answered a set of questions (i.e., job insecurity, trust in AI, acceptance of AI) to assess their different perceptions and reactions to the described scenario. The final section focused on demographic questions (i.e., age, gender, nationality, and education level).

The questionnaire concluded with a thank-you message to the participants for their collaboration, as well as my contact information should they wish to obtain more information about the study.

The complete questionnaire can be found in Appendix B.

3.4. Variable measurement

3.4.1. Independent variables

Experimental condition (Self-affirmation vs. Control): The independent variable in this study is the experimental condition, manipulated through a writing task based on the self-affirmation theory developed by Steele (1988) and already used in other studies and areas (Sweeney & Moyer, 2015; Escobar-Soler et al., 2023).

In this study, participants were randomly assigned to one of two conditions:

1. In the self-affirmation condition, participants chose a personal value they consider important in the work context (e.g., integrity, responsibility, collaboration) and wrote a brief reflection (three to five lines) on why they considered this value important and describing a situation in which this value was relevant in their professional life. This process follows the classic self-affirmation structure developed by Steele (1988) and adapted to the context under study (instead of selecting values related to life in general, participants selected values specific to the work context).
2. In the active control condition, participants were asked to write (also three to five lines) about their commute to work, a neutral task without self-reflective content. This type of task has already been used as a control task in other self-affirmation studies, in which participants write about everyday topics that do not involve reflection on personal values (e.g., Harvey & Oswald, 2000). These tasks are a common procedure to ensure that the exercise does not activate self-affirmation processes. Thus, the task used in this study aligns with tasks already established in the literature.

For results evaluation, an additional dummy variable was created, corresponding to the experimental manipulation, in order to identify the condition assigned to each participant: self-affirmation or control.

3.4.2. Dependent variables

Acceptance of AI: The dependent variable in this study is the acceptance of AI by workers in the workplace. This variable demonstrates the participants' intention to use AI systems, learn about them, and recommend them in an organizational context.

In this study, AI acceptance was measured using four items that assess participants' intention to use, learn about, and recommend the AI system described in the scenario. These items were inspired by existing literature on AI technological acceptance, namely the Behavioral Intention subscale of the UTAUT (Venkatesh et al., 2003), and influenced by the conceptual literature on acceptance and trust in AI in organizational contexts (Glikson & Woolley, 2020).

In the UTAUT, behavioral intention is measured by items such as "I would find the system useful in my work." Following this formulation, three of the study's items were formulated taking this specific context into account, such as "If my organization made this system available, I would use it."

The fourth item was included based on recent literature on AI acceptance, which adds recommendation behaviors as relevant components of the adoption of intelligent systems in organizations (Glikson & Woolley, 2020).

It was decided not to use a pre-existing scale in its entirety, such as the TAM (Davis, 1989) and other UTAUT subscales, because they include multiple points that were not the central focus of this study (e.g., perceived usefulness, technology anxiety, enabling conditions). Thus, the items were formulated to directly capture the intention to adopt the product in the presented work context.

The responses were evaluated on a Likert scale from 1 ("Strongly Disagree") to 5 ("Strongly Agree").

3.4.3. Mediating variables

Emotions: Emotions were measured using an adaptation of the Positive and Negative Affect Scale (PANAS), developed by Watson et al. (1988). The original scale consists of 20 items (10 positive and 10 negative affects), which are frequently used to measure immediate affective states in experimental studies. This allows capturing participants' immediate emotional responses right after the self-affirmation manipulation.

For this study, only eight items from the original scale were selected: four positive affect items (enthusiastic, interested, active, attentive) and four negative affect items (nervous, worried, scared, irritated). The remaining PANAS items were excluded, namely, the positive affects

strong, animated, proud, alert, inspired, determined and the negative affects disturbed, frightened, hostile, worried, embarrassed, agitated and distressed. This selection was made considering two reasons: (1) relevance in scenarios of technological change, prioritizing emotions associated with activation and alertness; and (2) reducing the duration of the questionnaire in order to decrease the questionnaire dropout rate. Participants indicated the intensity with which they felt each emotion on a scale of 1 (“Not at all”) to 5 (“Extremely”). Logically, higher values reflect higher levels of the corresponding emotion.

Perception of Job Insecurity: The perception of job insecurity was measured using a four-item scale developed specifically for this study. The wording of these items was inspired by classic measures of job insecurity used in previous studies, namely De Witte (1999), whose original scale includes four items. Due to the specific theme of the study, the items were formulated to reflect the context in question, maintaining the conceptual logic of the original items, but applying it to the experimental setting.

Considering this approach, an example of this inspiration could be: "I feel insecure about my future in this job." from De Witte (1999) to "With the introduction of this system, I am afraid that my job will become less relevant."

Responses were evaluated on a Likert scale from 1 ("Strongly Disagree") to 5 ("Strongly Agree"). Higher scores indicate a greater perception of job insecurity regarding the implementation of AI.

Trust in AI: Trust in AI was measured using four items developed specifically for this study, inspired by the definition of trust in AI systems proposed by Glikson and Woolley (2020), which identifies three fundamental dimensions: credibility, fairness, and impartiality.

Based on these dimensions, the items formulated in this study were inspired by formulations found in the literature on trust in technologies. For example, while Glikson and Woolley (2020) argue that user trust is fundamental to their development in organizations, one of the items used in this study was: "I would trust the recommendations provided by this system."

The responses were evaluated on a Likert scale from 1 (“Strongly Disagree”) to 5 (“Strongly Agree”).

3.4.4. Control variables

Familiarity with AI: Familiarity with AI was included as a control variable, given that previous studies demonstrate that prior contact with digital technologies tends to reduce uncertainty, increase confidence, and facilitate the adoption of new systems (Flynn & Goldsmith, 1999; Venkatesh et al., 2003; Glikson & Woolley, 2020).

Since there is no specific standard scale for familiarity with AI systems at work that simultaneously captures frequency of use, subjective comfort, and training received, three independent items were included that reflect these dimensions recognized in the literature: (1) Frequency of use (“How often do you use AI tools in your work?”), measured on a scale of 1 (“Never”) to 5 (“Frequently”), aligned with the idea of technological exposure as a way to reduce uncertainty (Flynn & Goldsmith, 1999); (2) Perceived comfort (“How comfortable do you feel using AI-based technologies?”), measured on a scale of 1 (“Not comfortable at all”) to 5 (“Very comfortable”), reflecting the role of initial comfort and trust in building trust in systems (Venkatesh et al., 2003); (3) Prior training (“Have you received AI training at work?”), answered in a dichotomous format (“yes”/“no”), consistent with the notion that prior knowledge reduces perceived risk and increases trust in AI (Glikson & Woolley, 2020).

Thus, although the items do not result from a standard scale, their formulation directly follows the principles present in the literature on technological familiarity and trust in AI systems.

Although scales exist to assess familiarity with technology, such as the Computer Self-Efficacy Scale (Compeau & Higgins, 1995), the Technology Readiness Index (Parasuraman, 2000), or prior experience items included in the UTAUT (Venkatesh et al., 2003), these measures focus primarily on digital technologies in general (e.g., computers, applications, or electronic devices) and not specifically on AI systems in the workplace context. For this reason, I considered that a simple terminological adaptation (e.g., replacing "computer" with "AI") would not capture relevant aspects of AI familiarity, especially in the workplace. Therefore, it was decided to develop specific items for this study that were analyzed separately.

Demographics: Several demographic variables were collected to characterize the sample and control potential individual differences among participants. The variables included were:

- Occupational sector: participants selected the sector where they work or have worked from a predefined list consisting of construction, education, energy, finance, health, manufacturing, retail, services, technology, and other.

- Gender: response options included male, female, non-binary/third gender, and I prefer not to answer.
- Age: open-ended numerical response, allowing participants to indicate their age in years.
- Education level: measured through ordered options, including less than a high school diploma, high school diploma, incomplete higher education but no degree, bachelor's degree, master's degree, medical degree, professional diploma, and other.
- Nationality: selected from a comprehensive list of nationalities provided in the questionnaire.

Manipulation effectiveness check: In order to verify the effectiveness of the manipulation, a verification item was included, in which participants indicated, on a scale of 1 (“Strongly Disagree”) to 5 (“Strongly Agree”), the extent to which the exercise made them reflect on important aspects of their professional life. The item used was: “The exercise made me reflect on what I consider important in my work”.

4. Results

4.1. Scale reliability

Given that most scales used in this study required linguistic or contextual adaptations, I decided to evaluate the reliability of each scale. In this sense, Cronbach's α was calculated for each scale, with the results presented in Table 1.

All scales show adequate levels, revealing reliability values that fall within the recommended standards ($\alpha \geq .70$). Thus, the job insecurity, trust in AI, and AI acceptance scales, as well as the measures of positive and negative emotions, demonstrated sufficient reliability to be used in the analyses of this thesis.

Table 1 - *Scale's reliability test results*

Scale	Number of items	Cronbach's α
Job insecurity	4	.87
Trust in AI	4	.82
AI acceptance	4	.89
Positive emotions	4	.83
Negative emotions	4	.80

4.2. Data preparation

Additional variables were created by calculating the average of the respective items, related to job insecurity, trust in AI, acceptance of AI, positive and negative emotions, in order to represent this variable.

The variable “familiarity with AI” was composed of three independent indicators (frequency of use, comfort, and prior training). Given that they do not constitute a unidimensional scale, these indicators were not aggregated or subjected to reliability tests, being analyzed individually and used only descriptively.

4.3. Descriptive statistics

In terms of the distribution of participants across the two experimental conditions, the sample was roughly divided between the self-affirmation condition ($n = 92$ or 46%) and the control condition ($n = 108$ or 54%). Regarding training in AI, 32% ($n = 64$) of participants indicated having received prior training, while 68% ($n = 136$) reported not having received any type of training. Table 2 presents the descriptive values of the main variables used in this study.

Table 2 - *Means and Standard Deviations (SD) of all the variables used*

Scale	N	M	SD
Job insecurity	200	2.98	1.10
Trust in AI	200	3.33	0.87
AI acceptance	200	3.86	0.90
Positive emotions	200	2.94	0.93
Negative emotions	200	2.11	0.94
Frequency of AI use	200	3.43	1.20
Comfort with AI	200	3.81	0.99

4.4. Bivariate correlations

The bivariate correlations between all variables included in the study, namely the main variables (job insecurity, positive emotions, negative emotions, trust in AI, and AI acceptance), experimental condition, and demographics, can be observed in Table 7 in Appendix C.

4.4.1. Relationships between primary variables and control/demographic variables

Regarding demographic variables, age shows negative correlations with job insecurity ($r = -.15$; $p = .047$) and with negative emotions ($r = -.15$; $p = .037$), suggesting that older participants tend to feel less insecure about work and fewer negative emotions towards AI. Age is not significantly related to AI acceptance ($r = .06$; $p = .396$) or trust in AI ($r = -.07$; $p = .357$).

As for nationality, Portuguese participants show significantly lower levels of job insecurity ($r = -.24$; $p < .001$) and slightly higher levels of trust in AI ($r = .16$; $p = .031$), as well as a positive, though not significant, trend towards greater AI acceptance ($r = .13$; $p = .086$).

The master's degree variable shows a positive correlation with positive emotions ($r = .21$; $p = .002$), indicating that participants with a master's degree report, on average, more positive emotions regarding AI.

Regarding the sector of activity, there are significant correlations between some sector main variables, but, in general, of reduced importance. For example, working in the financial sector is slightly associated with greater trust in AI ($r = .14$; $p = .046$).

It is also important to highlight that, at the bivariate level, positive emotions show a positive and statistically significant correlation with AI acceptance ($r = 0.23$; $p = 0.001$). However, as will be analyzed in the hypothesis testing section, this relationship is no longer significant when considered in a multivariate model that includes other mediators and control variables, suggesting that the association observed in the bivariate analysis may be explained by other factors included in the model.

4.5. Hypothesis testing

To test the hypotheses, the PROCESS macro model 4 (Hayes, 2017) was used, which is frequently employed in experimental research to simultaneously verify direct, indirect, and total effects. The independent variable corresponded to the experimental condition (0 = control;

1 = self-affirmation), while the dependent variable was the acceptance of AI. The mediating variables included were: (1) perceived job insecurity, (2) positive emotions, and (3) trust in AI. This procedure allows for simultaneously testing the individual effects of each mediator and assessing the existence of mediation by a single test, instead of an indirect step-wise procedure that stops immediately when one of the steps produces a non-significant result (Hayes, 2017).

The analysis was conducted with 5,000 bootstrap samples and a 95% confidence interval.

More information is available in the PROCESS output in Appendix D.

4.5.1. Effect of experimental condition on mediating variables

Before evaluating the main hypotheses, it was verified whether the experimental manipulation had an influence on the mediators.

Regarding H5a, self-affirmation had a significant and positive effect on the positive emotions felt by the participants, $b = 0.43, p < .001$. This result is informative, insofar as it demonstrates that the intervention produced the expected effect on the proposed mediator, although, by itself, it is not sufficient to establish the existence of a mediating effect.

Next, regarding H5b, the effect of self-affirmation on job insecurity was not significant, $\beta = .19, p = .228$. This result indicates that the intervention did not alter the perception of job threat.

Finally, with respect to H5c, there was no significant effect of self-affirmation on trust in AI, $\beta = .08, p = .530$. That is, the intervention did not contribute to altering the participants' trust levels in AI.

Thus, of the three mediators, only positive emotions were significantly influenced by the experimental condition.

4.5.2. Effects of mediators on AI acceptance

In order to evaluate the hypotheses related to the association between the mediators and the acceptance of AI, the multiple regression model calculated by PROCESS was analyzed, which included the experimental condition and the three mediators. This model showed that it was statistically significant, explaining 29% of the variation in AI acceptance, $R^2 = .29; F(4,195) = 19.77; p < .001$.

Regarding positive emotions, a positive but not significant effect was observed, $\beta = .11, p < .085$. Since the p-value is greater than .05, the effect is not significant, so H1, which predicted

that higher levels of positive emotions would be associated with greater acceptance of AI in the workplace, is not supported.

Job insecurity shows a negative effect, as predicted, but it is not significant, $\beta = -.05, p < .346$. Although the effect is consistent with H2, which predicted that higher levels of perceived job insecurity would be associated with lower acceptance of AI in the workplace, is not supported. Trust in AI showed a significant effect, $\beta = -.52, p < .001$, suggesting that higher levels of trust are associated with greater acceptance of AI, keeping all other variables constant. Thus, H3, which predicted that higher levels of trust would be associated with greater acceptance of AI in the workplace, is corroborated, demonstrating that trust is the main determinant of acceptance in the model.

4.5.3. Direct and Total Effects of Self-Affirmation on AI Acceptance

The total effect of self-affirmation on AI acceptance was not significant, $\beta = -.11, p < .378$, suggesting that, overall, the self-affirmation intervention did not have a significant impact on AI acceptance.

As for the direct effect, that is, when mediators are included in the model, it remained non-significant, $\beta = -.19, p < .091$.

Thus, H4 is not supported; self-affirmation did not increase AI acceptance, either directly or in total.

4.5.4. Mediation analyses

H5 proposed that the effect of the experimental condition (self-affirmation vs. control) on AI acceptance would occur through three possible mediators: positive emotions (H5a), job insecurity (H5b), and trust in AI (H5c).

The mediation analysis examined the indirect effects of the experimental condition on AI acceptance through each mediator. An indirect effect is considered significant when the bootstrap confidence interval does not include zero (Hayes, 2017).

Hypothesis H5 proposed that the effect of the experimental condition (self-affirmation vs. control) on AI acceptance would occur through three possible mediators: positive emotions (H5a), job insecurity (H5b), and trust in AI (H5c).

The mediation analysis examined the indirect effects of the experimental condition on AI acceptance through each mediator.

Starting with Hypothesis H5a, the indirect effect through positive emotions, although positive, was not significant, $b = .046$, 95% $CI [-.009; .109]$. The mediating effect was not strong enough, thus H5a is not supported. Still, the results reveal an interesting pattern: the experimental condition had a significant effect on increasing positive emotions, but the effect of these emotions on AI acceptance was only marginal.

Regarding Hypothesis H5b, the indirect effect through job insecurity was not significant, $b = -.009$, 95% $CI [-.040; .014]$. Consequently, H5b is not supported.

H5c was also not supported, as the indirect effect through trust in AI was not significant, $b = .040$, 95% $CI [-.081; .179]$. This result does not contradict the central role of confidence in AI acceptance, which proved to be a strong and significant factor in acceptance, but indicates that the self-affirmation intervention was not effective in altering confidence levels in AI.

Finally, the total indirect effect of the intervention on the mediators was assessed. The result was not significant ($b = .079$, 95% $CI [-.060; 0.237]$), reinforcing the absence of overall mediation.

To assess the veracity of the results, the mediation model was recalculated but with the inclusion of control and demographic variables. The inclusion of these controls did not significantly alter the main conclusions of the study: trust in AI remained the only statistically significant factor in AI acceptance ($\beta = 0.433$; $p < .001$), while self-affirmation continued to show no direct or indirect effects on technology acceptance ($\beta = -0.041$; $p = .708$). Similarly, the mediating variables were not significant as mediators of the relationship between the experimental condition and AI acceptance.

The inclusion of control variables increased the explanatory power of the model, explaining 46% of the variance in AI acceptance ($R^2 = .459$). This indicates that individual characteristics, such as familiarity ($\beta = 0.195$; $p = .0001$) and comfort with AI ($\beta = 0.183$; $p = .0023$), as well as age ($\beta = 0.013$; $p = .0395$), help explain why some participants accept the technology more than others.

However, the inclusion of these factors did not alter the main relationships tested in the study, showing that the results regarding the role of self-affirmation and mediators remained unchanged after the inclusion of additional variables.

More information is available in the PROCESS output in Appendix E.

5. Discussion

5.1. Summary of results

In this dissertation, my objective was to understand how a self-affirmation intervention can influence the acceptance of AI in the workplace, as well as to explore the mediating role of positive emotions, perceived job insecurity, and trust in AI in this relationship. To this end, five hypotheses were formulated, testing both the direct relationships between this intervention and AI acceptance, and the possible indirect effects through the mediation of variables (i.e., trust in AI, perceived job insecurity, and felt positive emotions).

Before the hypotheses were tested, bivariate correlations were analyzed between all study variables, including primary variables, control variables, and demographic variables.

These correlations revealed some important patterns. Positive emotions were positively correlated with trust in AI and acceptance of AI, while trust in AI showed the strongest association with acceptance of AI. Job insecurity was not significantly correlated with acceptance of AI and showed associations primarily with negative emotions. The experimental condition was positively correlated only with positive emotions, suggesting that self-affirmation acted primarily at the emotional rather than cognitive level. Among the demographic variables, some significant but minor correlations were observed.

Regarding hypothesis testing, the analysis was conducted using PROCESS, which allowed for an evaluation of the direct and indirect effects of self-affirmation. The results show, firstly, that the experimental manipulation had a significant effect on increasing participants' positive emotions, but did not significantly alter job insecurity or confidence in AI. Thus, the effectiveness of self-affirmation manifested only at the emotional level.

Moving on to the associations between the mediating variables and the dependent variable, trust in AI emerged as a strong and significant factor in the acceptance of AI, thus confirming H3 and reinforcing its importance. On the other hand, neither positive emotions nor job insecurity have a significant impact on acceptance, which led to hypotheses H1 and H2 not being supported.

Analysis of the direct and total effects of the experimental condition showed that self-affirmation did not increase acceptance of AI, either before or after the inclusion of mediators in the model; thus, H4 is not supported. Finally, the indirect effects of self-affirmation through positive emotions, job insecurity, and trust in AI were not statistically significant, since all bootstrap confidence intervals included zero. Thus, the mediation hypotheses H5a, H5b, and H5c were not supported.

Taken together, these results indicate that, although self-affirmation did manage to influence the participants' emotional state, this effect did not translate into greater confidence in the technology or greater acceptance of AI. Acceptance of AI appears to depend primarily on stable cognitive factors, especially trust in AI, and less on momentary emotional changes induced by a brief intervention.

5.2.Connection to the existing literature

The results of this study are partially consistent with what is already known in the existing literature on AI acceptance and the psychological processes associated with this acceptance. Aligned with previous studies, trust in AI emerged as the strongest factor in technological acceptance (Glikson & Woolley, 2020; Venkatesh et al., 2003). This result reinforces that trust functions as a central cognitive mechanism that allows individuals to assess the reliability and security of intelligent systems, reducing the uncertainty inherent in their use.

Regarding emotions, although classic studies argue that positive emotions can facilitate more open attitudes towards technological change (Fredrickson, 2001), this effect was not clearly confirmed in the context of AI acceptance. One possible explanation is that, in this case, the induced emotions were moderate and short-lived, having a limited impact when compared to stronger cognitive determinants, such as trust.

The perception of job insecurity, which is often associated with a sense of technological threat (Frey & Osborne, 2017), did not have a significant effect on AI acceptance. This result does not confirm some theoretical and popular predictions, but it is in line with recent research suggesting that the relationship between job threat and technology acceptance is more complex and is not direct, obvious, or linear (Makarius et al., 2020).

Finally, the literature on self-affirmation has shown that, in various contexts, when subjected to self-affirmation, people tend to become less defensive and more open to new ideas and information (Sherman & Cohen, 2006). However, in this study, self-affirmation did not have the same effect, suggesting that the influence of self-affirmation may depend on the type of threat, the level of emotional involvement, or the context in which it is applied.

5.3. Implications

The results of this study contribute to the literature on the acceptance of AI technologies, providing further evidence that psychological factors play an important role in how individuals react to the introduction of AI in the workplace. Specifically, trust in AI emerges as the factor that most directly influences its acceptance. This conclusion reinforces the idea already present in the literature that perceptions of reliability, usefulness, and control are fundamental to the acceptance of intelligent systems (Venkatesh et al., 2003).

Furthermore, the results regarding positive emotions suggest that emotional states may play an indirect role in the acceptance of AI. Although the mediation was only partial, this example demonstrates that more positive emotional experiences can lead to more favorable attitudes toward AI.

On the other hand, the absence of effects of job insecurity contradicts the idea, often discussed in the media and public discourse (Makarius et al., 2020), that the perception of a threat to employment automatically translates into resistance to AI. In this study, this relationship was not found, suggesting that job insecurity may not be an influencing factor in the acceptance of AI.

Finally, the fact that the self-affirmation intervention did not produce the expected effects suggests that, specifically on this issue, this type of mechanism alone may not be sufficient to change attitudes. These results contribute to the discussion on the limits of self-affirmation interventions.

From a practical standpoint, the results can guide organizations wishing to implement AI tools in the workplace. As already mentioned, trust emerged as the main determinant of AI acceptance. Thus, companies and managers should prioritize strategies that increase the perception of transparency, reliability, and understanding of the technology. The literature consistently demonstrates that trust in AI systems increases when users understand their

operation, limitations, and objectives (Glikson & Woolley, 2020; Zhang & Dafoe, 2019). Practices such as clear explanations about algorithms, verification mechanisms, and human supervision increase perceived predictability, reducing fears and strengthening acceptance (Davenport & Ronanki, 2018; Raisch & Krakowski, 2021).

Regarding the perception of job insecurity, the fact that this variable did not influence AI acceptance suggests that the introduction of intelligent technologies does not necessarily generate resistance among employees. However, it is crucial that organizations clearly communicate the role of AI in the workplace to avoid misinterpretations or unjustified fears. The literature has shown that a lack of information or unclear communication about the objectives of AI increases uncertainty and fears related to job replacement (Cave et al., 2019). Finally, the fact that the self-affirmation intervention did not increase acceptance of AI suggests that organizations should focus less on strategies that seek to change workers' momentary emotional states and more on conditions that effectively promote confidence and competence in the use of technology. The literature indicates that adequate training, initial mentoring, ongoing support, and intuitive interfaces are factors that significantly increase confidence and reduce resistance to the use of new technologies (Venkatesh et al., 2003; McKnight et al., 2011; Davenport & Ronanki, 2018).

5.4. Limitations and future studies

This research has some limitations that need to be considered. One of the most obvious limitations concerns the size and composition of the sample. Although the data collected were sufficient to perform the proposed analyses, the sample was relatively small ($N = 200$). Furthermore, it was a convenience sample, since participants were not randomly selected but based on availability and accessibility, therefore not representing the general working population. The sample was used due to the need to quickly obtain a viable sample, considering the time constraints associated with the project. It is possible to assume that a larger and more representative sample could have led to different conclusions or revealed effects not observed in this study.

Secondly, a hypothetical scenario was used in an attempt to reproduce what might happen in a real work environment. Although participants were asked to imagine that they were actually in the proposed situation, this does not completely replicate the reactions that would occur in a real context. The literature argues that different choices in constructing these scenarios, such as

the type of AI described, the degree of autonomy attributed to the system, and the language used, can influence perceptions of trust and security (Glikson & Woolley, 2020; Longoni et al., 2019; Lee & See, 2004). Thus, it is always possible to question whether the results might have been different if obtained in a real-world scenario, where factors such as situational pressure, team dynamics, or the real consequences of decisions can play a significant role in trust and acceptance (Mayer et al., 1995; Siau & Wang, 2018). Future studies could utilize more realistic contexts, such as more personalized simulations or even field experiments in organizational environments. Given that many companies are already implementing AI systems in their operations, this experiment would not be difficult to implement, and the results could be measured through real feedback from company employees.

Another limitation stems from the fact that the questionnaire must be kept relatively short, given the voluntary nature of participation. This decision was intentional, to avoid dropouts and ensure that participants answered carefully until the end. However, this limited the number of items available to measure some variables, which may have reduced the reliability of the scales. Of all the measures used, only the positive and negative emotions scale was based on a validated and widely used scale, the PANAS (Watson et al., 1988). However, even this scale had to be adapted and have its number of items reduced to ensure that the questionnaire remained concise and suitable for the voluntary, online data collection format. The other scales (job insecurity, trust in AI, and acceptance of AI) were constructed from items inspired by different studies, which may have limited the reliability and sensitivity of the measures. Although all scales showed good reliability ($\alpha > 0.80$), future research could opt for longer questionnaires or other more robust scales if data collection does not have such a rigid time limit.

Furthermore, it is possible that the two tasks requested (self-affirmation and active control) were not equivalent in terms of cognitive effort, emotion, or motivation. The self-affirmation task required participants to reflect on their personal values and their importance in professional life, which is a deeper, more introspective, and potentially more demanding activity. In contrast, the control task, which asked participants to describe their commute to work, was simpler, more automatic, and required much less involvement.

This difference in task demands may have led some participants in the self-affirmation condition to feel more tired, bored, or less motivated throughout the questionnaire, and this may have influenced their responses. However, the number of participants was relatively similar in

both conditions, so it is unlikely that any increase in tiredness or lower motivation alone explains the results. Another possibility is that the self-affirmation manipulation was too brief. Studies show that longer and more structured self-affirmation interventions tend to produce more robust effects, while shorter versions may not be strong enough to activate self-affirmation mechanisms (McQueen & Klein, 2006; Cohen & Sherman, 2014; Creswell et al., 2005).

5.5.Conclusion

This study sought to understand whether a brief self-affirmation intervention could influence the acceptance of AI in the workplace and how positive emotions, perceived job insecurity, and trust in AI could act as mediating mechanisms for this effect. Although self-affirmation did not produce the expected direct impact, the study revealed that it significantly increased positive emotions, although there was no subsequent indirect impact on AI acceptance. Among the variables studied, only trust emerged as a consistent and strong factor in acceptance.

The results suggest that momentary emotional states, even when successfully induced, may not be sufficient to change relatively stable attitudes towards AI. Additionally, the absence of effects on job insecurity aligns with recent literature that points to a more complex relationship than is often assumed in public discourse.

Despite the study's limitations, the results offer guidance for future research and highlight the importance of continuing to explore how different psychological factors change the way workers adapt to technological transformations. In summary, the main contribution of this dissertation is to demonstrate that the acceptance of AI depends not only on technological characteristics, but also on fundamental psychological processes, which are essential to promoting a more balanced, human, and conscious integration of AI in the workplace.

6. References

- Baron, R. M., & Kenny, D. A. (1986). The moderator-mediator variable distinction in social psychological research: Conceptual, strategic, and statistical consideration. *Journal of personality and social psychology*, 51(6), 1173. <https://doi.org/10.1037/0022-3514.51.6.1173>
- Bean, R. (2019). Why fear of disruption is driving investment in AI. *MIT Sloan Management Review*. <https://sloanreview.mit.edu/article/why-fear-of-disruption-is-driving-investment-in-ai/>
- Brockner, J., & Wiesenfeld, B. M. (1996). An integrative framework for explaining reactions to decisions: interactive effects of outcomes and procedures. *Psychological bulletin*, 120(2), 189. <https://doi.org/10.1037/0033-2909.120.2.189>
- Brougham, D., & Haar, J. (2017). Employee assessment of their technological redundancy. *Labour & Industry: a journal of the social and economic relations of work*, 27(3), 213-231. <https://doi.org/10.1080/10301763.2017.1369718>
- Brougham, D., & Haar, J. (2018). Smart Technology, Artificial Intelligence, Robotics, and Algorithms (STARA): Employees' perceptions of our future workplace. *Journal of Management & Organization*, 24(2), 239–257. <https://doi.org/10.1017/jmo.2016.55>
- Brougham, D., & Haar, J. (2020). Technological disruption and employment: The influence on job insecurity and turnover intentions: A multi-country study. *Technological Forecasting and Social Change*, 161, 120276. <https://doi.org/10.1016/j.techfore.2020.120276>
- Burrell, J. (2016). How the machine 'thinks': Understanding opacity in machine learning algorithms. *Big data & society*, 3(1), 2053951715622512. <https://doi.org/10.1177/2053951715622512>
- Cave, S., Coughlan, K., & Dihal, K. (2019, January). " Scary Robots" Examining Public Responses to AI. In *Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society* (pp. 331-337). <https://doi.org/10.1145/3306618.3314232>

- Cheng, G. H. L., & Chan, D. K. S. (2008). Who suffers more from job insecurity? A meta-analytic review. *Applied psychology*, 57(2), 272-303. <https://doi.org/10.1111/j.1464-0597.2007.00312.x>
- Cohen, G. L., & Sherman, D. K. (2014). The psychology of change: Self-affirmation and social psychological intervention. *Annual review of psychology*, 65(1), 333-371. <https://doi.org/10.1146/annurev-psych-010213-115137>
- Compeau, D. R., & Higgins, C. A. (1995). Computer self-efficacy: Development of a measure and initial test. *MIS quarterly*, 189-211. <https://doi.org/10.2307/249688>
- Creswell, J. D., Welch, W. T., Taylor, S. E., Sherman, D. K., Gruenewald, T. L., & Mann, T. (2005). Affirmation of personal values buffers neuroendocrine and psychological stress responses. *Psychological science*, 16(11), 846-851. <https://doi.org/10.1111/j.1467-9280.2005.01624.x>
- Davenport, T. H. (2018). The AI advantage: How to put the artificial intelligence revolution to work. *MIT Press*. <https://doi.org/10.7551/mitpress/11781.001.0001>
- Davenport, T. H., & Ronanki, R. (2018). Artificial intelligence for the real world. *Harvard business review*, 96(1), 108-116.
- Davis, F. D. (1989). Perceived Usefulness, Perceived Ease of Use, and User Acceptance of Information Technology. *MIS Quarterly*, 13(3), 319-340. <https://doi.org/10.2307/249008>
- De Visser, E. J., Pak, R., & Shaw, T. H. (2018). From ‘automation’ to ‘autonomy’: the importance of trust repair in human-machine interaction. *Ergonomics*, 61(10), 1409-1427. <https://doi.org/10.1080/00140139.2018.1457725>
- De Witte, H. (1999). Job Insecurity and Psychological Well-being: Review of the Literature and Exploration of Some Unresolved Issues. *European Journal of Work and Organizational Psychology*, 8(2), 155-177. <https://doi.org/10.1080/135943299398302>
- De Witte, H., Vander Elst, T., De Cuyper, N. (2015). Job Insecurity, Health and Well-Being. In: Vuori, J., Blonk, R., Price, R. (eds) *Sustainable Working Lives. Aligning Perspectives*

on Health, Safety and Well-Being. Springer, Dordrecht. https://doi.org/10.1007/978-94-017-9798-6_7

- Dwivedi, Y. K., Hughes, L., Ismagilova, E., Aarts, G., Coombs, C., Crick, T., ... & Williams, M. D. (2021). Artificial Intelligence (AI): Multidisciplinary perspectives on emerging challenges, opportunities, and agenda for research, practice and policy. *International journal of information management*, 57, 101994. <https://doi.org/10.1016/j.ijinfomgt.2019.08.002>
- Escobar-Soler, C., Berrios, R., Peñaloza-Díaz, G., Melis-Rivera, C., Caqueo-Urizar, A., Ponce-Correa, F., & Flores, J. (2024). Effectiveness of Self-Affirmation Interventions in Educational Settings: A Meta-Analysis. *Healthcare*, 12(1), 3. <https://doi.org/10.3390/healthcare12010003>
- Flynn, L. R., & Goldsmith, R. E. (1999). A short, reliable measure of subjective knowledge. *Journal of business research*, 46(1), 57-66. [https://doi.org/10.1016/S0148-2963\(98\)00057-5](https://doi.org/10.1016/S0148-2963(98)00057-5)
- Fredrickson, B. L. (2001). The role of positive emotions in positive psychology: The broaden-and-build theory of positive emotions. *American psychologist*, 56(3), 218. <https://doi.org/10.1037/0003-066X.56.3.218>
- Fredrickson, B. L. (2004). The broaden-and-build theory of positive emotions. *Philosophical transactions of the royal society of London. Series B: Biological Sciences*, 359(1449), 1367-1377. <https://doi.org/10.1098/rstb.2004.1512>
- Frey, C. B., & Osborne, M. A. (2017). The future of employment: How susceptible are jobs to computerisation?. *Technological forecasting and social change*, 114, 254-280. <https://doi.org/10.1016/j.techfore.2016.08.019>
- Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of management annals*, 14(2), 627-660. <https://doi.org/10.5465/annals.2018.0057>
- Haenlein, M., & Kaplan, A. (2019). A Brief History of Artificial Intelligence: On the Past, Present, and Future of Artificial Intelligence. *California Management Review*, 61(4), 5-14. <https://doi.org/10.1177/0008125619864925>

- Harvey, R. D., & Oswald, D. L. (2000). Collective guilt and shame as motivation for White support of Black programs 1. *Journal of Applied Social Psychology*, 30(9), 1790-1811. <https://doi.org/10.1111/j.1559-1816.2000.tb02468.x>
- Hayes, A. F. (2017). Introduction to mediation, moderation, and conditional process analysis: A regression-based approach. *Guilford publications*. <https://doi.org/10.1111/jedm.12050>
- Hoff, K. A., & Bashir, M. (2015). Trust in automation: Integrating empirical evidence on factors that influence trust. *Human factors*, 57(3), 407-434. <https://doi.org/10.1177/00187208145475>
- Jussupow, E., Spohrer, K., Heinzl, A., & Gawlitza, J. (2021). Augmenting medical diagnosis decisions? An investigation into physicians' decision-making process with artificial intelligence. *Information Systems Research*, 32(3), 713-735. <https://doi.org/10.1287/isre.2020.0980>
- Kaplan, A., & Haenlein, M. (2019). Siri, Siri, in my hand: Who's the fairest in the land? On the interpretations, illustrations, and implications of artificial intelligence. *Business horizons*, 62(1), 15-25. <https://doi.org/10.1016/j.bushor.2018.08.004>
- Koole, S. L., Smeets, K., Van Knippenberg, A., & Dijksterhuis, A. (1999). The cessation of rumination through self-affirmation. *Journal of Personality and Social Psychology*, 77(1), 111. <https://doi.org/10.1037/0022-3514.77.1.111>
- Lee, J. D., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human factors*, 46(1), 50-80. https://doi.org/10.1518/hfes.46.1.50_30392
- Legault, L., Al-Khindi, T., & Inzlicht, M. (2012). Preserving integrity in the face of performance threat: Self-affirmation enhances neurophysiological responsiveness to errors. *Psychological science*, 23(12), 1455-1460. <https://doi.org/10.1177/0956797612448483>
- Longoni, C., Bonezzi, A., & Morewedge, C. K. (2019). Resistance to medical artificial intelligence. *Journal of consumer research*, 46(4), 629-650. <https://doi.org/10.1093/jcr/ucz013>

- Makarius, E. E., Mukherjee, D., Fox, J. D., & Fox, A. K. (2020). Rising with the machines: A sociotechnical framework for bringing artificial intelligence into the organization. *Journal of business research*, 120, 262-273. <https://doi.org/10.1016/j.jbusres.2020.07.045>
- Manyika, J., Lund, S., Chui, M., Bughin, J., Woetzel, J., Batra, P., ... & Sanghvi, S. (2017). Jobs lost, jobs gained: What the future of work will mean for jobs, skills, and wages. *McKinsey Global Institute*
- Maslej, N., Fattorini, L., Perrault, R., Gil, Y., Parli, V., Kariuki, N., ... & Oak, S. (2025). Artificial intelligence index report 2025. *arXiv preprint arXiv:2504.07139*. <https://doi.org/10.48550/arXiv.2504.07139>
- Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *Academy of management review*, 20(3), 709-734. <https://doi.org/10.5465/amr.1995.9508080335>
- McAfee, A., & Brynjolfsson, E. (2017). Machine, platform, crowd: Harnessing our digital future. *WW Norton & Company* <https://doi.org/10.3917/futur.427.0121f>
- Mcknight, D. H., Carter, M., Thatcher, J. B., & Clay, P. F. (2011). Trust in a specific technology: An investigation of its components and measures. *ACM Transactions on management information systems (TMIS)*, 2(2), 1-25. <https://doi.org/10.1145/1985347.1985353>
- MCQueen, A., & Klein, W. M. P. (2006). Experimental manipulations of self-affirmation: A systematic review. *Self and Identity*, 5(4), 289-354. <https://doi.org/10.1080/15298860600805325>
- Moor, J. (2006). The Dartmouth College artificial intelligence conference: The next fifty years. *AI magazine*, 27(4), 87-87. <https://doi.org/10.1609/aimag.v27i4.1911>
- OECD. (2019). OECD Employment Outlook 2019: The Future of Work. *OECD Publishing*. <https://doi.org/10.1787/9ee00155-en>

- Parasuraman, A. (2000). Technology Readiness Index (TRI) a multiple-item scale to measure readiness to embrace new technologies. *Journal of service research*, 2(4), 307-320. <https://doi.org/10.1177/10946705002400>
- Philbeck, T., & Davis, N. (2018). THE FOURTH INDUSTRIAL REVOLUTION: SHAPING A NEW ERA. *Journal of International Affairs*, 72(1), 17–22. <https://www.jstor.org/stable/26588339>
- Raisch, S., & Krakowski, S. (2021). Artificial intelligence and management: The automation–augmentation paradox. *Academy of management review*, 46(1), 192-210. <https://doi.org/10.5465/amr.2018.0072>
- Rogers, E. M., Singhal, A., & Quinlan, M. M. (2014). Diffusion of innovations. *In An integrated approach to communication theory and research* (pp. 432-448). Routledge. <https://doi.org/10.4324/9780203710753-35>
- Russell, S., Norvig, P., & Intelligence, A. (1995). A modern approach. Artificial Intelligence. *Prentice-Hall, Englewood Cliffs*, 25(27), 79-80. <https://doi.org/10.1017/s0269888900007724>
- Saks, A. M., Uggerslev, K. L., & Fassina, N. E. (2007). Socialization tactics and newcomer adjustment: A meta-analytic review and test of a model. *Journal of vocational behavior*, 70(3), 413-446. <https://doi.org/10.1016/j.jvb.2006.12.004>
- Schepman, A., & Rodway, P. (2020). Initial validation of the general attitudes towards Artificial Intelligence Scale. *Computers in human behavior reports*, 1, 100014. <https://doi.org/10.1016/j.chbr.2020.100014>
- Seo, M. G., Barrett, L. F., & Bartunek, J. M. (2004). The role of affective experience in work motivation. *Academy of management review*, 29(3), 423-439. <https://doi.org/10.5465/amr.2004.13670972>
- Sherman, D. K., & Cohen, G. L. (2006). The psychology of self-defense: Self-affirmation theory. *Advances in experimental social psychology*, 38, 183-242. [https://doi.org/10.1016/S0065-2601\(06\)38004-5](https://doi.org/10.1016/S0065-2601(06)38004-5)

- Sherman, D. K., & Hartson, K. A. (2011). Reconciling self-protection with self-improvement. *Handbook of self-enhancement and self-protection*, 128, 151.
- Shin, D. (2020). User Perceptions of Algorithmic Decisions in the Personalized AI System: Perceptual Evaluation of Fairness, Accountability, Transparency, and Explainability. *Journal of Broadcasting & Electronic Media*, 64(4), 541–565. <https://doi.org/10.1080/08838151.2020.1843357>
- Siau, K., & Wang, W. (2018). Building trust in artificial intelligence, machine learning, and robotics. *Cutter business technology journal*, 31(2), 47.
- Steele, C. M. (1988). The psychology of self-affirmation: Sustaining the integrity of the self. *In Advances in experimental social psychology (Vol. 21, pp. 261-302)*. Academic Press. [https://doi.org/10.1016/S0065-2601\(08\)60229-4](https://doi.org/10.1016/S0065-2601(08)60229-4)
- Sweeney, A. M., & Moyer, A. (2015). Self-affirmation and responses to health messages: A meta-analysis on intentions and behavior. *Health Psychology*, 34(2), 149–159. <https://doi.org/10.1037/hea0000110>
- Vander Elst, T., De Witte, H., & De Cuyper, N. (2013). The Job Insecurity Scale: A psychometric evaluation across five European countries. *European Journal of Work and Organizational Psychology*, 23(3), 364–380. <https://doi.org/10.1080/1359432X.2012.745989>
- Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS quarterly*, 425-478. <https://doi.org/10.2307/30036540>
- Watson, D., Clark, L. A., & Tellegen, A. (1988). Development and validation of brief measures of positive and negative affect: the PANAS scales. *Journal of personality and social psychology*, 54(6), 1063. <https://doi.org/10.1037/0022-3514.54.6.1063>
- Wiesenfeld, B. M., Brockner, J., & Martin, C. (1999). A self-affirmation analysis of survivors' reactions to unfair organizational downsizings. *Journal of Experimental Social Psychology*, 35(5), 441-460. <https://doi.org/10.1006/jesp.1999.1389>

Wilson, H. J., & Daugherty, P. R. (2018). Human+ machine: Reimagining work in the age of AI. (No Title).

World Economic Forum. (2020). The future of jobs report 2020. *World Economic Forum*. https://www3.weforum.org/docs/WEF_Future_of_Jobs_2020.pdf

Zhang, B., & Dafoe, A. (2019). Artificial intelligence: American attitudes and trends. *Available at SSRN 3312874*. <https://doi.org/10.2139/ssrn.3312874>

7. Appendix

A. Study demographic characteristics

Table 3 - Sample demographic characteristics

		Count	%
Gender	Male	62	31
	Female	135	67.5
	Non-binary / third gender	1	0.5
	Prefer not to say	2	1.0
Education Level	Less than a high school diploma	4	2.0
	High school diploma	17	8.5
	Incomplete higher education, but no diploma	12	6.0
	Bachelor's degree	78	39.2
	Master's degree	80	40.2
	Doctorate	4	2.0
	Vocational diploma	1	0.5
	Other	3	1.5

Table 4- Mean age and SD recorded

	N	Mean	Std. Deviation
Age	186	27.67	8.924

Table 5 - Sample sector distribution

	Frequency	Percent
Construction	4	2.0
Education	24	12.1
Energy	4	2.0
Finance	21	10.6
Health	18	9.0
Manufacturing	5	2.5
Retail	17	8.5
Services	39	19.6
Technology	32	16.1
Other	35	17.6

Table 6 - Sample nationality distribution

	Frequency	Percent
Angola	1	0.5
Australia	3	1.6
Austria	1	0.5
Belgium	1	0.5
Brazil	1	0.5
Canada	1	0.5
Chile	2	1.1
China	1	0.5
Colombia	1	0.5
Finland	1	0.5
France	3	1.6
Germany	7	3.7
Greece	3	1.6
Guinea	1	0.5
India	35	18.5
Indonesia	4	2.1
Ireland	2	1.1
Italy	4	2.1
Lithuania	2	1.1
Luxembourg	1	0.5
Morocco	1	0.5
Mexico	1	0.5
Netherlands	4	2.1
New Zealand	1	0.5
Pakistan	8	4.2
Poland	3	1.6
Portugal	34	18.0
Romania	2	1.1
Russia	2	1.1
Spain	1	0.5
Switzerland	2	1.1
Turkey	1	0.5
Ukraine	2	1.1
United Arab Emirates	1	0.5
United Kingdom	10	5.3
United States	24	12.7
Other	17	9.0

B. Experiment Survey

Dear participant,

First of all, thank you for your time and availability. My name is Maria Inês Nunes, and I am writing my master's thesis in Business at the Portuguese Catholic University.

This study seeks to understand perceptions related to work and new technologies such as Artificial Intelligence. Participation in this questionnaire is completely voluntary, anonymous, and confidential, and will take approximately 5 to 8 minutes. At no point will you be asked for information that could personally identify you.

Please answer all questions honestly so that the study is as reliable as possible. You can stop the questionnaire at any time by simply closing it, without any consequences.

The responses collected will be used exclusively for academic purposes. If you have any questions, please contact me by email at: s-marianunes@ucp.pt

By clicking the arrow, you confirm that you have read this information, understand the purpose of the study, and agree to participate.

Thank you very much for your cooperation!

Do you currently work, or have you ever worked (internships, part-time jobs, full-time jobs, summer jobs, etc.)?

- ☐ Yes, I currently work
- ☐ Yes, I have worked before
- ☐ No, I have never worked

Please, indicate the name of the company you currently work for or have worked for.

Self-affirmation - Please do the following task: Choose only one value from the list below that you consider most important in the context of work (e.g., integrity, competence, helping others,

responsibility, resilience, fairness, collaboration). Write 3–5 lines explaining why this value is important in your work and describe a specific episode where you experienced it (what happened, what you did, impact).

Control - Please do the following task: Describe your usual commute to work in 3–5 lines (e.g., how you commute, what you do on the way).

For each of the following statements, select the option that best represents your level of agreement or disagreement. (1 = Strongly disagree; 7 = Strongly agree):

The writing assignment made me reflect on what I consider important in my work.

- ☐ Strongly disagree
- ☐ Somewhat disagree
- ☐ Neither agree nor disagree
- ☐ Somewhat agree
- ☐ Strongly agree

After completing the writing assignment, I felt more aware of my personal values.

- ☐ Strongly disagree
- ☐ Somewhat disagree
- ☐ Neither agree nor disagree
- ☐ Somewhat agree
- ☐ Strongly agree

Indicate how you feel at this moment in relation to each word (1 = Not at all; 5 = Extremely).

	None at all	A little	A moderate amount	A lot	Extremely
Enthusiastic	☐	☐	☐	☐	☐
Nervous	☐	☐	☐	☐	☐
Interested	☐	☐	☐	☐	☐
Attentive	☐	☐	☐	☐	☐
Anxious	☐	☐	☐	☐	☐
Proud	☐	☐	☐	☐	☐
Worried	☐	☐	☐	☐	☐
Bored	☐	☐	☐	☐	☐

Please imagine that the management at the organization you work for, (if you don't currently work there, please imagine that you do) is implementing a new Artificial Intelligence (AI) system and sends the following email:

Dear colleagues,

As part of our ongoing digital transformation, we will soon implement a new Artificial Intelligence (AI) system designed to support employees in their daily activities. This system will assist in analyzing work data, prioritizing tasks, and decision-making, for example, identifying more efficient options, automating repetitive tasks such as organizing information, filling out reports, or scheduling activities. We believe that this technology will help increase efficiency, reduce errors, and improve collaboration between teams. However, we recognize that this change may also alter the way some tasks are performed, and, in some cases, certain roles may evolve or become redundant. Our goal is to ensure that this transition benefits both the company and our employees, helping everyone focus on higher-value and more creative tasks.

— The Management Team

Select the option that best represents your level of agreement or disagreement: (1 = Strongly disagree; 5 = Strongly agree):

It was easy to imagine myself in the situation described.

- ☐ Strongly disagree
- ☐ Somewhat disagree
- ☐ Neither agree nor disagree
- ☐ Somewhat agree
- ☐ Strongly agree

Considering the scenario described above, select the option that best represents your level of agreement or disagreement regarding the impact of AI on your professional future: (1 = Strongly disagree; 5 = Strongly agree):

	Strongly disagree	Somewhat disagree	Neither agree nor disagree	Somewhat agree	Strongly agree
I am concerned about possible job losses due to Artificial Intelligence.	☐	☐	☐	☐	☐
With the introduction of this system, I fear that my job will become less relevant.	☐	☐	☐	☐	☐
The presence of Artificial Intelligence would increase uncertainty about my professional future.	☐	☐	☐	☐	☐
I feel that Artificial Intelligence could reduce the need for my role and replace me, and this worries me.	☐	☐	☐	☐	☐

Considering the scenario described before and how you felt towards the proposed implementation of artificial intelligence, select the option that best represents your level of agreement or disagreement: (1 = Strongly disagree; 5 = Strongly agree)

	Strongly disagree	Somewhat disagree	Neither agree nor disagree	Somewhat agree	Strongly agree
If my organization made this system available, I would use it.	☐	☐	☐	☐	☐
If I had access, I would make regular use of Artificial Intelligence in the performance of my duties.	☐	☐	☐	☐	☐
I intend to learn how to use this system.	☐	☐	☐	☐	☐
I would recommend the use of the system to my team.	☐	☐	☐	☐	☐

Considering the scenario described before and how you felt towards the proposed implementation of artificial intelligence, select the option that best represents your level of agreement or disagreement: (1 = Strongly disagree; 5 = Strongly agree)

	Strongly disagree	Somewhat disagree	Neither agree nor disagree	Somewhat agree	Strongly agree
I would trust the recommendations provided by this system.	☐	☐	☐	☐	☐
Artificial Intelligence would act fairly in decisions related to my work.	☐	☐	☐	☐	☐
I can rely on Artificial Intelligence to reduce operational errors.	☐	☐	☐	☐	☐
Decisions made by Artificial Intelligence would be impartial.	☐	☐	☐	☐	☐

How often do you use AI tools in your work?

- ☐ Never
- ☐ Rarely
- ☐ Sometimes
- ☐ Regularly
- ☐ Frequently

Select the option that best represents your comfort level with Artificial Intelligence technologies: (1 = Extremely uncomfortable; 5 = Extremely comfortable):

- ☐ Extremely uncomfortable
- ☐ Somewhat uncomfortable
- ☐ Neither comfortable nor uncomfortable
- ☐ Somewhat comfortable
- ☐ Extremely comfortable

I have received training on AI at work.

- ☐ Yes
- ☐ No

What sector do you work or have you worked in?

- ☐ Construction
- ☐ Education
- ☐ Energy
- ☐ Finance
- ☐ Health
- ☐ Manufacturing
- ☐ Retail
- ☐ Services
- ☐ Technology
- ☐ Other

Please specify your gender:

- ☐ Male
- ☐ Female
- ☐ Non-binary / third gender
- ☐ Prefer not to say

How old are you?

What is the highest level of education you have completed or the highest degree you have received?

- ☐ Less than a high school diploma
- ☐ High school diploma
- ☐ Incomplete higher education, but no diploma
- ☐ Bachelor's degree
- ☐ Master's degree
- ☐ Doctorate
- ☐ Vocational diploma
- ☐ Other

What is your nationality?

select an option --- 

Thank you for participating!

This study analyzes how different forms of personal reflection (such as writing about values that are important to you) can influence perceptions about the adoption of Artificial Intelligence at work.

The writing tasks were used only to create comparison groups, and there are no "right" or "wrong" answers.

We emphasize that all responses are anonymous and confidential. If you have any questions or would like more information, you can contact me at: s-marianunes@ucp.pt

Thank you very much for your time and cooperation!

C. Bivariate correlations

Table 7 - Bivariate correlations

		1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17
1- Job insecurity	Person Correlation		0,118	,308***	,214**	0,053	0,086	0,084	0,023	0,008	-0,012	0,107	0,031	-0,064	-,238***	0,057	-0,035	-,146*
	Sig. (2-tailed)	-	0,097	<,001	0,002	0,459	0,228	0,238	0,751	0,912	0,869	0,133	0,662	0,366	<,001	0,435	0,629	0,047
2- Positive emotions	Person Correlation	0,118		,229**	,295***	,229**	,229**	0,069	0,077	0,057	-0,034	0,024	-0,105	,214**	0,017	0,079	-0,083	0,095
	Sig. (2-tailed)	0,097		0,001	<,001	0,001	0,001	0,333	0,279	0,427	0,633	0,741	0,139	0,002	0,819	0,28	0,255	0,196
3- Negative emotions	Person Correlation	,308***	,229**		,142*	0,029	0,103	0,025	0,049	0,096	-0,119	-0,02	-0,025	-0,041	-0,125	0,038	-0,038	-,153*
	Sig. (2-tailed)	<,001	0,001		0,044	0,683	0,148	0,725	0,495	0,177	0,094	0,779	0,721	0,563	0,086	0,608	0,602	0,037
4- Trust AI	Person Correlation	,214**	,295***	,142*		,517***	0,045	0,136	-0,089	,142*	0,022	0,037	0,017	0,066	,157*	-0,056	-0,081	-0,068
	Sig. (2-tailed)	0,002	<,001	0,044		<,001	0,53	0,054	0,213	0,046	0,756	0,604	0,806	0,357	0,031	0,44	0,266	0,357
5- AI Acceptance	Person Correlation	0,053	,229**	0,029	,517***		-0,063	0,058	-0,114	0	-0,023	0,091	0,047	0,122	0,125	-0,044	-0,012	0,063
	Sig. (2-tailed)	0,459	0,001	0,683	<,001		0,378	0,418	0,11	0,996	0,75	0,2	0,507	0,086	0,086	0,552	0,875	0,396
6- Condition	Person Correlation	0,086	,229**	0,103	0,045	-0,063		-0,033	-0,003	,174*	-0,001	-0,049	0,069	-0,074	-0,008	0,034	-0,089	0,017
	Sig. (2-tailed)	0,228	0,001	0,148	0,53	0,378		0,643	0,967	0,014	0,991	0,49	0,334	0,301	0,912	0,638	0,222	0,818
7- Gender	Person Correlation	0,084	0,069	0,025	0,136	0,058	-0,033		-,182**	0,016	-0,031	,149*	0,016	0,024	-0,013	-0,022	-0,047	,148*
	Sig. (2-tailed)	0,238	0,333	0,725	0,054	0,418	0,643		0,01	0,821	0,659	0,036	0,828	0,739	0,86	0,765	0,52	0,044
8 - Sector education	Person Correlation	0,023	0,077	0,049	-0,089	-0,114	-0,003	-,182**		-0,127	-,183**	-,162*	-0,066	0,055	-0,042	-0,047	-0,04	-0,027
	Sig. (2-tailed)	0,751	0,279	0,495	0,213	0,11	0,967	0,01		0,073	0,01	0,022	0,352	0,443	0,566	0,526	0,585	0,72
9 - Sector finance	Person Correlation	0,008	0,057	0,096	,142*	0	,174*	0,016	-0,127		-,170*	-,150*	0,058	-0,083	0,129	0,123	-0,07	-0,022
	Sig. (2-tailed)	0,912	0,427	0,177	0,046	0,996	0,014	0,821	0,073		0,017	0,034	0,417	0,245	0,078	0,093	0,338	0,764
10 - Sector services	Person Correlation	-0,012	-0,034	-0,119	0,022	-0,023	-0,001	-0,031	-,183**	-,170*		-,216**	0,043	-0,071	-0,123	-0,059	0,016	-0,01
	Sig. (2-tailed)	0,869	0,633	0,094	0,756	0,75	0,991	0,659	0,01	0,017		0,002	0,552	0,318	0,092	0,42	0,823	0,891
11 - Sector technology	Person Correlation	0,107	0,024	-0,02	0,037	0,091	-0,049	,149*	-,162*	-,150*	-,216**		-0,101	,142*	-0,029	0,002	0,124	0,086
	Sig. (2-tailed)	0,133	0,741	0,779	0,604	0,2	0,49	0,036	0,022	0,034	0,002		0,156	0,046	0,693	0,983	0,091	0,242
12 - Bachelor	Person Correlation	0,031	-0,105	-0,025	0,017	0,047	0,069	0,016	-0,066	0,058	0,043	-0,101		-,658***	0,088	-0,118	-0,025	-0,118
	Sig. (2-tailed)	0,662	0,139	0,721	0,806	0,507	0,334	0,828	0,352	0,417	0,552	0,156		<,001	0,227	0,106	0,731	0,108
13 - Master	Person Correlation	-0,064	,214**	-0,041	0,066	0,122	-0,074	0,024	0,055	-0,083	-0,071	,142*	-,658***		-0,042	,226**	-0,017	,219**
	Sig. (2-tailed)	0,366	0,002	0,563	0,357	0,086	0,301	0,739	0,443	0,245	0,318	0,046	<,001		0,566	0,002	0,816	0,003
14 - Portugal	Person Correlation	-,238***	0,017	-0,125	,157*	0,125	-0,008	-0,013	-0,042	0,129	-0,123	-0,029	0,088	-0,042		-,223**	-,179*	0,055
	Sig. (2-tailed)	<,001	0,819	0,086	0,031	0,086	0,912	0,86	0,566	0,078	0,092	0,693	0,227	0,566		0,002	0,014	0,463
15 - India	Person Correlation	0,057	0,079	0,038	-0,056	-0,044	0,034	-0,022	-0,047	0,123	-0,059	0,002	-0,118	,226**	-,223**		-,182*	-0,101
	Sig. (2-tailed)	0,435	0,28	0,608	0,44	0,552	0,638	0,765	0,526	0,093	0,42	0,983	0,106	0,002	0,002		0,012	0,178
16 - US	Person Correlation	-0,035	-0,083	-0,038	-0,081	-0,012	-0,089	-0,047	-0,04	-0,07	0,016	0,124	-0,025	-0,017	-,179*	-,182*		0,105
	Sig. (2-tailed)	0,629	0,255	0,602	0,266	0,875	0,222	0,52	0,585	0,338	0,823	0,091	0,731	0,816	0,014	0,012		0,161
17 - Age	Person Correlation	-,146*	0,095	-,153*	-0,068	0,063	0,017	,148*	-0,027	-0,022	-0,01	0,086	-0,118	,219**	0,055	-0,101	0,105	
	Sig. (2-tailed)	0,047	0,196	0,037	0,357	0,396	0,818	0,044	0,72	0,764	0,891	0,242	0,108	0,003	0,463	0,178	0,161	

D. PROCESS

Run MATRIX procedure:

***** PROCESS Procedure for SPSS Version 4.2 *****

Written by Andrew F. Hayes, Ph.D. www.afhayes.com

Documentation available in Hayes (2022). www.guilford.com/p/hayes3

Model : 4

Y : A_mean

X : c_dummy

M1 : JI_mean

M2 : PE_mean

M3 : T_mean

Sample

Size: 200

OUTCOME VARIABLE:

JI_mean

Model Summary

R	R-sq	MSE	F	df1	df2	p
,0856	,0073	1,2003	1,4629	1,0000	198,0000	,2279

Model

	coeff	se	t	p	LLCI	ULCI
constant	2,8935	,1054	27,4466	,0000	2,6856	3,1014
c_dummy	,1880	,1554	1,2095	,2279	-,1185	,4945

OUTCOME VARIABLE:

PE_mean

Model Summary

R	R-sq	MSE	F	df1	df2	p
,2294	,0526	,8194	10,9940	1,0000	198,0000	,0011

Model

	coeff	se	t	p	LLCI	ULCI
constant	2,7454	,0871	31,5187	,0000	2,5736	2,9171
c_dummy	,4258	,1284	3,3157	,0011	,1726	,6791

OUTCOME VARIABLE:

T_mean

Model Summary

R	R-sq	MSE	F	df1	df2	p
,0447	,0020	,7607	,3963	1,0000	198,0000	,5297

Model

	coeff	se	t	p	LLCI	ULCI
constant	3,2917	,0839	39,2207	,0000	3,1262	3,4572
c_dummy	,0779	,1237	,6295	,5297	-,1661	,3219

OUTCOME VARIABLE:

A_mean

Model Summary

R	R-sq	MSE	F	df1	df2	p
,5371	,2885	,5860	19,7653	4,0000	195,0000	,0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	2,0525	,2632	7,7989	,0000	1,5334	2,5715
c_dummy	-,1903	,1119	-1,7013	,0905	-,4109	,0303
JI_mean	-,0480	,0508	-,9444	,3461	-,1483	,0523
PE_mean	,1088	,0629	1,7299	,0852	-,0152	,2329
T_mean	,5171	,0664	7,7828	,0000	,3861	,6481

***** TOTAL EFFECT MODEL *****

OUTCOME VARIABLE:

A_mean

Model Summary

R	R-sq	MSE	F	df1	df2	p
,0627	,0039	,8079	,7813	1,0000	198,0000	,3778

Model

	coeff	se	t	p	LLCI	ULCI
constant	3,9144	,0865	45,2580	,0000	3,7438	4,0849
c_dummy	-,1127	,1275	-,8839	,3778	-,3642	,1388

***** TOTAL, DIRECT, AND INDIRECT EFFECTS OF X ON Y *****

Total effect of X on Y

Effect	se	t	p	LLCI	ULCI
-,1127	,1275	-,8839	,3778	-,3642	,1388

Direct effect of X on Y

Effect	se	t	p	LLCI	ULCI
-,1903	,1119	-1,7013	,0905	-,4109	,0303

Indirect effect(s) of X on Y:

Effect	BootSE	BootLLCI	BootULCI
TOTAL	,0776	,0763	-,0603
JI_mean	-,0090	,0131	-,0395
PE_mean	,0463	,0301	-,0098
T_mean	,0403	,0655	-,0807

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
95,0000

Number of bootstrap samples for percentile bootstrap confidence intervals:
5000

----- END MATRIX -----

E. PROCESS with control variables

Run MATRIX procedure:

***** PROCESS Procedure for SPSS Version 4.2 *****

Written by Andrew F. Hayes, Ph.D. www.afhayes.com
Documentation available in Hayes (2022).

Model : 4

Y : A_mean
X : c_dummy
M1 : JI_mean
M2 : PE_mean
M3 : T_mean
M4 : NE_mean

Covariates:

Familia Comfort Training Sector Gender edu nation age_num

Sample

Size: 178

OUTCOME VARIABLE:

JI_mean

Model Summary

R	R-sq	MSE	F	df1	df2	p
.3460	.1197	1.1230	2.5382	9.0000	168.0000	.0093

Model

	coeff	se	t	p	LLCI	ULCI
constant	4.2091	.7196	5.8490	.0000	2.7884	5.6297
c_dummy	.2346	.1614	1.4536	.1479	-.0840	.5532
Familia	.2171	.0743	2.9226	.0039	.0704	.3637
Comfort	-.1307	.0900	-1.4515	.1485	-.3084	.0470
Training	-.2975	.1812	-1.6415	.1026	-.6553	.0603
Sector	.0180	.0297	.6059	.5454	-.0406	.0766
Gender	-.1592	.1640	-.9703	.3333	-.4830	.1647
edu	-.1502	.0759	-1.9796	.0494	-.2999	-.0004
nation	-.0021	.0054	-.3921	.6955	-.0128	.0085
age_num	-.0069	.0098	-.7019	.4837	-.0262	.0125

OUTCOME VARIABLE:

PE_mean

Model Summary

R	R-sq	MSE	F	df1	df2	p
.3990	.1592	.8076	3.5354	9.0000	168.0000	.0005

Model

	coeff	se	t	p	LLCI	ULCI
constant	1.3056	.6103	2.1393	.0339	.1008	2.5104
c_dummy	.5041	.1369	3.6836	.0003	.2339	.7743
Familia	.1164	.0630	1.8485	.0663	-.0079	.2408
Comfort	.1310	.0763	1.7164	.0879	-.0197	.2817
Training	.0044	.1537	.0286	.9772	-.2990	.3078
Sector	-.0437	.0252	-1.7371	.0842	-.0934	.0060
Gender	.0428	.1391	.3077	.7587	-.2318	.3175
edu	.1295	.0643	2.0129	.0457	.0025	.2565
nation	-.0022	.0046	-.4852	.6282	-.0113	.0068
age_num	.0099	.0083	1.1909	.2354	-.0065	.0263

OUTCOME VARIABLE:

T_mean

Model Summary

R	R-sq	MSE	F	df1	df2	p
.3639	.1324	.6797	2.8498	9.0000	168.0000	.0038

Model

coeff	se	t	p	LLCI	ULCI
-------	----	---	---	------	------

constant	2.1729	.5599	3.8812	.0001	1.0676	3.2782
c_dummy	.2249	.1256	1.7913	.0750	-.0230	.4728
Familia	.1622	.0578	2.8069	.0056	.0481	.2763
Comfort	.1331	.0700	1.9009	.0590	-.0051	.2714
Training	-.0865	.1410	-.6135	.5404	-.3649	.1919
Sector	.0190	.0231	.8212	.4127	-.0266	.0646
Gender	.0125	.1276	.0977	.9223	-.2395	.2644
edu	-.0018	.0590	-.0305	.9757	-.1183	.1147
nation	.0023	.0042	.5476	.5847	-.0060	.0106
age_num	-.0023	.0076	-.3067	.7595	-.0174	.0127

OUTCOME VARIABLE:

NE_mean

Model Summary

R	R-sq	MSE	F	df1	df2	p
.2559	.0655	.8624	1.3075	9.0000	168.0000	.2363

Model

	coeff	se	t	p	LLCI	ULCI
constant	3.4198	.6306	5.4230	.0000	2.1749	4.6648
c_dummy	.1194	.1414	.8445	.3996	-.1598	.3986
Familia	.0411	.0651	.6309	.5290	-.0874	.1696
Comfort	-.0640	.0789	-.8116	.4182	-.2198	.0917
Training	-.3154	.1588	-1.9856	.0487	-.6289	-.0018
Sector	-.0311	.0260	-1.1945	.2340	-.0824	.0203
Gender	-.0394	.1438	-.2739	.7845	-.3232	.2444
edu	-.0614	.0665	-.9235	.3571	-.1926	.0698
nation	.0033	.0047	.7062	.4811	-.0060	.0127
age_num	-.0126	.0086	-1.4639	.1451	-.0296	.0044

OUTCOME VARIABLE:

A_mean

Model Summary

R	R-sq	MSE	F	df1	df2	p
.6774	.4589	.4588	10.6972	13.0000	164.0000	.0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	.4510	.5308	.8495	.3968	-.5972	1.4991
c_dummy	-.0405	.1078	-.3754	.7079	-.2534	.1724
JI_mean	-.0268	.0518	-.5174	.6056	-.1291	.0755
PE_mean	.0323	.0617	.5238	.6011	-.0895	.1542
T_mean	.4327	.0660	6.5513	.0000	.3023	.5631
NE_mean	.0240	.0607	.3961	.6926	-.0958	.1439
Familia	.1949	.0496	3.9275	.0001	.0969	.2929
Comfort	.1833	.0593	3.0922	.0023	.0663	.3004
Training	.0192	.1178	.1628	.8709	-.2134	.2518
Sector	.0060	.0193	.3116	.7557	-.0321	.0442
Gender	.0808	.1052	.7683	.4434	-.1269	.2886
edu	.0075	.0499	.1505	.8806	-.0911	.1061
nation	-.0014	.0035	-.4175	.6769	-.0083	.0054
age_num	.0132	.0064	2.0753	.0395	.0006	.0258

***** DIRECT AND INDIRECT EFFECTS OF X ON Y

Direct effect of X on Y

Effect	se	t	p	LLCI	ULCI
-.0405	.1078	-.3754	.7079	-.2534	.1724

Indirect effect(s) of X on Y:

Effect	BootSE	BootLLCI	BootULCI
TOTAL	.1102	.0699	-.0283 .2458
JI_mean	-.0063	.0139	-.0378 .0211
PE_mean	.0163	.0368	-.0649 .0855
T_mean	.0973	.0574	-.0150 .2150
NE_mean	.0029	.0110	-.0157 .0318

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:

95.0000

Number of bootstrap samples for percentile bootstrap confidence intervals:

5000

----- END MATRIX -----