

An ethical approach to AI aversion in the case of medical devices

ZOLTÁN SOMOSI¹ – NOÉMI HAJDÚ²

During technological development, many researchers have pointed out that various information systems are not used by companies and, in some cases, also by consumers. Technology acceptance models, which have been widely used in recent decades, have been developed to capture this and identify the reasons why it happens. Theoretical and practical findings related to artificial intelligence show a fluctuating trend from time to time, which may even be due to aversion and the resulting resistance. In our literature review, we highlight nine problem areas. These include unemployment, inequality, humanity, anomaly, security, controllability, interest, singularity, and legal disputes. We continued the search based on keywords on the ScienceDirect portal. We worked our way through the articles sorted by relevance, followed by extended research on Generation Z students. The first phase of the survey was based on a sample of 131 people, which showed that respondents thought collaboration between artificial intelligence and humans was possible. The survey was conducted again six months later to determine whether the experience level influenced respondents' perceptions of artificial intelligence.

Keywords: AI, aversion, ethical approach, medical devices.

JEL code: M31.

Introduction

Many statistical features related to artificial intelligence have been shared within the scientific framework, which is essential for a thorough understanding of this field (Friedrich et al. 2020). Various researchers have pointed out that artificial intelligence can reduce the errors resulting from human activity, improve efficiency and speed up work (Mathew 2020). In addition, it can increase productivity in low-income countries, and it has an impact on various sectors such as e-governance, e-finance, e-education and, not coincidentally, the e-health sector (Sachs 2019), which is analysed in this study. However, the problem-solving ability is evident not only in its comprehensiveness but also in its speed, accuracy, and consistency (Choi 2021). Given its many positive attributes, it is not surprising that, with \$25.6 trillion, it can contribute to the growth of the global

¹ PhD student, University of Miskolc, Faculty of Economics, Institute of Marketing and Tourism, zoltan.somosi@uni-miskolc.hu.

² PhD, Associate Professor, University of Miskolc, Faculty of Economics, Institute of Marketing and Tourism, hajdu.noemi@uni-miskolc.hu.

economy (Chui et al. 2023). The PwC report (2024) also shows that, in 2023, almost three quarters of American companies actively used AI-supported tools in their activities. There have also been interesting changes in the consumer sector. According to the report, ChatGPT, which is freely available to general users worldwide, reached 1 million users 5 days after its launch in 2022 (surpassing all previously known platforms), and in 2 months, it reached a database of 100 million users (Statista 2023).

However, it would be wrong to assume that this trend has no downsides, as several researchers have also noted. Mathew (2020) pointed out that artificial intelligence can become dangerous over time by taking over unplanned tasks, which can have negative consequences, such as traffic congestion due to faulty traffic control or the danger posed by autonomous weapons. Webb (2019) outlines the possibility of the devaluation of skilled labour and the loss of jobs. Artificial intelligence can also have a negative impact on scientific work and opportunities for collaboration. In the long term, it can also lead to problems with power and the dissemination of knowledge (Clarke–Whittlestone 2022). We believe that the latter could stem from hallucinations that are common in large-scale artificial intelligence-based language models. Emsley (2023) believes that the hallucinations are a design flaw that reaches the user as false (erroneous) information. According to Leiser et al. (2023), Large Language Models (LLMs) such as ChatGPT can generate information that contradicts factual knowledge and contains false statements. He also addressed the fact that this information can influence users and cause them to change their attitudes and general beliefs.

As a complement to the technology acceptance model (UTAUT) developed by Venkatesh et al. (2012), Jain, Garg and Khara (2022) investigated the aversion to AI after finding that AI was becoming more prevalent in organisations. Using the partial least squares (PLS) method, they showed in their study that aversion is a moderating factor for the adoption of the technology. To assess the influence of ethical issues on reluctance, in the current study, we conducted a questionnaire survey among Generation Z students and, additionally, we analysed the ethical issues raised by various authors using a keyword search on ScienceDirect. Other moderating variables, including age, experience, and gender, are included in the UTAUT base model. According to Venkatesh et al. (2012) and Dwivedi et al. (2017), a significant proportion of researchers did not include these moderating variables in their studies, meaning that only the base of the model was considered.

Literature review

The UTAUT model is one of the most widely used models in artificial intelligence research. To prove this, we conducted a query on the ScienceDirect portal, which collects scientific papers, using the following keywords to narrow down the field of artificial intelligence: “artificial intelligence” OR “AI”. In relation to technology acceptance, we searched for: “TAM” OR “Technology Acceptance Model” OR “TAM 2” OR “Technology Acceptance Model 2” OR “TAM 3” OR “Technology Acceptance Model 3” OR “UTAUT” OR “Unified theory of acceptance and use of technology” OR “UTAUT 2” OR “Unified theory of acceptance and use of technology 2”. When collecting the statistical data, additional models were excluded with the expression AND NOT, so factual data were collected for some theories. The comprehensive review included a total of 12,540 studies that referred to the original version of the Technology Acceptance Model (TAM). In contrast, the refined second iteration listed 975 results. Of the TAM 3 models from 2008, 753 studies were included in the platform. The UTAUT model was used in 2,341 studies, while its refined version was used only 204 times. When analysing the adoption models in different domains, the aggregated data showed a clear shift when focusing exclusively on the artificial intelligence domain. In this context, the TAM 1 type model included 1,447 studies, TAM 2 only 5 studies, the theory of TAM 3 included 16 studies published on the ScienceDirect platform, while UTAUT 1 included 464 studies and UTAUT 2 included 21 studies.

Venkatesh et al. (2012) found that the UTAUT model has an explanatory power of 56% in terms of intention to use, while it reaches 40% in the case of actual use. As Keszey and Zsukk (2017) emphasise, the advantages of this model lie in the combination of eight previous theories and the consideration of important individual factors. Its application can be interpreted primarily in a workplace environment, as it was developed specifically for this purpose. Although the UTAUT model is already widely used due to its simplicity and greater explanatory power compared to the eight models on which it is based, even when interpreted for consumers, Dwivedi et al. (2017) pointed out certain limitations. Similarly, Venkatesh et al. (2012) felt that the researchers omitted a significant portion of the moderating variables, so their study was based on only a fraction of the model.

Dwivedi et al. (2019) claim that support and assistance in creating framework conditions can promote a positive attitude towards a technology and

thus reduce resistance to artificial intelligence. Dwivedi et al. (2019) found a positive correlation between social support and the reduction of aversion, like favourable framework conditions. Van Esch et al. (2019) claim that the fear of missing out (FOMO) is a social effect of technology use. In their study, Dwivedi et al. (2017) showed that a person's attitude towards adopting a new technology is positively influenced by its perceived usefulness. In other words, the perceived difficulty in using a technology or system is the expected cost (Venkatesh et al. 2003). As Rouidi et al. (2022) found, both the expected effort and the perceived ease of use of the Technology Adoption Model (TAM 1) can be assessed with the same predictors. These consist of the following factors: self-efficacy, pre-existing knowledge, interoperability with other systems that have been or can be developed, etc. It follows that all test factors are influenced by an aversion to AI.

The moderating effects of gender, age, volunteering, and experience on expected performance, expected effort, social impact, and facilitating conditions were demonstrated by Venkatesh et al. (2003) in a study focusing on the development and validation of the UTAUT. Regarding experience, Venkatesh et al. (2003) found the following results:

- The expected effort is higher for individuals with less experience,
- The social impact is stronger for people with less experience,
- The effect of favourable conditions increases with experience.

We believe that the level of experience found as a moderating variable in the UTAUT model influences the aversion to artificial intelligence, as the benefits offered by AI can compensate for any aversion that may occur. However, we can also talk about aversion from other angles. Political conservatism is related to distrust of artificial intelligence (Castelo–Ward 2021) and aversion to algorithms, i.e. not using algorithms. It also impairs the interaction between humans and AI and thus reduces the benefits of AI (Mahmud et al. 2022). Humans judge unethical decisions made by AI as less ethical than human decisions, leading to gaps in moral judgements, punishments, and behavioural responses to AI (Zhang et al. 2023).

According to PriceWaterhouseCoopers' 2022 Business Survey on AI and Analytics, the true value of these technologies lies in their ability to improve the customer experience (40%), increase efficiency (44%), make better decisions (41%) and develop new goods and services (40%). Conversely, artificial

intelligence is expected to lead to the loss of 85 million jobs (Jovanovic 2022). In addition to the human element, the World Economic Forum has drawn attention to other topics that fall under the following categories:

- Unemployment: How will artificial intelligence change the way labour is released? Will employment increase? What is the composition of the labour force?
 - Inequality: What is the best way to distribute the wealth generated by AI? Will it change the equality of people? Will it promote recognition?
 - Humanity: How will machines influence people's interactions and behaviours? Is it possible for humans and machines to coexist harmoniously?
 - Anomaly: How can we prevent errors caused by AI?
 - Security: It means protecting and reinforcing the artificial intelligence system.
 - Controllability: Is it possible for humans to control the creation and application of artificial intelligence?
 - Interest: Since humans design artificial intelligence systems, is complete independence really a guarantee? What role do the various interest groups play?
 - Singularity: Will the system be able to wake up?
 - Legal disputes: Who will be liable for the actions of AI and in what form?
- (World Economic Forum 2016).

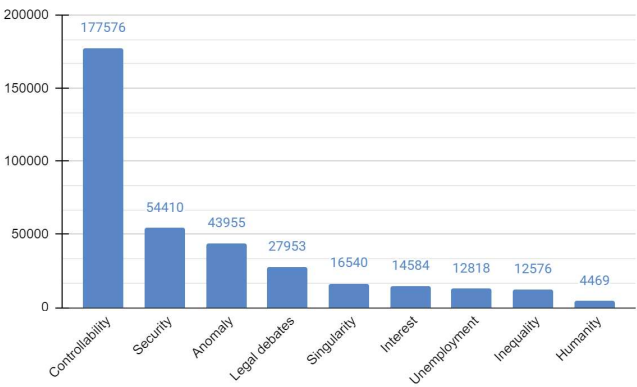
Nine relevant categories were identified and searched for on ScienceDirect, using appropriate keywords. No additional restrictions were placed on the listing of studies, such as year of publication, type, title, topic, or accessibility. The relevance score of ScienceDirect was used to evaluate the listed publications. The first 25 results were always considered, and the articles were then carefully reduced to the most relevant ones after analysing the title and abstract. If there were more than three relevant articles, the first three were included in the ranking to save space. The scope and keywords that formed the basis of the search are listed in Table 1.

All keywords for the combined groups contained the English term for artificial intelligence in addition to another term that best described the group. New keyword ideas were considered according to the suggestion of ChatGPT. The keywords for the different categories on the ScienceDirect portal resulted in a different number of hits (Figure 1).

Table 1. Definition of categories and keywords

Category	Keywords
Unemployment	artificial intelligence employment
Inequality	artificial intelligence equality
Humanity	artificial intelligence humanity
Anomaly	artificial intelligence safety
Security	artificial intelligence security
Controllability	artificial intelligence controllability
Interest	artificial intelligence interest
Singularity	artificial intelligence singularity
Legal disputes	artificial intelligence legal disputes

Source: Own editing



Source: Own editing

Figure 1. Number of studies in AI social impact categories (pcs)

The summary of the search for the various keywords shows that the topic of controllability has the most publications (177,576 articles), followed by a significantly smaller number of topics related to security (54,410 articles). Subsequently, the topics of disorder, legal issues and singularity show a decreasing trend. We find it interesting that the topic of humanity and artificial intelligence is one of the least researched when looking at the information found on the portal.

To make it easier to understand, we would like to explain these categories using a practical example from a specific industry. For example, medical technology is an excellent example regarding the many applications of AI.

- Medical technology has always played a leading role in innovation.
- AI can contribute to the development of new technologies that offer new approaches and solutions for the diagnosis and treatment of diseases.
 - With the help of medical applications, AI can contribute to early diagnoses, more effective treatments, and surgical interventions, which can ultimately lead to saving lives or improving the quality of life of those living with diseases.
 - In the field of healthcare, there are extremely complex tasks with a large amount of data. The application of medical technology usually requires a large amount of data and information about diseases, diagnoses and treatments that can be effectively managed by AI.
 - The application of AI can contribute to a more efficient use of healthcare resources, such as diagnostic equipment, healthcare professionals and time.
 - Developments in medical technology have a direct impact on people's health and wellbeing, which has important societal and social implications.

The authors' considerations were listed under the heading "Artificial intelligence employment" and then categorised according to their relevance. According to the study conducted by Yan-ping and Ai-qin (2022), there is a positive correlation between AI and the employment of young people. To achieve correlation, they make the following demands: (1) the state must play an important role in containing the epidemic and creating jobs; (2) there is great scope for increased cooperation between colleges and universities; and (3) the solution must be understood in terms of business innovation. In line with previous studies, Yang (2022) investigated the effects of artificial intelligence (AI) on productivity and employment in Taiwan. She found that, while AI increases employment, it also changes the composition of the labour force and favours workers with higher levels of education. Universities, companies, and the government were also cited as possible support to develop AI technologies.

In the case of medical technology, AI is mostly used as a supporting component, e.g. for analysing imaging procedures and diagnostic interventions. In this way, it creates new areas of knowledge and expands the work of specialists rather than replacing them. The introduction of AI may lead to the emergence of new specialisms, such as experts in the development, maintenance, and monitoring of AI systems. The medical technology workforce is changing; IT and AI specialists could play an important role alongside traditional healthcare professionals.

Listed under the keyword “artificial intelligence equality” and then sorted by relevance, the authors’ views were as follows. Pierson et al. (2021) examined the role of artificial intelligence in creating racial equality in the United States and showed that AI can promote equality. Part of their conclusion includes naming companies and research institutes as the solvers. O’Connor and Booth (2022) analysed the promotion of equality using AI in healthcare. They concluded that the education and training of nurses is key to improving and equalising care by using the potential that AI already has. The first 25 results for this keyword did not include any research on equality.

Medical technology applications can reduce health inequalities by improving access to healthcare. By enabling faster and more accurate diagnoses, they can improve equity in patient care. Artificial intelligence algorithms have the potential to improve health awareness by assisting in the early detection of disease and signposting preventative health measures.

Listed under the keywords “artificial intelligence humanity, artificial intelligence society” and then sorted by relevance, the following conclusions can be drawn. The conflict between human freedom and intelligent machines was analysed by Coombs et al. in 2021. They concluded that humans have a duty to guide artificial intelligence, promote critical thinking and empower people to make decisions. Instead of an institutional perspective, they take a problem-oriented approach. In their analysis of the effects of robots and artificial intelligence on the economy and society, Haenlein and Kaplan (2021) argue that international law is necessary to ensure the coexistence of humans and AI. Lamotte (2020) examined the relationship between society, AI and enlightenment and showed how AI companies are already having a major impact on people’s lives. Lamotte (2020) claimed that “Facebook knows who we are, Amazon knows what we want, and Google knows what we think” (Lamotte 2020. 17431), in contrast to the findings of Coombs et al. (2021), Haenlein and Kaplan (2021) and other researchers who concluded that global regulation is necessary. There may be disputes about whether it is better to use less AI or to regulate these giant corporations. It has been proven that artificial intelligence improves employment.

Medical technology applications offer the opportunity to improve the interaction between humans and machines, e.g. between patients and healthcare providers, which can facilitate the expression of compassion and empathy. The

supportive position of AI in medical technology enables the coexistence of human expertise and automation, with humans still having the final say in diagnosis and treatment.

After sorting them by relevance, the articles found on the term “artificial intelligence safety” are presented below. The safety of AI from the perspective of those affected was assessed by Sujan et al. in 2022. According to Johnson’s (2022) exploration of the basis for desirable and safe behaviour, metacognition could be the key to safety. It is the responsibility of the engineers and developers of AI systems to integrate and utilise it.

As medical technology works with sensitive data, it is crucial to ensure the security of healthcare data and the integrity of AI systems. Access restrictions, data encryption and other security measures are examples of enhanced protection to ensure that only authorised individuals can view sensitive patient data. Healthcare systems must safely integrate AI applications. These applications must work with the current healthcare infrastructure and must not jeopardise the overall stability and security of the system.

The following articles were highlighted after being indexed with the term “artificial intelligence stakeholders” and then sorted by relevance. Miller (2022) describes the function of stakeholders in artificial intelligence initiatives and points out the potential for AI systems to cause harm to people, communities, and the environment. He concludes that moral, ethical, and sustainable development requires a comprehensive strategy.

In the face of artificial intelligence, Garrett and Young (2022) argue that stakeholders’ participation in HIV research is critical. Like Miller (2022), they assign responsibility to stakeholders, saying that policy makers are responsible for protecting people’s rights and safety, while developers, investors and doctors are responsible for research.

To ensure that the use of AI in medical technology benefits everyone, a wide range of stakeholders, including patients, medical experts and developers, are actively involved in the decision-making process. Finding a middle ground between the various stakeholders is essential for the long-term and effective application of AI in medical technology. Promoting an open dialogue and considering many points of view can support the development, acceptance, and effective application of artificial intelligence. After analysing the various stakeholders, we identified the following benefits:

- The use of AI to raise standards in healthcare and improve accessibility is very attractive to patients.
- To ensure that AI systems fulfil clinical requirements, doctors and specialists are involved in their development and use.
- AI can streamline research procedures and help in the discovery and development of new medicines.
- The use of AI has the potential to increase system performance and improve standards of patient care.

Below you will find the most important articles found when entering the keyword “artificial intelligence rights”, which were then sorted by relevance. To uphold ethical principles and fundamental rights, Stahl et al. (2022) have developed a set of 12 criteria that address the areas that the European Artificial Intelligence Agency hopes to control. These include a map of places to be banned and a request to the European Parliament. Felice et al. (2022) have analysed how AI affects ethics, rights and lives. They believe that it is necessary to educate the future generation about AI. They claim that AI is a tool, not a goal. To ensure that fundamental rights are upheld, they advise focusing on legislation, political regulation, and the promotion of education. Beaumier (2022) states in his book review that there is currently no clear definition of whether a non-human being can have rights or not. While some support this holistic approach, others are against it. In his opinion, it is important to guarantee the rights and justice of robots.

Legal issues are raised by the question of who is responsible for errors, misdiagnoses or other problems resulting from the use of AI. It is important to clarify who is responsible: the developers, the doctor, the institution, or other parties involved. Patients should be involved in this decision-making process and informed about how AI may influence the diagnosis and course of therapy. To ensure patient safety, data confidentiality and ethical and legal practice, legal considerations for medical applications are crucial in the development and use of AI.

The “controllability of artificial intelligence” is the key to taking ethical considerations into account. The ability to control decisions is seen as a basic need (Morley et al. 2019). Lack of controllability can lead to unpredictable outcomes, which raises ethical and safety issues (Musiolik 2021). One of the biggest challenges in achieving controllability is the inadequacy of the data needed to model the system (Banno et al. 2021). According to Yampolskiy (2022), artificial

intelligence has uncontrollable and significant consequences for the future of humanity due to its uncontrollability.

In connection with the topic of “artificial intelligence singularity”, we first clarify the implications of the concept for the topic of aversion with a conceptual definition. Singularity is based on the notion that artificial intelligence and associated technologies will develop at an exponential rate, leading to a tipping point where machine intelligence will surpass human thought (Goode 2018). Using an ethical framework, Schrader–Ghosh (2018) tried to focus on the protection of human rights and wellbeing rather than on superior machine intelligence.

The “anomaly of artificial intelligence” was the central element of data mining. Attempts were made to map it as efficiently as possible, for which various methods were developed. While there is extensive work on anomaly detection, the literature does not address the explanation of anomalies, so a comprehensive review is needed in this area (Tchaghe et al. 2022).

In the literature review, we pointed out that the UTAUT model was one of the most appropriate models to assess the acceptance of artificial intelligence, whether in an employee or a consumer setting, as Keszey and Zsukk (2017) argue that it has the fewest limitations among technology acceptance models. Venkatesh (2003) also created four moderating variables to build the model. Although it is often omitted (Dwivedi et al. 2017), the experience variable piqued our interest as a moderating factor. Jain et al. (2022) also added aversion to artificial intelligence to the base model of UTAUT, an issue that is more important today than ever.

As a result, and after processing the literature presented, we believe that the topic of aversion can be effectively conceptualised by mapping the ethical topics, and we suggest the use of a possible scale.

Methodology

The UTAUT model is a highly robust and comprehensive framework for testing technology adoption, particularly in the field of artificial intelligence. Its ability to integrate a wide range of factors and moderating variables makes it particularly suitable for the employee and consumer context. The application of the model goes beyond purely theoretical studies and offers practical insights into improving AI adoption and addressing the ethical and societal issues involved.

Given the widespread use of UTAUT in AI research and its proven explanatory power, it is recommended that this model also be used in future

studies. In addition, further exploration of ethical considerations and the inclusion of variables such as experience and counter-perception may contribute to a better understanding of AI adoption and use.

Based on the literature review, the main questions of our research were outlined, which include the UTAUT model, AI aversion, and the moderating variable of experience, which is part of the UTAUT model.

Research question 1: How open-minded Generation Z is about AI, and how do they respond to its use and to questions about humans and their rights?

Research question 2: Do the benefits of AI outweigh the aversion that may arise according to the findings of de Bruijn et al. (2013)?

Research question 3: Does a higher level of experience with artificial intelligence reduce aversion to its use?

To answer these questions, during the period from September 5 to October 5, 2023, we conducted a survey among Generation Z students. The sampling was based on the snowball principle and can therefore be categorised as a random sample. The questions were created partly with a Likert scale and partly with a semantic difference scale. The question “Do you feel uncomfortable interacting with and using AI?” is based on a Likert scale. The question “Can there be cooperation between humans and machines?” is based on a semantic scale. The questionnaire was available online, and we used the Google Form.

To investigate the effects of the moderating variable *experience* on dislike and acceptance, we repeated the survey six months later, from March 1 to April 1, 2024, with the same target group.

Results

A total of 132 responses were received to the first questionnaire. The average age of the respondents was between 21 and 25, with the youngest respondent being 15 and the oldest 30 years old. The age difference between them was 15 years. As Generation Z was born between 1995 and 2009, we excluded respondents who did not fall within this age range and adjusted the data based on age. Of the 131 respondents in the adjusted data table, 91 identified as female, 37 identified as male and 3 did not wish to disclose their gender identity.

For each question, the following results were determined on a Likert scale of 1 to 5, with 1 standing for complete disagreement with AI and 5 for complete agreement with AI.

Respondents indicated that they had no concerns about the use of and interaction with AI. On a 5-point scale, the most common score for this type of dislike was 2, and the average was 2.51. The question of whether the use of AI could have a negative impact on our lives was one of the concerns associated with its use. Most respondents disagreed, but the average score (2.85) was close to the mean. When it comes to maintaining a corporate culture, things are different. Most respondents felt that the use of AI would change the current state of culture. Of those who responded to the survey, 55 felt that AI was getting out of hand and only 12 per cent felt the opposite.

Table 2 summarises the results of the impact assessment for the categories mentioned in the literature review. The average scores and directions are shown in the second and third columns.

Table 2. Results of the first survey on the perceptions of artificial intelligence

Category	Average	Directions
Unemployment	3.1	Job losses
Inequality	3.6	Increasing inequality
Anomaly	2.4	We are not prepared
Security	3.6	Not reliable
Controllability	3.6	Uncontrollable
Contact	3.5	Not independent
Singularity	4.1	Unable to wake up
Legal disputes	3.4	Unknown culprits
Humanity	4.8	Cooperation can be established

Source: Own editing

According to respondents, AI is to blame for increasing inequality and job losses, and we are ill-prepared for it. Artificial intelligence lacks independence, controllability, and reliability. However, it can promote cooperation between AI and humans, and it lacks self-confidence, even if most respondents do not know the identity of those responsible.

The survey was repeated six months later, using the same questionnaire. This time, 131 people responded (the previously excluded participant was not interviewed).

Respondents expressed concerns about the use of and interaction with artificial intelligence. Again, the dislike was often rated with a score of 2, although

the average fell to 2.37. The answer to the question of whether AI would make our lives worse was the same as in the previous questionnaire: no. There is no recognisable change in culture or loss of control; opinions remain the same.

Table 3. Results of the second survey on the perceptions of artificial intelligence

Category	Average	Directions
Unemployment	3.3	Job losses
Inequality	3.4	Increasing inequality
Anomaly	2.6	We are not prepared
Security	3.6	Not reliable
Controllability	3.2	Uncontrollable
Contact	3.1	Not independent
Singularity	4.3	Unable to wake up
Legal disputes	3.4	Unknown culprits
Humanity	4.9	Cooperation can be established

Source: Own editing

The average scores for unemployment, inequality, disorder, controllability, interest, uniqueness, and humanity differed slightly, according to the results of the five-point scale (Table 3). There were no changes in the areas of litigation and security. Six months later, the study found that, overall, attitudes had not changed significantly. In line with the perception of chaos, respondents saw unemployment as a greater threat. There was also a negative shift in the questions on interests, which could explain the more pessimistic views on controllability.

Conclusions

The industrial revolution has allowed AI to reach a new level of development, and after the aggressive push of Industry 4.0, the sustainable form of Industry 5.0 has also influenced the expansion of the field (Alexa et al. 2022). This has led to several ethical dilemmas that have been raised by various authors. According to the results of our study among Generation Z students, respondents show aversion to AI, which may discourage them from adopting the technology. The level of experience as a moderating variable increased the previously positively rated items and decreased the negatively rated items. We expected this increase to occur within a six-month period due to the increasing use of one of the more accessible AI-enabled software applications. We need to extend our model to include the

entire UTAUT base as well as the moderating factors to explain the changes if we are to demonstrate a discernible change. We would like to answer the research questions as follows:

Research question 1: Generation Z will be open-minded about AI, will not resist its use and will respond positively to a range of questions about humans and their rights. While students in the sample are happy to use artificial intelligence, they oppose it in certain areas and did not give positive responses to ethical questions about humanity.

Research question 2: The benefits of AI outweigh the aversion that may arise, as de Bruijn et al. (2013) found. The Generation Z university students included in the sample felt that the benefits of AI only outweighed the aversion in some cases when evaluated in relation to ethical issues. This also applies to the issue of collaboration.

Research question 3: A higher level of experience reduces aversion to artificial intelligence. A higher level of experience in the areas of equality, controllability and independence reduces aversion when we approach the topic of ethical considerations.

The reviewed literature shows that AI is being used in several areas, including employment, inequality, chaos, security, controllability, participation, singularity, legal issues, and humanity. Several authors have raised concerns about resistance to the technology (Jain et al. 2022). A sample of 132 Generation Z students at the University of Miskolc participated in an online survey; after adjustment, 131 valid responses were left. The results of the survey show that, in some cases, Generation Z university students hold the same views as those described in the literature. However, the results also suggest that respondents are not concerned that AI will develop its own consciousness, and they believe that technology and humans can work together. When we conducted our study again six months later with the same participants, the factors that had received a low rating were ranked even lower and the factors that had received a high rating were ranked even higher. The likelihood of AI being self-aware was even lower and the likelihood of it working together with humans was higher.

This study, like all others, has several weaknesses. First, the topics analysed may not cover all important ethical or societal concerns related to artificial intelligence, as the literature review is not methodical. Furthermore, there is no evidence that the topics we selected are among the best researched. For this

reason, the purpose of our field survey is to validate theories that will serve as a basis for future research, not to support or refute the theories they contain. To summarise, the following factors have to be considered when interpreting the results of the study:

- The literature review is not systematic, and the topics covered do not necessarily reflect the quantitative research directions of theorists.
- The research and research results have been tested on a small sample that is not representative, so conclusions can be drawn about the sample.
- The study uses basic statistical analyses.

Our aim is to conduct a more comprehensive and representative survey and to validate the results through a systematic literature review.

References

- Alexa, L.–Pîslaru, M.–Avasilcăi, S. 2022. From Industry 4.0 to Industry 5.0 – An Overview of European Union Enterprises, In: Draghici, A.–Ivascu, L. *Sustainability and Innovation in Manufacturing Enterprises*. Singapore: Springer, 221–231.
- Banno, I.–Azuma, S. I.–Ryo Ariizumi, Asai, T.–Imura, J.-I. 2021. Data-driven Estimation and Maximization of Controllability Gramians. *60th IEEE Conference on Decision and Control*.
- Beaumier, G. 2022. Book review. *Earth System Governance* 11, 100129.
- Castelo, N.–Ward, A. F. 2021. Conservatism predicts aversion to consequential Artificial Intelligence. *PLOS One* 16(12), e0261467.
- Choi, C. Q. 2021. 7 Revealing Ways AIs Fail: Neural Networks can be Disastrously Brittle, Forgetful, and Surprisingly Bad at Math. *IEEE Spectrum* 58(10), 42–47.
- Chui, M.–Hazan, E.–Roberts, R.–Singla, A.–Smaje, K.–Sukharevsky, A.–Yee, L.–Zemmel, R. 2023. *The economic potential of generative AI: The next productivity frontier*. <https://www.mckinsey.com/capabilities/mckinsey-digital/our-insights/the-economic-potential-of-generative-ai-the-next-productivity-frontier#introduction>, downloaded: 08.05.2024.
- Clarke, S.–Whittlestone, J. 2022. A Survey of the Potential Long-term Impacts of AI. *Proceedings of the 2022 AAAI/ACM Conference on AI, Ethics, and Society*, 192–202.
- Coombs, C.–Stacey, P.–Kawalek, P.–Simeonova, B.–Becker, J.–Bergener, K.–Carvalho, J. Á.–Fantinato, M.–Garmann-Johnsen, N. F.–Grimme, C.–Stein, A.–Trautmann, H. 2021. What is it about humanity that we can't give away to intelligent machines? A European perspective. *International Journal of Information Management* 58, 102311.
- De Bruijn, G. J.–Gardner, B.–van Osch, L.–Sniehotta, F. F. 2013. Predicting Automaticity in Exercise Behaviour: The Role of Perceived Behavioural Control, Affect, Intention, Action Planning, and Behaviour. *International Journal of Behavioral Medicine* 21(5), 767–774.

Dwivedi, Y. K.–Rana, N. P.–Janssen, M.–Lal, B.–Williams, M. D.– Clement, M. 2017. An empirical validation of a unified model of electronic government adoption. *Government Information Quarterly* 342, 211–230.

Dwivedi, Y. K.–Rana, N. P.–Jeyaraj, A.–Clement, M.–Williams, M. D. 2019. Re-examining the Unified Theory of Acceptance and Use of Technology (UTAUT): Towards a Revised Theoretical Model. *Information Systems Frontiers* 21(3), 719–734.

Emsley, R. 2023. ChatGPT: these are not hallucinations – they’re fabrications and falsifications. *Schizophrenia* 9(1), 52.

Felice, F. D.–Petrillo, A.–Luca, C. D.–Baffo, I. 2022. Artificial Intelligence or Augmented Intelligence? Impact on our lives, rights and ethics. *Procedia Computer Science* 200, 1846–1856.

Friedrich, S.–Antes, G.–Behr, S.–Binder, H.–Brannath, W.–Dumpert, F.–Ickstadt, K.–Kestler, H.–Lederer, J.–Leitgöb, H.–Pauly, M.–Steland, A.–Wilhelm, A.–Friede, T. 2020. Is there a role for statistics in artificial intelligence? *Advances in Data Analysis and Classification* 16(4), 823–846.

Garett, R.–Young, S. D. 2022. The importance of diverse key stakeholders in deciding the role of artificial intelligence for HIV research and policy. *Health Policy and Technology* 11(1), 100599.

Goode, L. 2018. Life, but not as we know it: A.I. and the popular imagination. *Culture Unbound* 102, 185–207.

Haenlein, M.–Kaplan, A. 2021. Artificial intelligence and robotics: shaking up the business world and society at large. *Journal of Business Research* 124, 405–407.

Jain, R.–Garg, N.–Khera, S. N. 2022. Adoption of AI-Enabled Tools in Social Development Organizations in India: An Extension of UTAUT Model. *Frontiers in Psychology* 13, 893691.

Johnson, B. 2022. Metacognition for artificial intelligence system safety – An approach to safe and desired behavior. *Safety Science* 151, 105743.

Jovanovic, B. 2022. *55 Fascinating AI Statistics and Trends for 2022*. <https://dataprot.net/statistics/ai-statistics/>, downloaded: 08.02.2024.

Keszey, T.–Zsuk, J. 2017. Az új technológiák fogyasztói elfogadása. A magyar és nemzetközi szakirodalom áttekintése és kritikai értékelése. *Vezetéstudomány – Budapest Management Review* 48(10), 38–47.

Lamotte, M. 2020. Enlightenment, Artificial Intelligence and Society. *IFAC-PapersOnLine* 532, 17427–17432.

Leiser, F.–Eckhardt, S.–Knaeble, M.–Maedche, A.–Schwabe, G.–Sunyaev, A. 2023. From ChatGPT to FactGPT: A participatory design study to mitigate the effects of large language model hallucinations on users. *Proceedings of Mensch und Computer* 81–90.

Mahmud, H.–Islam, A. K. M. Najmul–Ahmed, S. I.–Smolander, K. 2022. What influences algorithmic decision-making? A systematic literature review on algorithm aversion. *Technological Forecasting & Social Change* 175, 121390–121390.

- Mathew, A. 2020. The Peril of Artificial Intelligence. *2020 Fourth International Conference on Inventive Systems and Control (ICISC)*, 919–924.
- Miller, G. J. 2022. Stakeholder roles in artificial intelligence projects. *Project Leadership and Society* 3, 100068.
- Morley, J.–Floridi, L.–Kinsey, L.–Elhalal, A. 2019. From What to How: An Initial Review of Publicly Available AI Ethics Tools, Methods and Research to Translate Principles into Practices. *Science and Engineering Ethics* 26(4), 2141–2168.
- Musiolik, G. 2021. Predictability of AI Decisions. In: Musiolik, T. H.–Cheok, A. D. *Analyzing Future Applications of AI, Sensors, and Robotics in Society*. Hershey, PA: IGI Global Scientific Publishing, 17–28.
- O'Connor, S.–Booth, R. G. 2022. Algorithmic bias in health care: opportunities for nurses to improve equality in the age of artificial intelligence. *Nursing Outlook* 70(6), 780–782.
- Pierson, E.–Cutler, D. M.–Leskovec, J.–Mullainathan, S.–Obermeyer, Z. 2021. An algorithmic approach to reducing unexplained pain disparities in underserved populations. *Nature Medicine* 27(1), 136–140.
- PriceWaterhouseCoopers 2022. *Become a leader in AI and Analytics*. <https://www.pwc.com/us/en/tech-effect/ai-analytics/ai-business-survey.html>, downloaded: 12.12.2022.
- Rouidi, M.–Elouadi, E. A.–Hamdoun, A.–Choujtani, K.–Chati, A. 2022. TAM-UTAUT and the acceptance of remote healthcare technologies by healthcare professionals: a systematic review. *Informatics in Medicine Unlocked* 32, 1–14.
- Sachs, J. D. 2019. Some Brief Reflections on Digital Technologies and Economic Development. *Ethics & International Affairs* 33(02), 159–167.
- Schrader, D. E.–Ghosh, D. 2018. Proactively Protecting Against the Singularity: Ethical Decision Making in AI. *IEEE Security & Privacy* 16(3), 56–63.
- Stahl, B. C.–Rodrigues, R.–Santiago, N.–Macnish, K. 2022. A European Agency for Artificial Intelligence: protecting fundamental rights and ethical values. *Computer Law & Security Review* 45, 105661.
- Statista 2023. *Topic: AI-powered online search*. <https://www.statista.com/topics/10825/ai-powered-online-search/#topicOverview>, downloaded: 03.05.2024.
- Sujan, M. A.–White, S.–Habli, I.–Reynolds, N. 2022. Stakeholder perceptions of the safety and assurance of artificial intelligence in healthcare. *Safety Science* 155, 105870.
- Tchaghe, Y. V.–Smits, G.–Pivert, O. 2022. Anomaly explanation: A review. *Data & Knowledge Engineering* 137, 101946–101946.
- Van Esch, P.–Black, J. S.–Férolie, J. 2019. Marketing AI recruitment: The next phase in job application and selection. *Computers in Human Behavior* 90, 215–222.
- Venkatesh, V.–Thong, J. Y. L.–Xu, X. 2012. Consumer Acceptance and Use of Information Technology: Extending the Unified Theory of Acceptance and Use of Technology. *MIS Quarterly* 36(1), 157–178.

Venkatesh, V.–Morris, M. G.–Davis, G. B.–Davis, F. D. 2003. User Acceptance of Information Technology: Toward a Unified View. *MIS Quarterly* 27, 3.

Webb, M. 2019. The Impact of Artificial Intelligence on the Labor Market. *SSRN Electronic Journal*, 3482150.

World Economic Forum 2016. *Top 9 ethical issues in artificial intelligence*. <https://www.weforum.org/agenda/2016/10/top-10-ethical-issues-in-artificial-intelligence/>, downloaded: 03.12.2022.

Yampolskiy, R. V. 2022. On the Controllability of Artificial Intelligence: An Analysis of Limitations. *Journal of Cyber Security and Mobility*, 321–404.

Yan-ping, L.–Ai-qin, Q. 2022. Replace or create: Analysis of the Relationship between the Artificial Intelligence and Youth Employment in Post Epidemic Era. *Procedia Computer Science* 202, 217–222.

Yang, C.–H. 2022. How Artificial Intelligence Technology Affects Productivity and Employment: Firm-level Evidence from Taiwan. *Research Policy* 51(6), 104536.

Zhang, Y.–Wu, J.–Yu, F.–Xu, L. 2023. Moral Judgments of Human vs. AI Agents in Moral Dilemmas. *Behavioral Sciences* 13(2), 181.
