

Review

Not peer-reviewed version

Assessing the Societal Risks of AI's Rapid Advancement: Innovation or Threat to Safety?

[Oma Zaib](#)^{*} and Abbas Raza Ali

Posted Date: 17 September 2024

doi: 10.20944/preprints202409.1217.v1

Keywords: Artificial Intelligence (AI); Deep Fakes; Autonomous Vehicles (AV); Societal Safety; Cyber Security; Regulatory Frameworks



Preprints.org is a free multidiscipline platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This is an open access article distributed under the Creative Commons Attribution License which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

Review

Assessing the Societal Risks of AI's Rapid Advancement: Innovation or Threat to Safety?

Oma Zaib ^{1,*} and Abbas Raza Ali ²

¹ Beaconhouse College Programme (BCP), Beaconhouse, Liberty, Lahore, 54770, Pakistan.

² Artificial Intelligence, PricewaterhouseCoopers, 2 Glass Wharf, Temple Quay, Bristol, United Kingdom; abbas@gmail.com

* Correspondence: oma123zaib@gmail.com

Abstract: The rapid acceleration of artificial intelligence (AI) has sparked extensive debate about its impact on the safety and security of society. While AI has shown transformative potential, particularly in areas such as healthcare and independent structures where it complements productivity and minimizes human error, accuracy and protection. However, the misuse of artificial intelligence poses real risks, along with threats such as deep counterfeiting, computer cyber-attacks and the manipulation of social behavior through disinformation. These risks vary from privacy and cyber security breaches to moral dilemmas around liability. The unchecked improvement of AI is expected to threaten the social order, so the relationship between innovation and governance is essential so that AI-related threats do not overshadow its benefits. This article offers a detailed assessment and review of the literature that studies the potential dangers that AI poses to society. The cumulative effects underscore the urgent need for comprehensive regulatory frameworks to mitigate these threats.

Keywords: artificial intelligence (AI); deep fakes; autonomous vehicles (AV); societal safety; cyber security; regulatory frameworks

1. Introduction

The most important conceptual developments are the dazzling advances in Artificial Intelligence (AI) which has aroused concerns regarding its impact on societal safety and privacy. Expected to become a \$190 billion sector by the year 2025, AI (Artificial Intelligence) patterns human intelligence by machines and has experienced significant growth. This use of AI technology has spurred immense growth in spheres like health-care and transportation. Advancements in AI-powered technologies like autonomous vehicles and complex diagnostic tools that bolster efficiency and safety stands as a testimony to the benefits Earth could gain from innovation. Autonomous vehicles are said to have fewer accidents than human-driven cars, while AI driven diagnostic systems increases the accuracy rate in disease detection. Despite all these advancements, the emergence of AI has caused for alarm. The technology's ability to create and spread fake news presents a major barrier for accurate-decision making as well. Scenarios like manipulated social media campaigns or deepfake videos show how people can also pervert AI to manipulate opinions and spread falsehoods. Secondly, there has been the increasingly known hurdle of privacy as AI tools may allow for data harvesting and destruction without proper consent. Given this rapid evolution of AI, it is important to analyze whether the acceleration will be for the betterment or if new risks and problems emerge as a result. This dual nature of AI — beneficial and malevolent at the same time, elicits a critical discussion around how its benefits can be reaped while minimizing accompanying costs. This balance is crucial to ensure that the solutions developed by AI are advancing technologies for good and safe applications, while also not losing sight of ethical concerns.

2. Literature Review

AI, or Artificial Intelligence, refers back to the simulation of human intelligence processes via machines, specifically pc structures. The latest improvement of AI has sparked a debate in specific sectors of life. Its speedy pace is envisioned to attain a hundred ninety. Sixty one billion dollars with the aid of 2025, at a compound annual boom charge of 36.62 percent [1]. AI has been broadly used for the reason that beyond few many years for solving troubles in numerous fields of lifestyles [2]. Some argue that AI has been an irreplaceable asset in the medical field, specially since AI-powered microscopes that can automatically detect malaria parasites in blood samples had been developed by means of researchers from Sub-Saharan Africa [3]. Moreover, there is additionally a robust support for AI due to the speedy lower of car accidents because of the usage of self sustaining motors. The chance of a crash main to an injury is notably lower in AVs (11.7%) than in human driven (30.3%) [4]. However, like each development, there are a few downsides as nicely. One of the main worries being the social manipulation because of the increase in AI, which impairs an individual's ability to make decisions due to the affect of fake records [5]. During the Philippines' election in 2022, Ferdinand Marcos Jr., mobilized a TikTok troll military to coax the younger Filipinos' voters [6]. Furthermore, the breach of privateness has raised greater troubles. These elements have sparked a debate regarding the effect AI development has. Thus, the main question arises, is the increase in AI a fine improvement? Let's first take into account the perspective that the growth in AI is a fine development. The advantageous impact of AI is in reality seen within the subject of medical, pondered in AI healthcare marketplace accomplishing an predicted fee of \$11 billion in 2021 [7]. In a look at, a crew from Germany, France and United States proved that AI-based Computer-Aided Diagnosis (CAD) System outperforms human-primarily based prognosis of cancerous skin lesions [8]. Engaging researchers from diverse international locations offers global illustration, and helps put off capacity biases - by using supplying broader know-how, one-of-a-kind perspectives, cross-cultural validation and generalizability, thereby substantially improving the credibility and authenticity of the supply. The argument about positive impact of AI on healthcare is similarly supported by another study finished by using IBM- International Business Machine (IBM) [9] where it turned into deter- mined that 64% of sufferers are comfortable with using AI for continuous get entry to to statistics, furnished by means of an AI-primarily based Virtual Nurse. As in line with the National Academies of Sciences, Engineering, and Medicine (NASEM), the researchers observed that there may be a scarcity of nurses because of will increase in the populace and the high charge of leaving the nursing career by means of the young nurses [10]. In this kind of state of affairs AI is again playing its position and presenting digital nurses to conquer the dearth of nurses. IBM, main in AI because of its extremely good contributions and NASEM, a famend institute providing evidence-based steerage to the government and public, both are identified for their reliability and robust credentials of their respective domain names. Furthermore, it's far argued that Vehicles having the capability to navigate and function without human manipulate, referred to as Autonomous Vehicles (AV), which might be some other spinoff of the state-of- the-art advancements in AI leading to much less injuries. A observe published in the Journal of Advanced Transportation known as "Safety of Autonomous Vehicles" in comparison the variety of injuries caused by people versus AI. Over three.7 kilometers, out of 128 accidents, 63% of the accidents have been due to autonomous riding. How- ever, most effective 6% had been directly linked to AV while 94% have been because of 1/3 parties, inclusive of pedestrians, cyclists, motorcyclists, and conventional motors [11]. Thus, demonstrating that AV is safer as 94% injuries had been due to human beings. The study, authored through Jun Wang, Li Zhang, Yanjun Huang, and Jian Zhao, emphasizes the validity of their argument because of their expertise. They exam- ine the potential shortcomings of AVs, despite statistical proof of their protection. In any case, the have a look at doesn't just highlight issues but also proposes solutions, improving the reliability in their argument. This argument of AI having a tremendous impact on AV is likewise supported by way of Swiss Re and Waymo's collaborative study. Swiss Re's baselines, based totally on 600,000 claims and a hundred twenty five billion miles, allow powerful contrast [12]. The Waymo Driver, in assessments without human intervention, precipitated no damage claims

compared to 1. Eleven per million miles for human drivers, and had significantly fewer belongings harm claims at 0. Seventy eight in line with million miles versus three.26 for human drivers proving how AV has a fine impact on the society. The record changed into composed through authentic organizations which have the manner to acquire applicable facts and have a duty to uphold their recognition of imparting correct facts, which makes their argument extra credible. The evidence provided is likewise relevant which similarly highlights how using AV is important to reduce injuries which generates nice social consequences. However, critics contend that as AI maintains to advance in numerous fields, it poses a hazard to societal protection via facilitating the spread of incorrect information. This misinformation can cause impaired choice-making as people may be manipulated socially, highlighting concerns regarding the capability bad influences of AI on society. In 2020, researchers at the University of Sussex and an advertising employer made a deep-fake video of preceding U.S. President Barack Obama, highlighting AI's capacity risk by spreading misinformation and dissolving consider in dependable sources [13]. This underscores the want for moral policies and guidelines to keep away from mis- use. The proof is dependable because it comes from an academic placing and involves collaboration with a reputable establishments [14]. Hany Farid, a professor at the University of California, Berkeley, has notably researched deep fakes and their societal implications, specializing in strategies to mitigate their threats. Collaborations with respectable institutions just like the Electronic Frontier Foundation (EFF) and the Brookings Institution make certain the reliability and rigor of deepfake research [15]. This emphasizes how the development of Artificial Intelligence leads to societal decline, as respected individuals have progressed to underscore this hazard for the general public. Another instance is the Facebook/Cambridge Analytica case which suggests how facts misuse can effect ideals and movements [16]. Information from thou- sands and thousands of Facebook users changed into collected without authorization to make psychographic profiles for centered political advertisements. This affected the 2016 U.S. Election and Brexit vote, raising worries approximately social manipulation, statistics security, and political impact The observe where this turned into posted, is a strong piece of evidence because of the reliable affiliations and various understanding of its authors, the multidisciplinary method, and its book in 2019 thru a peer-reviewed technique, indicating relevance and scholarly validation. But what weakens is that the writer does not completely discover counter arguments or spotlight any realistic solutions. Secondly AI has made hacking simple, threatening our on line records, presence, and safety. An instance is Conficker, additionally called Kido, Downadup, and Downup, which became notorious as a worm that infected tens of millions of computers around the arena in the past due 2000s [17]. Exploiting vulnerabilities in Microsoft Windows, it spread over structures, creating a botnet managed by means of cyber- criminals. The worm should disable protection packages, take touchy statistics, and launch DDoS assaults against websites. In spite of efforts to mitigate its effect, Conficker stays energetic on several systems, posing a continual hazard to safety. Conficker highlights the link among cybersecurity dangers like malware and societal security. It underscores the project of guarding AI systems from cyber threats [18]. Conficker's persevering nearness serves as a clear name for proactive measures and time-honored collaboration to stand up to the risks posed by using pernicious on-display characters in the virtual domain. Conficker's presence requires proactive measures and common collaboration to confront malicious actors within the digital area. Be that as it is able to, the paper does no longer cope with improving popularity-based blacklisting for identifying Conficker or provide counterarguments. In spite of the truth that it covers 25 million victims, giving valuable insights, no empirical proof is utilized. Generally, the study helps us recognize the significance and seriousness of this circumstance and the negative impact it has on our society Another comparable case is a major phish ing scheme focusing on Google customers in May 2017, encouraging them to click on false e mail links [19]. These emails postured as actual invitations to collaborate on Google Docs, main recipients to a faux Google login page where hackers stole their login records. Google diagnosed this phishing attack as a actual safety threat after broad media coverage. Thorough examinations of the phishing emails and the phony website offered verifiable proof of the cybercriminal's strategies. The severity of the scenario became emphasised by using

Google's activate motion, which included the suspension of malicious accounts and the creation of similarly security measures [20]. The proof from the Association for Computing Machinery (ACM) is reliable due to its thorough peer overview method, making certain accuracy and pertinence. ACM's appeared reputation as a leading authority in computing solidifies the reliability of its facts. Overall, this incident highlights how simply one click can reason a person to lose the whole lot, from his personal info to his debts, the whole lot is under threat, resulting in an unfavourable final results for society. Both views on AI pos ing a threat to societal protection made logical claims which had been subsidized via credible sources to bolster their argument. Those arguing in want of AI included researchers from Germany, France and the United States which increases its generalisability, allows a comparative evaluation and offers pass cultural validity. It additionally hyperlinks to how AI plays an immediate and vast role in assisting and assisting suf- ferers. The statistical and numerical data provided, quantitatively validates the large fantastic impact of AI in clinical and healthcare, strengthening the argument with the aid of demonstrating vast financial boom and relevance within this zone. Qualitative and Quantitative proof showcases the impact the AV cars have added in decreasing accidents. A weak spot is evident in generalizing AI popularity information without accounting for demographics, cultural nuances, or specific situations that could affect patient consolation, thereby diminishing the precision and applicability of the offered statistics. The competition's worries about AI risks benefit traction through the Face- book/Cambridge Analytica scandal which serves as a compelling real-global example of AI's potential for social manipulation and political affect, bolstering concerns about its risks. Moreover, incorporating examples just like the Conficker trojan horse and the Google phishing scheme offers tangible evidence of AI's cybersecurity threats as exact specifics decorate persuasiveness and shed light on broader societal implications. A top notch shortcoming is the absence of exploration into present or proposed AI legislation, crucial for understanding administrative structures. Also, both viewpoints lack particular methodological information, doubtlessly undermining the findings' reliability. Despite those weaknesses, the compelling evidence from each sides leads to the belief that favorable AI effects are feasible with suitable rules.

3. Methodology

The current study assumes a mixed-method approach by combining qualitative inputs through an extensive literature with support from quantitative data collected via direct survey. In order to find this topic, a lot of different sectors are going to be affected in not only healthcare but cyber security and autonomous systems as well so by showing public opinions about the societal risks/benefits involved with AI compared to that of expert opinion on these three topics.

3.1. Phase 1: Literature Review

We conducted a systematic literature review of academic and industry reports that cover possible safety-related applications around AI. Boys also phrased the challenges in AI centered on medical diagnosis, autonomous automobile or cyber security — with an emphasis of using AI as a tool to firm workflow stability and decrease human error. It also examined problematic uses such as deep fakes, time-wasting algorithms that promote inactivity or foster addiction (e.g. of gambling), social manipulation — and other unintended applications like automated cyber attacks. From top AI, computer science and ethics journals etc. as well few real-world use cases of implementation of AI systems were included in these studies The literature review knowledgeable the survey layout through identifying the primary benefits and dangers that had to be explored similarly. It additionally set the muse for information regulatory needs and moral demanding situations associated with AI's unchecked growth.

3.2. Phase 2: Survey Design and Data Collection

A move-sectional survey is designed and developed to quantitatively investigating the perceptions of AI's impact on societal safety. The survey design concentrated on each specialists and the overall public.

3.2.1. Survey Structure:

The survey contained 30 gadgets, divided into three main sections:

- Demographic Information: Participants' age, gender, career, and level of familiarity with AI technology.
- AI Perceptions: Likert-scale questions assessing respondents' views on AI's role in enhancing societal protection, its capacity risks, and regions wherein law is essential.
- Ethical and Security Concerns: Questions focused on awareness of AI misuse (e.G., deep fakes, computerized cyber-attacks) and agree with in AI-driven systems (e.G., scientific diagnostics, self reliant vehicles).

3.2.2. Target Population:

The survey aimed to capture statistics from distinct populations:

- AI Experts: Researchers, developers, and lecturers in AI and associated fields. Participants had been recruited via academic and expert networks.
- General Public: Individuals with various levels of familiarity with AI. The survey was disseminated through social media, on-line boards, and email lists.

3.2.3. Sampling Method:

A combination of proportional sampling for AI professionals and random sampling for the overall public changed into employed. An overall of 60 members (n=60, with 20% specialists and 80% from the overall public) completed the survey. The sample length became decided based totally on power analysis to make certain sufficient statistical energy for significant comparisons.

3.3. Phase 3: Data Analysis

Survey responses were analyzed using both quantitative and qualitative methods:

3.3.1. Quantitative Analysis:

Descriptive statistics, including means and standard deviations, were calculated to summarize the general trends in AI perceptions. Inferential statistics (e.g., t-tests, ANOVA) were used to compare differences between experts and the general public regarding the perceived benefits and risks of AI. Correlations were also examined to identify relationships between demographic factors (e.g., familiarity with AI) and perceptions of societal safety.

3.3.2. Qualitative Analysis:

Open-ended responses were coded using thematic analysis to identify common themes related to ethical concerns, governance, and AI's societal implications. Recurring concerns, such as the potential for AI misuse in misinformation and privacy violations, were categorized and analyzed to provide deeper insights into participants' views.

3.4. Questionnaire

The survey questions capture an insight into people's perspective towards AI, its potential threats and its impact on various aspects of life. Statistics gained from our survey regarding AI, its societal impact and public attitudes towards AI, revealed a wealth of insights reflecting the increasing complexity of the AI's role in our lives. The aim was to have a deep insight into how

people perceive AI, its potential threats especially in critical areas of day to day life which is associated potential life threats.

3.4.1. AI familiarity and Impressions

The survey targeted how many respondents are familiar with AI and its applications, with an overall expression of positive, negative or neutral impression. However, the significant differences in comfort levels and perceptions of the potential risks of AI are also captured in the questionnaire.

3.4.2. Perceived Threats and Comfort Levels

Questions are asked about AI's threat to society, with a scale capturing different opinions. The proportion of respondents concerns about the impact of artificial intelligence and emphasized the need for thoughtful consideration of its societal implications is also captured via questioning. Despite these concerns, there was a relatively high level of comfort with AI applications in critical areas such as medicine and surgery, indicating confidence in the potential benefits of AI when used responsibly, hence questions are designed to address these concerns as well.

3.4.3. Fear and Trust

The use of artificial intelligence in autonomous vehicles has raised significant concerns, reflecting ongoing debates about safety and reliability in this emerging field. Trust in artificial intelligence in handling personal data and privacy was another area of concern, which is tackled by introducing questions related to trust and skepticism, underscoring, and the importance of robust data protection measures.

3.4.4. Regulation and Future Impact

A clear idea of strict regulations in the development and deployment of AI, indicating a strong desire for oversight to ensure the ethical and safe use of AI technologies is also incorporated by survey questions. Difference of opinion is captured on whether AI will have a more positive or negative impact on society in the next decade, highlighting the need for continued dialogue and research to navigate the evolving AI landscape. Overall, the survey underscores the public engagement and transparent discourse about the future of AI. The findings will offer valuable guidance for technologists, policymakers, and educators as we work to harness the benefits of AI to address potential risks.

3.4.5. Main Questions

- Q1.** How familiar are you with AI and its applications?
- Q2.** What is your overall impression of AI?
- Q3.** To what extent do you think AI poses a threat to society?
- Q4.** How comfortable are you with the idea of AI being used in critical areas like medicine and surgery?
- Q5.** How concerned are you about the use of AI in autonomous vehicles?
- Q6.** Do you trust AI to handle personal data and privacy?
- Q7.** Should there be strict regulations in the development and deployment of AI?
- Q8.** Do you think AI will have a more positive than negative impact on society over the next 10 years?

3.5. Ethical Considerations

This study adhered to strict ethical hints. Informed consent was acquired from all members, making sure that they understood the reason of the study and that their participation became voluntary. Data turned into anonymized to shield individuals' privacy, and the study acquired ethical approval from the relevant institutional evaluation board. The blended-strategies approach

lets in for a radical investigation into both theoretical discussions and real-global perceptions of AI’s societal effect.

By combining insights from the literature evaluate with empirical statistics from the survey, this look at presents a nuanced understanding of the balance between AI’s advantages and risks, as well as the regulatory measures essential to safeguard societal protection.

4. Results

The survey results are compiled and illustrated as pie-chars and histograms (Figures 1–8) for data analysis and visualization. The survey affords a compelling photo of public sentiment on the question, “Is the latest increase in AI posing a chance to societal safety?” Responses suggest that while AI has emerge as an increasing number of acquainted, there are giant concerns approximately its implications. When requested about familiarity with AI (Q1), the bulk of respondents (fifty one 7%) pro-nounced slight understanding, suggesting that many are privy to AI however might not completely apprehend its complexities as illustrated in Figure 1. Only 15% indicated a deep familiarity, at the same time as 26.7% had been handiest barely familiar, hinting that there’s nonetheless a gaining knowledge of curve for a sizable part of the populace. Impressions of AI (Q2) similarly screen a cautious outlook, with 55% of respondents holding a negative view of AI, even as only 28, as detailed in Figure 2. 3% have a nice impression, and 16.7% remain impartial. This terrible sentiment can also stem from issues about AI’s capacity risks to society. These concerns are echoed in responses to whether AI poses a hazard to society (Q3), presented in Figure 3. Nearly 35% rated AI as a slight chance (on a scale of one-five), even as 26.7% leaned toward a higher hazard belief (four on the scale), and 15% expressed a robust notion that AI represents a major threat. Combined, this indicates extra than three-quarters of respondents see AI as at the least a slight danger to societal protection. The survey additionally probes evaluations at the application of AI in important regions, which include remedy (Q4), Figure 4 and motors (Q5), Figure 5. The use of AI in medical regions is met with warning, with 26.7% of respondents score it at the lowest safety degree, while best 10% sense particularly steady about AI’s function in such touchy applications. Similarly, in relation to AI in vehicles, 30% of respondents expressedsome consolation with AI involvement, however the majority remained hesitant, with 26.7% score it poorly. Trust and privacy worries are highlighted in Q6, Figure 6, wherea powerful 60% of respondents said they do now not consider AI in relation to privacy. 3% were uncertain, leaving most effective 16.7% feeling assured in AI’s managing of touchy records. Finally, the overwhelming help for AI regulations (Q7), Figure 7, in which 68.3% of respondents prefer guidelines, underlines the public’s desire for oversight to mitigate the perceived dangers of AI, with 0% opposing regulation, and 26.7% uncertain however open to the idea. In precis, whilst the latest AI boom has introduced familiarity, it has also raised extensive concerns regarding societal protection, trust, and the capacity need for law, particularly in crucial regions like remedy and self sustaining cars. These information suggest that while AI holds promise, its effect on society is regarded with a combination of warning and skepticism.

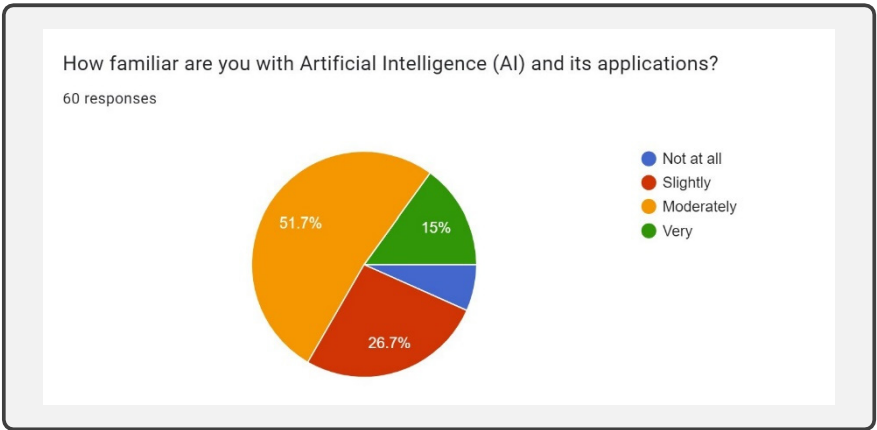


Figure 1. Q1. A graphical representation illustrating participants’ responses on a 1 to 5 scale, assessing their familiarity with AI and its applications.

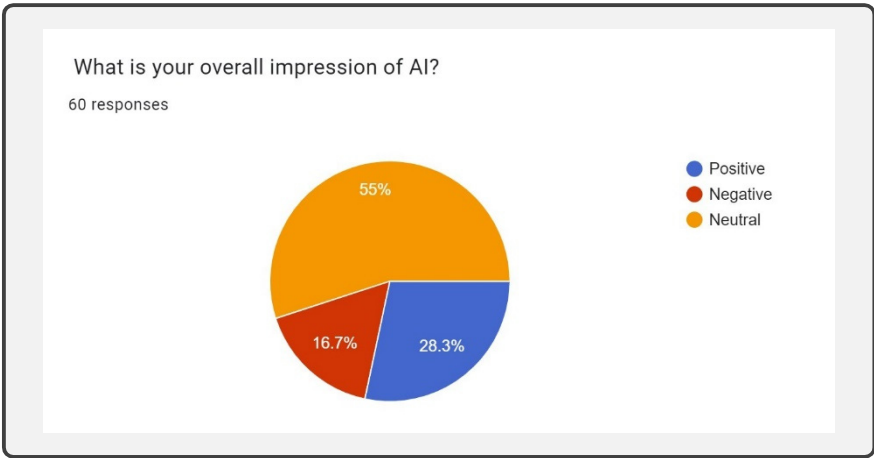


Figure 2. Q2. Participants’ responses (positive, negative or neutral), for their impression of AI.

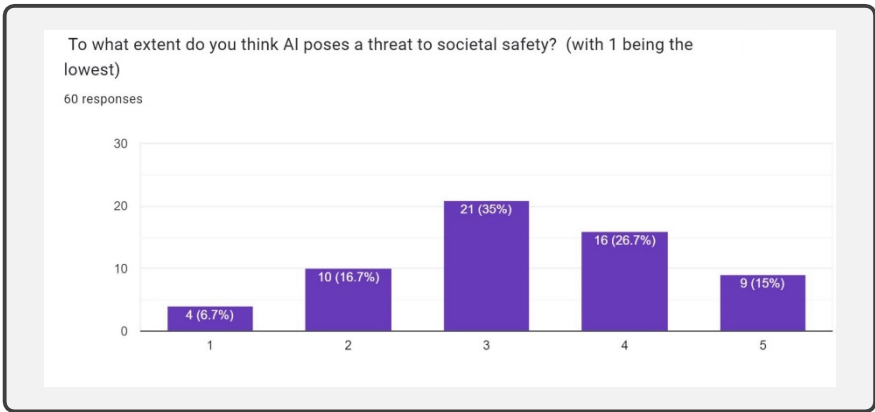


Figure 3. Q3. A histogram representing participants’ responses on a scale of 1 to 5 regarding their perception of AI’s applications, posing a threat to societal safety.

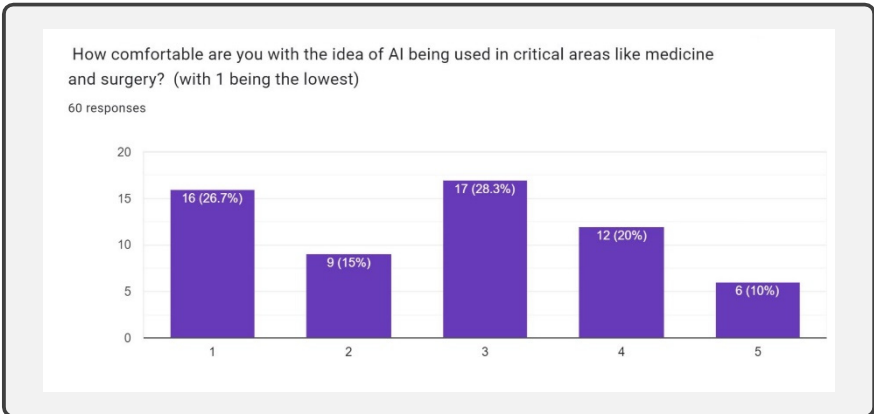


Figure 4. Q4. An illustration representing the peoples’ awareness regarding their day to day encounter with the presence of AI, and their trust in AI within critical fields such as medicine and surgery, as revealed by a sub-sample of the population.

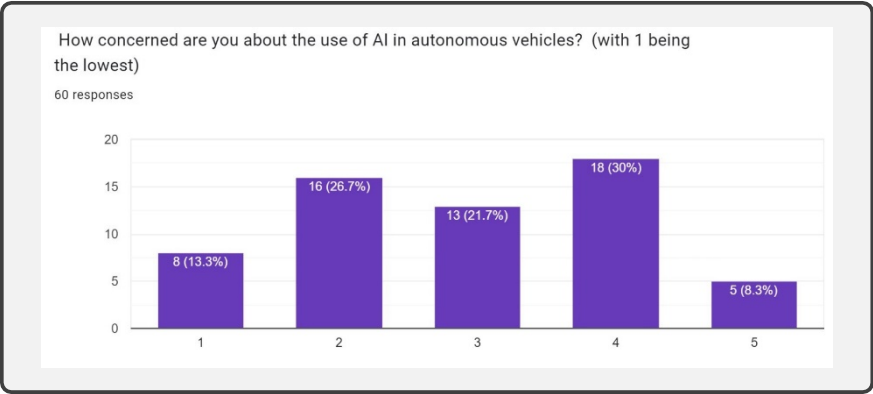


Figure 5. Q5. Histograms of participant’s concern level regarding autonomous vehicles (AV) being driven by AI.

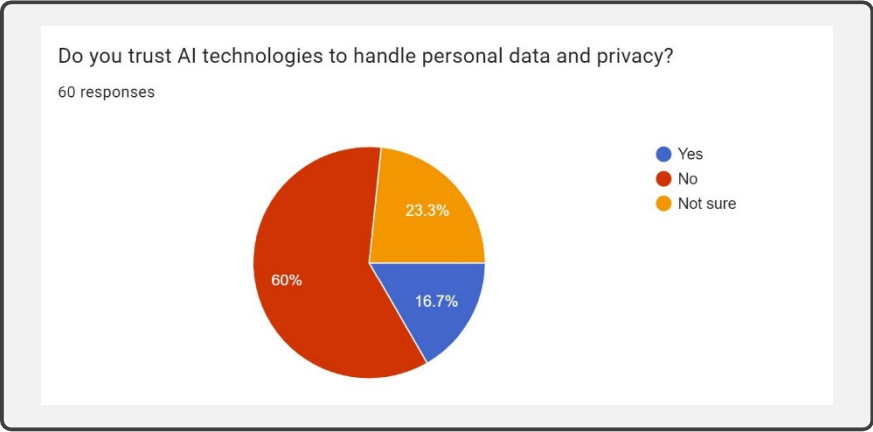


Figure 6. Q6. Graphical illustration of participants’ trust level regarding the use of AI and data privacy.

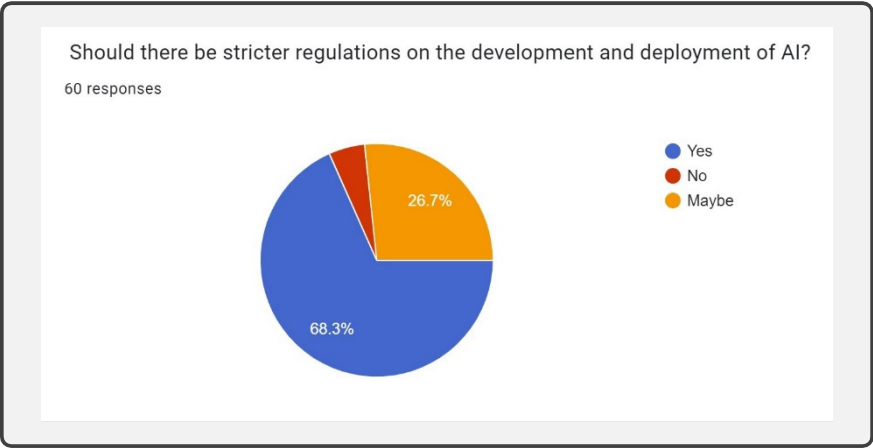


Figure 7. Q7. The results of assessment of society’s need regarding strict regulations and policies for the use of AI.

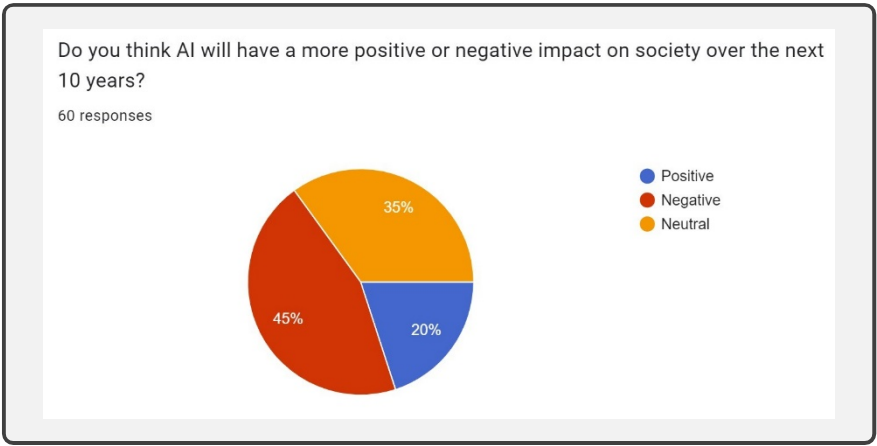


Figure 8. Q8. Results of participants’ perception regarding the future of AI and its impact over the next 10 years.

4.1. *Scaling*

To scale all responses to 3 values, different statistical methods are applied considering the number of options in each question. Here’s the breakdown to handle each case:

- 1. **For questions with 4 or 5 values:** a weighted average is applied for questions with 4 and 5 values to reduce the data to 3 values.
- 2. **For questions already having 3 values:** Questions with 3 values were kept as-is.

Weighted Average: The Equation 1 is used for converting the original values ranging from 1 to 4 and 1 to 5 into 3 categories while keeping the distribution of data intact.

$$WeightedAverage = \frac{\sum(value \times weight)}{\sum weight} \tag{1}$$

5. **Discussion**

The recent increase in the development of artificial intelligence caused both excitement and concerns about the impact on social security. AI promises a transformer training the benefits of sectors such as health care, transport, industries, etc.The increase in literature suggests that rapid development can be important. There is a risk if they are not properly controlled. Ethical dilemma, data confidentiality, moving work, and the potential inappropriate use of AI in important areas is generally referred to. Researcher emphasize the transparency in the system, responsibilities, and the necessity of human supervision time to help public interests without impairing security and security. The survey findings are consistent with these concerns, revealing a public opinion that reflects both wariness and skepticism about the role of artificial intelligence in society. While there is a general knowledge of AI and awareness of its potential, many express doubts about its reliability, particularly in areas such as healthcare and self-driving cars. The call for regulation is strong as the public increasingly demands frameworks that can ensure that AI is developed and deployed responsibly.

Considering the major factors as below:

- 1. **Factor 1 - Scale 1 to 3:** Represent respondents’ perception of AI as a threat, where:
 - Scale 1: Low threat perception.
 - Scale 3: High threat perception.
- 2. **Factor 2 - Expert Opinion:** 20% of the respondents are experts, and they predominantly see AI not as a threat

The general distribution, as illustrated in Table 1, analysis shows certain trends in experts’ and non-experts’ opinions and reveals a wonderful difference within the belief of AI as a societal

threat between specialists and non-professionals. Experts, who make up 20% of the respondents, predominantly view AI as posing a low chance. We see that in Q1, in which 16.65% originally saw AI as a low hazard, about 13.32% of this got here from non-professionals. Similarly, for Q2, 55% viewed AI as low risk, but after accounting for professional responses, 44% of non-specialists nevertheless perceived AI as less threatening. In contrast, non-professionals expressed greater challenge in Q3 and Q4, with a tremendous percent (approximately 35% in Q3 and 30% in Q4) seeing AI as a moderate to excessive chance. This analysis shows that while specialists remain assured in AI’s safety, non-experts have a more cautious stance, with a larger proportion of them viewing AI as a mild to large risk to society. The hole among professional and non-professional perception highlights the want for higher public knowledge and conversation concerning AI’s real-world applications and obstacles.

Table 1. Final Scaled Data Table.

Qs	Scale 1	Scale 2	Scale 3
Q1	16.65	51.7	15
Q2	16.7	55	28.3
Q3	11.7	35	20.85
Q4	20.85	28.3	15
Q5	20	21.7	19.15
Q6	23.3	60	16.7
Q7	23.3	60	16.7
Q8	35	43	20

6. Conclusions

Before delving into this subject, my comprehension of AI was rather limited, shaped by common beliefs heavily influenced by cinematic portrayals, implying the potential dominance of AI over the world. However, our view underwent a transformation when we delved into the positive impact of AI in the medical sector and fewer accidents caused by AVs. Convincing evidence and statistical findings have revealed a prevailing comfort with using AI. However, upon further research, we discovered ethical issues and manipulative tendencies associated with AI, which made me question whether AI is positive or not. Another survey found that despite its merits, AI is inadvertently posing challenges to society as a whole, an unfavorable development. In conclusion, AI can be a beneficial advancement and we should encourage its use. However, it is necessary to introduce strict regulations and rules to prevent its misuse. AI, which is inherently neutral, is completely dependent on how it is used, which highlights the critical need for responsible governance.

In conclusion, the literature and survey data agree that while AI has great potential to improve lives, it also poses significant risks. As technology continues to advance, the need for ethical guidelines, regulatory oversight, and public trust is ever more important. Only with careful management can the benefits of artificial intelligence be reaped without jeopardizing societal security.

References

1. Simplilearn.com: Top Artificial Intelligence Stats You Should Know about in 2023. <https://www.simplilearn.com/artificial-intelligence-stats-article#:~:text=The%20global%20AI%20market%20is>. Accessed: 2024-09-12 (2021)
2. Simon, H.A.: Studying human intelligence by creating artificial intelligence: When considered as a physical symbol system, the human brain can be fruitfully studied by computer simulation of its processes. *American Scientist*
3. AI Impact on Healthcare in Developing Countries. <https://www.linkedin.com/pulse/ai-impact-healthcare-developing-countries-jose-claudio-terra-phd>. Accessed: 2024-09-12
4. Autonomous Vehicles Cause Fewer Injuries and Fatal Crashes. *Warp News*. Accessed: 2024-09-12 (2023). <https://www.warpnews.org/transportation/autonomous-vehicles-causes-less-injuries-and-fatal-crashes/>

5. William Marcellino, e.a.: The Rise of Generative AI and the Coming Era of Social Media Manipulation 3.0: Next-Generation Chinese Astroturfing and Coping with Ubiquitous AI. <https://www.rand.org/pubs/perspectives/PEA2679-1.html>. Accessed: 2024-09-12 (2023)
6. Ong, J.C.: Philippine elections 2022. *Contemporary Southeast Asia* **44**(3), 396–403 (2022)
7. Stewart, C.: AI in healthcare market size worldwide 2021-2030 (2023)
8. Presse, A.F.: Computer learns to detect skin cancer more accurately than doctors. *The Guardian* **29** (2018)
9. Polachowska, K.: 5 Medical Challenges That Can Be Solved with AI in Healthcare - Neoteric. <https://neoteric.eu/blog/5-medical-challenges-that-can-be-solved-with-ai-in-healthcare/>. Accessed: 2024-09-12 (2019)
10. Kathleen Sanford DBA, e.a.: Virtual Nursing: Improving Patient Care and Meeting Workforce Challenges (2023)
11. J. Wang, Y.H.J.Z.F.B. L. Zhang: Safety of autonomous vehicles. *Journal of advanced transportation*, 1–13 (2020)
12. Luigi Di Lillo, e.a.: Comparative safety performance of autonomous-and human drivers: A real-world case study of the waymo one service. arXiv preprint arXiv:2309.01206 (2023)
13. Deep Fake video of Barack Obama shows new disinformation frontier. <https://www.youtube.com/watch?v=AmUC4m6w1wo>. Accessed: 2024-09-12 (2020)
14. Farid, H.: Deep fake detection: Current challenges and future directions. *ACM Transactions on Multimedia Computing, Communications, and Applications* **16**(3), 45–58 (2020)
15. Institution, E.F.F.E..B.: Combating Deepfake Disinformation: A Collaborative Research Initiative. <https://www.eff.org/deepfake-research>. Accessed: 2024-09-12 (2021)
16. Cambridge Analytica and Facebook: The Scandal so Far (2018)
17. Seungwon Shin, e.a.: A large-scale empirical study of conficker. *IEEE Transactions on Information Forensics and Security* **7**(2), 676–690 (2011)
18. Daniel Bilar, G.C., Murphy, J.: Adversarial dynamics: the conficker case study. In: *Moving Target Defense II: Application of Game Theory and Adversarial Modeling*, pp. 41–71. Springer, ??? (2012)
19. Vanessa Gomes, B.A. Joaquim Reis: Social engineering and the dangers of phish- ing. In: *2020 15th Iberian Conference on Information Systems and Technologies (CISTI)* (2020). IEEE
20. Nakamura, A., Dobashi, F.: Proactive phishing sites detection. In: *IEEE/WIC/ACM International Conference on Web Intelligence* (2019)

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.