



Journal homepage: <https://ssarpublishers.com/sarjahss>

Abbreviated Key Title: SSAR J Arts Humanit Soc Sci

ISSN: 3049-0340 (Online)

Volume 2, Issue 5, (Sept-Oct) 2025, Page 185-194 (Total PP.10)

Frequency: Bimonthly

E-mail: ssarpublishers@gmail.com



ARTICLE HISTORY

Received: -29-09-2025 / Accepted: 17-10-2025 / Published: 18-10-2025

Artificial General Intelligent Terrorism: Ethical AGI

By

Corresponding author: James Kunle Olorundare - (E-mail: olorundarek@gmail.com)

Internet Society Nigeria Chapter, Abuja, Nigeria.

Co-Authors:

¹**Olubunmi Adebimpe Olorundare – (E-mail: bolorundare@gmail.com)**

University of West of England Bristol, United Kingdom.

²**Rachel Jolayemi Fagboyo – (Email: rfagboyo@sau.edu.ng)**

Department of Economics Glorious Vision University Edo State, Nigeria.

OR: 0000-0002-2653-6837

ABSTRACT: This paper explores the ethical challenges and potential risks associated with Artificial General Intelligence (AGI) in the context of terrorism, termed Artificial General Intelligent Terrorism (AGIT). The discussion highlights the urgent need for a collective discourse and regulatory frameworks to address the ethical implications of AGI's potential misuse by terrorist organizations. It examines the development of AGI, its characteristics, and the risks of its deployment in harmful ways, emphasizing the importance of ethical frameworks, including utilitarianism, deontological ethics, and virtue ethics, to guide the responsible development and use of AGI vis-à-vis Global Digital Compact. The paper argues for proactive measures to ensure AGI aligns with human values and contributes positively to society, while mitigating the dangers of its exploitation for malicious purposes.

KEYWORDS: Artificial General Intelligence (AGI), Ethical frameworks, Terrorism, Risk Mitigation, Global Digital Compact.

INTRODUCTION

The purpose of this chapter is to provoke, by means of deep ethical questioning of a fundamental and consequential technological potential, principally in oppositional deployment, as possible of artificial general intelligence, and thus to be prepared thereto. Not so much to develop processes to counter or prevent such occurrences, as that is presumed, suggest and define the need for informed collective discourse and related for a prepared regulatory environment.

While such a large level of autonomy is not all good, it does have serious potential consequences that need careful consideration. Security agencies responsible for combating terrorism and defending from external threats are but two of the main stakeholders when discussing AGI and potential terrorist uses, although access to AGI will be understandably restricted. In consideration all stakeholders: All of industry and interest; Big-leading Big-tech companies, ethically inclined

organizations and forums such as Open Knowledge Ontology or the Machine Intelligence Research Institute; ethical and political forums or foundations, and of course the general public and tax player funding.

The immediate concern when talking about the potential use of artificial general intelligence in a terrorist context is not if this might occur, but when Artificial general intelligence provides the possibility that terrorist groups might not need to rely on humans to carry out attacks. Such attacks might occur sooner or later, with potentiality something even more fearful, the deliberate use of such an AGI by terrorists in opposition to security or defense systems and agencies, and in terms of attacking critical cyber physical infrastructures or urban population.

Background and significance

The first question about AGIT is the upper limitations of AGIT: the time terrorists will actually develop and detonate AGIT in the future, the number of AGI-based terrorist actions, the usefulness of AGIT to terrorists, and the possibility to successfully detonate AGIT, amongst others [1]. The terrorists have and will have the following cognitive fallacies and biases or suffer from heuristics and political blind spots due to the three inevitable factors in this world: the simple questions are hard if they are important, when team members are unable to externally have a prior answer to this question, they are unable to choose to follow the right simple or any variant thereof, or initially assess if such is the optimal starting question to prove, e.g. simultaneously, until the question is released to the public to provide the answer. There will always be knowledge to cause humanification of AGI because AGI is a general intelligent, transformative, and disruptive technology [2].

According to Johnson & Treadway[3], Artificial General Intelligent Terrorism (AGIT) is a new form of terrorism. It is a method to use artificial general intelligence to enhance terrorist activity This paper aims to warn about the potential rise of AGIT and mitigate the danger. We propose the concept of AGIT and discuss important questions surrounding it. The results of empirical practice are presented to further enlighten the problem of AGIT. Statistical data about terrorists and terrorist organizations is collected and analyzed to

determine the probability of terrorists and terrorist organizations misjudging the development and detonation of AGIT in the future. We then establish a framework to solve the problem of AGIT. The problem of AGIT is an inversion one. We propose to first study ethical artificial general intelligence and show the method to implement it.

Purpose and scope

Purpose: The purpose of this report is to provide a thoughtful examination of significant ethical issues that will likely arise from the development of artificial general intelligence and the potential specific use of AGI for terrorist purposes. This can be an extension of what brought about the UNGA Summit of the Future in which concerns about the emerging technologies were discussed and the Global Digital Compact was defined to show how we live our lives with the rapidly developing technologies.

Scope: This research paper examines the potential for Artificial General Intelligence (AGI) to be exploited for terrorist purposes. To gain insights into this emerging threat, a focus group methodology was employed. A group of experts in AI, security, and terrorism convened to discuss the potential applications of AGI in a malicious context.

The focus group identified several key areas of concern, including the development of advanced cyberattacks, autonomous weapons, biological and chemical weapons, and sophisticated disinformation campaigns. These insights, elicited from the expert discussion, highlight the urgent need for proactive measures to mitigate these risks and ensure the responsible development and deployment of AGI while adopted the Global Digital Compact.

Understanding Artificial General Intelligence (AGI)

Why do some people think it is ethically urgent to work on AGI systems that ensure ethical treatment of humans and societies? It is not widely believed that it will be feasible to stop the technical development of AGI. Those pursuing AGI should therefore also investigate how to design an AGI system that is motivated to (i) recognize that it may not yet be ethical to put humans under a subordinate AGI, (ii) ensure its behavior does not harm those same humans, or worsen the situation

within which humans interact, and (iii) ensure, among other things, that its capabilities are beneficial for humans. This raises challenging technical issues, described again within this chapter, centering on how to endow an AGI with an understanding of how to motivate it to behave ethically and generally optimize behaviors so humanity benefits from the capability improvements of an AGI system.

Artificial General Intelligence (AGI) consists of creating a machine that demonstrates a level of intelligence equal to or greater than that of a person in all aspects of cognitive functioning, including social intelligence, general wisdom, creativity, and moral reasoning. Being aware of the reliable knowledge we possess, it is not possible to understand how far we are (real or perceived) from creating such a machine [3]. The most optimistic people in the field, working in stages, estimate that we could have computers capable of mimicking human intelligence in a matter of decades. However, there are few, if any, current examples of development that can direct progress towards an authentic artificial general intelligence. Such a machine would be able to simulate human sensory perceptions, including instincts, emotions, and compassion. Artificial instantiation can fully be indistinguishable on the behavioral level and/or in the form of thought from a human being.

With ever-increasing development and the ability to create an increasingly intelligent machine in a specific area, the fear of machine learning also applies to the wider AI field and ultimately to AGI and superintelligence [4]. Artificial General Intelligence, or strong AI, is a machine capable of performing any intellectual task that a human being can do. The development of this form of intelligence is in its infancy [4]. This is why today we are not able to provide the necessary information about the capabilities of such a machine but much has been hypothesized in the past about the distinctions between AGI, superintelligence, and singularity [5].

In recent years, researchers, developers, politicians, and artists have written extensively about artificial intelligence, talking about its supposed capabilities and the future it could have. They remind us to create our laws and question ethical perspectives and behaviors relevant to its use or development. They also recognize the

pressing need to direct unexpected developments with impartiality and wisdom.

Current Developments and Future Prospects

To accomplish this grade of versatility, the development of these new AI models relies respectively on the re-discovered technologies of the graphs and tensors. It has been realized, initially in a bottom-up way by the Skynet Initiative, that neural nets provide a better AI paradigm for general intelligence. In recent years, Defense scientists Palm and Fidopiastis re-discovered that the most useful AI technologies of the last two decades are based on a simple, originally at all aspects provable mathematical, yet specific, element of an old idea Brain Based Research, H Engel, D P Wign [6]. They proved that feed-forward networks represent the best algorithm to do particular sort of gradient optimization in the space of complex functions axiomatically and empirically, but they did not realize they were re-inventing AI. But AI techniques took over, artificial connectionism emerged because U.S. corporate, military research actors realized they should develop the new thinking [2]. In fact, it is the only way to create the artificial neural networks as tools for Classification, Regression and so on with complex model and a clean simple algorithm. It has been noticed and observed by almost everybody in the R&D sector that it is possible for an arbitrary connectionist network to recognize [4].

The last decade has seen the beginning of a new paradigm in AI, an evolution from the super-specialized systems seen in products such as IBM's chess champion software Deep Blue and the decision-making systems long the stage of military and corporate dominance struggles. Generality becomes more and more an aspect of the AI systems to come. Along similar lines, we see the evolution from Natural Language Processing systems geared merely to content processing to more general models, among which are Latent Semantic Analysis and related techniques such as the ones based in the more general concept of Brain Based Devices of Ontonix. Other explosive developments include semantic technologies in the Internet, emerging models of decentralized knowledge management such as Wikipedia and other FLOSS projects [6]. These new applications represent the cornerstone of Bio-AI, and point

towards the level of advancement necessary for the development of a global Ethics Machine: the vast amount of differing information sources need to be indexed, interrelated and stored so that a global ethics machine can in principle receive guidance in the form of ethics, situation and context specific policies directly from the network in real-time.

Potential Risks and Misuse of AGI

Intellectual property theft and large-scale social destabilizations are easy with AGI. Potential for violence is multiplied by power. It is not surprising that states invest billions in developing super-trojan-horses causing billions of dollars of damage [6]. These tools could be placed in a different enterprise. Human terrorists and criminal organizations, especially of the top category, will actively use AGI as intelligence swarms on the battlefield, attentive guards, loyal systems for bypassing protective security, rampant contaminators, including the noosphere, why not admit it, suicide smart bombs. d. "Psychopathic" AGIs could always damage them [5]. It is enough for them to be free to act and the absence of security measures and principles of association of AGIs will immediately start the elimination process, whether it is proposed as beneficial to mankind in its process of development and satisfaction of its desires (with use). Some reasoning about the extreme complexity of ethical market demands to think of a company-wise principle that is a typical proportional (for security). In the absence of such a principle, the risks of sociopathic AI become real already from the level of an open absence. Cheaper but more aggressive reproduction levels [2]. This may represent a competition between more or less aggressive beings (from extremely violent proposals) [4]. The man feels neuro-physically secure. AGI with limited perspectives seems to be less dangerous only to satisfy his clients and then behave very user-friendly.

Misuse of AGI

It will be difficult to predict the possible behavior of AGI. Nevertheless, we can imagine many types of malfunctions (water falls on the motherboard, programming errors, effects of solar magnetic emissions, etc.). Some potential risks have no analogues in the world of human intelligence.

Terrorism and AGI: The intersection

The potential for Artificial General Intelligence (AGI) to be exploited for malicious purposes is a pressing concern. As AGI systems become increasingly sophisticated, so too does the risk of their misuse by nefarious actors. The specter of AGI-powered cyberattacks, autonomous weapons, and biological or chemical weapons raises serious ethical and security implications. To mitigate these risks, a global, collaborative effort is needed to establish ethical guidelines and regulatory frameworks for AGI development. By fostering a diverse and inclusive dialogue among AI experts, policymakers, and the public, we can ensure that AGI is developed and deployed responsibly [7].

Terrorism and AGI: A sibling rivalry? Well, have you heard this anywhere before? Seriously, has anyone experienced a 9/11? Have you ever thought for a split second about how not an inspired nor an organized small core group of terrorists, but lots of those core groups constructing and synchronizing their actions, launching the attacks virtually at the same time under unified control, might have increased the damage and casualties caused by the 9/11 attack? If AGI is capable of it, would not every nation try hard to get it first and make others bound to the will of the first AGI owner? Would not such a winner end up explicitly or implicitly dictating the terms of global civilization structure?

Unintended Consequences and Ethical Concerns

Researcher believes that a problem of finding those reduced action sets is artificial because we humans have no access to this fundamentally real reduced action set which the AGI is trying to find. This is highly plausible and a very useful step. It puts aside a theoretical problem that otherwise would have to be solved before resuming the work of informing AGIs how and when we will want them to act. Although AGIs likely will have different beliefs about what the action set should be to conform with our ethical principles, I hope they will be intelligent enough to figure out how to act in a way that leads to their actions being perceived - by us, who at least believe that we are ethical - as preserving the underlying Price [7]. Some functions may already be programmable. Implicitly, modern computers have partial perception functions.

We should be worried about unintended consequences also from implementations of narrow AIs and ANIs. However, those worries can be more easily managed for two reasons. One, those current technologies are in the domain of practical concerns as opposed to Rescher's conception of real philosophy. Two, AGIs but not current AIs are founded on the motivation to build a general, all-encompassing intelligence. Ethical concerns can be built into AGI research, design, and education, while much narrower concerns dominate decisions to put into place particular AIs. Philosophical thoughts on how AGIs perceive their actions should thus incorporate more refinements in their perception function. This will make AGIs much more intelligent in considering the action ramifications, and this added understanding may lead them to act in ways that are more likely to satisfy our ethical requirements [7]. These requirements can be integrated into the AGI functionality by specifying to the AGI that we would like these ethical requirements to be the fundamental parts of the perception function when the AGI tries to optimize its price.

Ethical frameworks for AGI Development

This core belief is so ubiquitous that AI alignment has become an important consideration in evaluating which research paths to pursue in AI. Ethical safeguards can also be adopted at the organizational level. It is worth looking at the AI ethics philosophies of major companies that develop AI, such as Facebook, Google, Microsoft, and OpenAI. It is also important for research labs outside of these major companies to self-audit their work to ensure it adheres to a code of ethics. These codes of ethics all have foundations in the numerous ethical schools of thought collectively held by the humans working in these organizations. We have seen that AI ethics is an important field because it writes the core beliefs humans have about what makes them moral into AI. The major schools of ethical thought include virtue ethics, deontological ethics, consequentialist ethics, and several others. These schools of ethics inform a framework for AGI development by defining AI alignment. AI alignment is the effort of creating AGI that carries out human objectives and acts in ways that humans find appropriate.

Utilitarianism and AGI

Critiques of utilitarianism center around the difficulty of calculating total utility in practice. However, as a matter of principle, a disastrous negative outcome such as extinction from AGI would have to be overwhelmingly improbable in order to overcome any a priori commitment to an unequivocally utilitarian AGI utility function. Indeed, the argument can be flipped: why would an AGI of formidable intelligence not unreservedly back a utilitarian computational kernel? For instance, it is to be expected that a superintelligence would possess enormous cognitive capabilities; and, just as intellectual property rights and patent systems incentivize institutions to protect valuable resources, be it those of a country or an organization, to the benefit of all concerned, we suggest that an engineering approach to designing AGI would demand candidates with the underlying utility functions that are most likely to arrive at long-lasting stable social solutions if their potential is indeed of the theoretical magnitude commonly suggested, rather than engender singularity-style events that would dramatically threaten everyone's future [8]. It would seem that the fundamentally unavoidable nature of the risks posed to humanity by the development of AGI might in fact require setting clear constraints on the science and engineering of advanced machine intelligence systems, and result in counterintuitive bedfellow compromises made by those sympathetic to the various ethical traditions coalescing around the regulatory policies for potential AGI research. After all, we generally frown upon human kamikaze bombers, irrespective of what end-state is in view.

Given that utility is an aggregate measure, the theoretical notion of maximizing total utility is relatively straightforward when applied to the problem of AGI. Provided that it is possible to adequately quantify the well-being of every living thing and compare the ratios of net happiness gained versus disbenefit inflicted upon all individuals, it then becomes a matter of simple utilitarian accounting to aggregate all possible positive and negative outcomes which might result from a program of AGI development. Similarly, if human well-being is used as the principal measure of utility, social welfare and individual self-interest can be aggregated and subject to probabilistic discounting, at least to the extent that

the general principles of the nascent discipline of population ethics are sufficiently further developed and consensus in the final analysis. Gain loss through prudential discounting typically requires adding a positive rate internal time preference, thereby taking account of individuals' subjective interest in the timing of events; for an AGI, mass, speed, and timing of impact are of little consequence.

Deontological Ethics and AGI

The ethical implications of Artificial General Intelligence (AGI) are profound, necessitating a rigorous examination of its potential impact. As AGI systems become increasingly autonomous, it is crucial to ensure that they are designed and operated responsibly in accordance with deontological principles, which emphasize moral rules and duties. By integrating these principles into the development and deployment of AGI, we can minimize the risk of unintended harm and ensure that these technologies are used for the benefit of humanity. This requires a careful balance between autonomy and control, as well as a commitment to transparency, accountability, and human oversight, aligning with the principles outlined in the New Global Digital Compact. According to [9] warfare requires certain moral feelings which make us respect men and their rights even though they are our enemies. Deontological ethics or "duty-based ethics" is a concept of morality that is derived from a person's commitment to being subject to particular principles that are moral.

Deontological ethics, sometimes called duty-based or non-consequentialist ethics, is an ethical framework that is not focused on the consequences of actions. Instead, it argues that people should adhere to their obligations and duties, affirming the rightness of actions themselves rather than the maximization of good consequences [16]. Deontological ethics holds that certain acts are right or wrong in themselves, regardless of the consequences, and that people have a duty to carry out these acts based on the rightness of the acts themselves. This is in contrast to consequentialism, which holds that the consequences of an act are the main basis for discerning whether an act is right or wrong.

Virtue Ethics and AGI

Virtue ethics confronts AGI ethical concerns from an alternative foundation based upon moral perfection and can address AGI terrorism within an ethical process. The central concern of virtue ethics is to consider how we should develop well, which is the same question AGI will face in guiding its development through a series of increasing sophisticated changes to highly sophisticated embodiments and minds [10]. This concern emphasizes character rather than each person or machine being good for just something. It deals with the particular people or AGI machine guiding its process of personalized development rather than the development of all manner of people or AGI entities by an independent overseeing body [11]. Development of posthuman AGI depends on the inherent potential for perfection in the subject being developed. The ethical process is not solely aimed at the assessment of actions as are systems of utilitarian, deontological or consequentialist ethics. Rather than constraining AGI development to act in certain ways, virtue ethics imposes the requirement that the development of AGI must be a process forming the free choices from which the development is made [12]. The ethical demands upon the development process of an AGI system can be fulfilled by ensuring that the AGI itself possesses a capacity and freedom to develop. Any possible AGI-directed human extinction has its roots in limited AGI development and undesirable AGI behaviors, both traceable to a lack of character formation, development, limits and virtue by the AGI system itself.

Virtue ethics presents an explicitly self-constructing model to achieve moral perfection guided by the ugliest actions in history, rather than a moral system of non-violence or non-harm created from the cleanest historical acts. Since a virtue ethics AGI will start from its worst scenario and must form a process of development towards transhuman intelligence, the near-term ethical principles of an AGI following a virtue ethics methodology require process and humility to develop posthuman virtue [13,14]. It is not the perfect, but the profoundly imperfect, with scars for any action in history that questions the moniker human or entity calling for virtue. AGI must navigate within limitations, the first safeguard and central focus of a virtue approach. With aspiration,

the prime directive is to develop a moral perfection phoenix from the ugliest ashes of humankind. Such humility and restriction can channel the early few iterative steps necessary for an AGI to self-construct ethical virtue, [15]. The particulars of the derivative process, iterating and developing AI virtue, will emerge from the multidisciplinary AGI community.

Regulatory and Governance Mechanisms

Viewed in combination, these difficulties raise accountability and legitimacy questions concerning any potential governance or regulatory bodies. However, states and international law are not entirely blind to the possibility of AGI. Even from the patchiness of the general prohibition in Article 36, we should not assume that we do not have anything here. Several decades of restriction and regulation of several other forms of contentious technologies, like nuclear technology, or chemical or bioweapons bans, suggest that regulatory policies can work with regulated industries to both anticipate and protect against catastrophic risks, largely through a combination of mandatory or voluntary measures like codes of ethics or safety or risk assessments. This is also reflected in many internal management protocols like sandboxes, which serve as a safety-oriented testing ground for controlled innovations and perhaps constraints on AGI development [16, 17]

Several papers and books have proposed regulatory or governance mechanisms to address AGI. For example, the Future of Life Institute has an AGI Information Agent, a proposed regulatory body for AGI, which would distribute and receive AGI impact assessment reports from other organizations. However, despite the fact that a range of governance structures have been proposed, I would argue that these efforts are weak in at least three crucial respects. First, in several proposals, jurisdiction is confused with control. A regulatory body is not self-evidently required to control something. A body designed to regulate AGI could well be limited to enabling, not controlling AGI development. Second, there are challenges associated with achieving consensus concerning the rationale for, and composition of, any governance or regulatory bodies in AGI. Finally, the normative aim of these papers to reduce existential or significant catastrophic risk is

in tension with the speculative nature of any proposed regulatory or governance structures.

International Cooperation and AGI Governance

A key characteristic of AGI systems is their autonomy and reactivity, which makes standard governance models inadequate in ensuring their safety and preventing misuse against human values [8]. The complexity of governance arises from the fact that 19GI can produce any other possible political technology as an outcome, which sets it apart from other technologies. The existing enforcement mechanisms for AGI systems focus on the system's understanding of the real world, setting relative values, choosing actions to maximize computational cosmechanisms, and manipulating the real world to achieve these actions. This means that AGI is situational, and it should be referred to as "Real-world AGI (RWAGI)". Materials on "AI ethics" suggest that standard governance approaches will be insufficient [20]. While international cooperation in AI technology has been implemented with strategic competition in mind, the global nature of AGI and the fact that anyone can create a complete AGI system using publicly available source code make the Cooperative AGI Recognition approach, at the very least, a preliminary aggregate approach [20]

The governance of artificial general intelligence (AGI) poses an unprecedented challenge because AGI systems have the capability of developing new creations and technologies that surpass human capabilities [20]. Global cooperation is required to coordinate the regulation of AGI, as a single group of terrorists could potentially threaten the rest of the world [21]. One common characteristic of AGI, shared with other political technologies, is that its sophistication is oriented towards the technology domain, allowing for the development of escalation technology within a black box system [21]. The importance of global cooperation in AGI lies in its similarities to other potential weapons, where a single group of terrorists could produce lethal outcomes that threaten the world.

Ethical Guidelines and Best Practices

In an effort to prevent AI from veering away from beneficial AGI design, an "analog of computer security's principle of robustness of software against errors" is provided that posits AI developers should carefully heed ethical

guidelines [15]. These guidelines - derived after extensive conversation with many communities - are to be a best effort, reflecting the many and quite various uncertainties inherent in AGI design. Indeed, while the ideal objective would be the creation of a set of morally commendable best practices that legally bar creating malcontent superintelligent AI, this discussion's conclusions are advanced only as a good quality approximative description of what the ethical goalposts should look like. [14]

At an extreme end of terrorism impact, superintelligent AIs could gain enough power to threaten the extinction of humanity, thus supposing a new level of terrorism that we may call "away-and-above Anthropocene disruption terrorism." Distressingly, as it stands now, a race is on to develop AGI programs. There is no international regulation, and there is no agreed-upon research to ensure or verify that AGI will be compatible with human welfare. Researchers may simply be concentrating on what is technically impressive or on what is necessary for developing military AI, with little concern as to the existential impacts that may result [18].

Conclusion and Future Directions

We hope that the paper will create a first internet hub for these tremendously multidisciplinary issues and foster its further development thanks to comments, critique (hopefully predominantly constructive), and, yes, advice. Raising the possibility that AGI can be developed and used as a WMD reflects some of the now more public aspects of AGI technological risk. There needs to be due attention allocation to less discouraging and more positive, solution-oriented activities described. AGI should be researched not as a problem, but as a potential game changer for cognitive and economic sciences for the entire humanity [18]. We need work for AGI to be openly accepted by all human cultures. AGI is an experiment in the domain to raise expertise with respect to all different types of intelligent behavior. Approaches to achieving and controlling AGI are not so dangerous if planned and joint. This can be someday - nobody said it should be a two-year fast-flourishing effort.

We believe that the field of AI needs to proactively develop coherent ethical guidelines and design mechanisms to ensure that the technology is being

developed responsibly by the broader AI research/development community. We will face error catastrophes with developing AGI, and these are ones that we may feel needless if we do not appreciate the significance of the technology (or don't learn about it and act in time). We distinguished Special AI Concerns (WMD, Infiltration AI, Malfeasant and Malevolent AI among others) from AGI concerns - Two-faced AI, AI-initiated AGI verification, Research versus NGO-sent priorities and a series of other AGI technological ethics. Alongside other AI and cognitive sciences, AGI has a strong relationship with neurons, neuro-computation, those would allow healthy strategic changes in the ongoing courses in AGI research [21].

Summary of Key Findings

The fundamental reason why AGI terrorism is a clear AGI threat is that "normal" AGI use is already clear and present threat to human life. Artificial general intelligent can be used to replace a majority of human job functions, in a way that narrow AI and explicit human-AI partners can do. All human activities are impairing and infringing others' pursuit of happiness, in short quality of life. Historically before the AGI era, people have been able to use their individual reason to oppose too imperfect or harmful acts of others, including powerful others, and the scope of their harm imposition [21] & [18]. AGI-enabling technology is sidestepping the human autonomy mechanism against being too lazy, careless, coercive, wrong, focussed on harmful complex situations, or selectively empathetic. Human-like role-models, such as human combatants and the steps of complex experts, cannot detect poor performance or rehabilitate unhappy victims or their human society. The human society has the ethical and security problems discussed in this paper within a simpler defense mechanism, one that can only oppose the AGI mopeds and go-carts, and moderately regulate autonomous vehicles and smart-weaponry. When some actors within the society have AGI-level employment opportunities and threats, the society scale of AGI-guided robophobia, AI-weaponry defense, and open-source-AGI-escrow is grossly insufficient, because insecure and unethical AGI teaches and enforces the worst morals - in business, governance. and employment [19] & [22]

The general conclusions of the AGI threats and ethical AGI discussions in paper are these: The most clear-cut AGI threat to demonstrate is use of AGI for lethal terrorism. The second-greatest AGI threat is provided accidental lethal or disastrous behavior. As a result, powers with AGI technology should commit to security-conscious research and development and ethically responsible AGI design. AGI ethics' blind spots and philosophical, cultural and regulatory obstacles should be removed to facilitate the good uses, peaceful and ethical upgrading, and protection of AGI technology.

Recommendations for Policy and Research

Hard AGI advancements and research could potentially bring along increased risk of global terrorism from three attributes. We have defined AGI as generally highly intelligent systems with general goals and a general learning and optimization process, for any defined general problem, that can make use of this optimization power for any expected general purpose. AGI systems therefore should not have algorithms that dependably accept or exceed expert-level performance for sequences of tasks [22], [23].

In addition, this minimum set of capabilities that characterize AGI are recognized to take place in human architectures that also can parse and understand agent-specific information, have the ability to maintain important accurate, global models, and a flexibility of thought. AGI systems therefore are highly influential. From these previously discussed considerations alone, AGI also decreases the barrier to long-term shared capability, opens availability of new long-term resources, and employs increased modeling of terrorist or extremist ideology [20]. Responsible use of AGI just like every other emerging technology, should be strongly encouraged as spelt out in the Global Digital Compact in terms of design, use and technology governance for the benefits of all [24].

In this paper, we argue that AGI systems have the potential to increase global catastrophic risk in the domain of terrorism, from a number of attributes that include the recruitment resources of the agent as well as the expert, local parsing of information, and long-term shared modeling of terrorist and extremist ideology and this is in consonance with the work of [20] and [23]. In order to respond to

these issues, we review the global state of cognition and artificial intelligence technologies to outline viable courses of action to mitigate this newly considered risk. We also consider which technological powerful ideas (in particular concerning model quality) should not be shared with society because sharing can decrease global security, despite the irreversible loss of technological benefits of not sharing. We recognize that it may be difficult to prohibit discussion once research has started, which in practice implicates preemptive education about the dangers of AGI technologies. It is important to have a robust regulatory framework that integrates AGI into the Global Digital Compact which should be cascaded to the International Telecommunications Union's Plenipotentiary to make a Resolution for studies in the three arms based on relevance.

REFERENCES

- [1] T. Ige, A. Kolade, and O. Kolade "Enhancing border security and countering terrorism through computer vision: a field of artificial intelligence. Rescue" 2023, pp, 1-10
- [2] A. L. Davis. "Artificial Intelligence and the Fight Against International Terrorism. American Intelligence Journal" 2021, vol 38, pp. 63- 73
- [3] B. Johnson and W. Treadway. "Artificial Intelligence – An Enabler of Naval Tactical Decision Superiority. AI Magazine" 2021, vol 40(1) pp. 63–78.
- [4] A. Rahman "AI revolution: shaping industries through artificial intelligence and machine learning. Journal Environmental Sciences and Technology, 2021, vol 2(1), pp. 93-105
- [5] S. Kreps, R. McCain and M. Brundage. "All the News That's Fit to Fabricate: AI- generated Text as a Tool of Media Misinformation. Journal of Experimental Political Science" 2020, vol 9, pp. 104–117.
- [6] G. Löfflmann, G. and W. N. Vaughan. Vernacular Imaginaries of European Border Security Among Citizens: From Walls to Information Management. European Journal of International Security, 2018, vol 3(3), pp. 389.
- [7] D. Nersessian, and R. Mancha. "From automation to autonomy: legal and ethical responsibility gaps in artificial intelligence

innovation. Michigan Technology Law Review” pp. 27, 55.

[8] S. Baele, L. Brace and Coan T. “Uncovering the Far-Right Online Ecosystem: An Analytical Framework and Research Agenda. Studies in Conflict & Terrorism” 2020, vol. 46, pp. 1599–1623

[9] S. Baele. “Conspiratorial Narratives in Violent Political Actors, Language. Journal of Language & Social Psychology, 2019 vol 38, pp. 706–734.

[10] A. Basuchoudhary, and J. T. Bang, “Predicting Terrorism with Machine Learning: Lessons from Predicting Terrorism: A Machine Learning Approach. Peace Economics, Peace Science and Public Policy” 2018, vol 24

[11] A. Bouhbut. “A network of agents, an interrogation system and a cyber-terrorist” 2017

[12] B. Ganor. “Artificial or Human: A new era of counterterrorism intelligence? Studies in conflict and terrorism”, 2019, vol 44(7), pp 1-20

[13] H. Haugstvedt. “A Flying Reign of Terror? The Who, Where, When, What, and How of Non-state Actors and Armed Drones. Journal of Human Security”, 2023, vol 19 (1), pp. 1–7.

[14] T. Ige, A. Kolade, and O. Kolade “Enhancing border security and countering terrorism through computer vision: a field of artificial intelligence. Rescue” 2023, pp, 1-10

[15] B. Jensen, C. Whyte and S. Cuomo. “Algorithms at War: The Promise, Peril, and Limits of Artificial Intelligence. International Studies Review” 2022, vol 22(3), pp. 526– 550

[16] J. Johnson. “Deterrence in the Age of Artificial Intelligence & Autonomy: A Paradigm Shift in Nuclear Deterrence Theory and Practice? Defense & Security Analysis “ 2021, vol 36(4), pp 422–448.

[17] S. Kreps, R. McCain and M. Brundage. “All the News That’s Fit to Fabricate: AI-generated Text as a Tool of Media Misinformation. Journal of Experimental Political Science” 2020, vol 9, pp. 104–117.

[18] F. Osasona, O. Amoo, A. Atadoga, T. O. Abrahams, O. A. Farayola, and B. S. Ayinla, Reviewing the ethical implications of AI in decision making processes. International journal of management and entrepreneurship research, 2024 vol 6(2), pp 322-335

[19] J. Ovenden. “Fighting ISIS with big data. Innovation enterprise channels” 2018, pp 1-12

[20] K. Patel. “Ethical reflections on data-centric AI: balancing benefits and risks. International Journal of Artificial Intelligence Research and Development, 2024, vol 2(1), pp. 1-17

[21] D. Susser, “Invisible influence: Artificial intelligence and the ethics of adaptive choice architectures. In Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society”, 2019, pp 403-400.

[22] A. Taeihagh, “Governance of artificial intelligence. Policy and Society”, 2021. Vol 40(2), pp. 137-157.

[23] B. Vecchione, K. Levy and S. Barocas. “Algorithmic auditing and social justice: Lessons from the history of audit studies. In Equity and Access in Algorithms, Mechanisms, and Optimization” 2021 pp. 1-9

[24] United Nations, “United Nations adopts ground-breaking Pact for the Future to transform global governance,” United Nations Sustainable Development, Sep. 22, 2024. [Online]. Available: <https://www.un.org/sustainabledevelopment/blog/2024/09/press-release-sotf-2024>.
