

Do we have Procreative Obligations to AI Superbeneficiaries?

Sherri Conklin, University of California Santa Barbara, (United States)
International Conference on Computer Ethics: Philosophical Enquiry (CEPE) 2023,
Chicago, IL

Short Abstract: This paper concerns itself primarily with questions about our obligations to AI superbeneficiaries – entities with inherently valuable interests that exceed those of humans in terms of quality and/or quantity. Specifically, this paper deals with questions about whether we have any obligations to bring AI superbeneficiaries into existence, especially if it turns out that human well-being might very well be at stake. I employ an anti-natalist argument to establish that we have all-things-considered moral obligations against bringing AI superbeneficiaries into existence because of the existential risk they pose to their own survival as well as to the survival of humanity.

Keywords: AI Ethics, Superbeneficiaries, moral status, anti-natalism

Introduction

This paper employs an anti-natalist argument to establish that we have moral obligations against bringing AI superbeneficiaries into existence because of the existential risk they pose to their own survival as well as to the survival of humanity. AI superbeneficiaries are entities with inherently valuable interests that exceed those of humans in terms of quality and/or quantity. According to some ethical theories, AI superbeneficiaries threaten the future well-being of humans because, if we do bring AI superbeneficiaries into existence, we will most likely have moral obligations to promote the interests of these AI at the expense of human interests. Moreover, it could turn out that, once AI superbeneficiaries have been created, humans will have additional obligations to bring many more into existence (i.e., we will have procreative obligations to AI superbeneficiaries). I argue, on anti-natalist grounds, that, even if we have some *pro tanto* procreative obligations to these entities, we have all-things-considered obligations against bringing AI superbeneficiaries into existence.

AI Superbeneficiaries

One might reasonably hypothesize that artificial intelligences will eventually gain moral status. An entity has moral status insofar as its interests have inherent value [1]. Ethicists have long debated over which interests have the most inherent value, but there seems to be consensus around a set of physical and psychological capabilities, including (but by no means limited to) the capacity to derive enjoyment from life (e.g., to experience happiness or pleasure) or the capacity to suffer (i.e., to experience pain), as well as the capacity to engage in rational deliberation about what is good for us and others, and the capacity to exercise our will in the pursuit of those goods without interference from others [2, 3].

The strength of humanity's different obligations to an entity with moral status seems to depend on which interests have the most inherent value, as well as the degree to which the relevant interests are properly attributable to an entity [4]. An entity might, for example, be more or less susceptible to pleasure and pain, more or less rational, and more or less concerned with the extent to which its plans are thwarted. Where an entity lacks the relevant interests, we would have fewer and weaker obligations relating to those interests whereas we would have more and mightier obligations towards an entity presenting with those interests. At present, we have about as much reason to think that emerging AI will share few human-interpretable interests and possess correspondingly low moral status as we have to think that the quality and quantity of AI interests will far exceed our own – perhaps radically changing our notions of what it means to be a full moral agent and relegating humanity's interests to secondary consideration, in the same way that we often treat the interests of animals as secondary to our own.

With regard to the first point, we might think this because AI will be equipped with physiological and psychological structures that are different from that of humans. Although we have good reason to think mammals share many of our interests because their nervous systems are similar to our own, an AI might be equipped with physiologies sharing no such resemblance. For example, while an AI would almost certainly rely on a physical substrate, an AI mind that is distributed and

duplicated across millions of physical devices would most likely be able to tolerate decimation before experiencing a decline in capacity, might not ever become aware of any such damage or might be able to rely on back-up versions of itself before any such damage began to matter. Pain is a tool that animals developed for relaying information about bodily damage to our brains, but it seems unlikely that a distributed intelligence would have any need to develop this capacity, even if equipped with a suitable sensory apparatus, since extensive damage to an AI's physical structures might be of relatively little consequence to its survival or overall well-being.⁵ Such entities might not, then, have interests relating to pain or pleasure, and perhaps there will be similar differences between AI and humans in other areas relevant to developing interests with inherent value that would further reduce the moral status of AI.

Regarding the second point, some AI might instead develop into what Shulman & Bostrom [2020] call superbenebeneficiaries – entities with inherently valuable interests that exceed those of humans in terms of quality and/or quantity and which obtain correspondingly greater moral status as so-called superpatients [6]. While humans might have relatively few obligations towards AI as entities with low moral status (discussed above), we would most likely have more and mightier obligations towards AI superbenebeneficiaries, perhaps more obligations than we have towards our fellow humans. As a result, we could end up with scenarios where the well-being of AI superbenebeneficiaries should be promoted at the expense of human well-being, with the further risk that human well-being might be permissibly neglected entirely if such AI were sufficiently efficient at deriving well-being from the world's resources.

Procreative Obligations to AI Superbenebeneficiaries

This paper concerns itself primarily with questions about our obligations to AI superbenebeneficiaries. Specifically, this paper deals with questions about whether we have any obligations to bring AI superbenebeneficiaries into existence, especially given the concern that human well-being might very well be at stake [7]. Hereafter, I refer to any obligations to bring an entity into existence as procreative obligations. Whether we have any *prima facie* procreative obligations to AI seems to hang on our views about which interests have the most inherent value. If we follow Kantian-style reasoning about the inherent value of exercising our capacities for rational deliberation about the good, the answer is probably that we have no such obligations. Kantian obligations derive from the requirement to promote the rational well-being of existing entities with moral status and most likely extend to possible future entities only insofar as they guide actions that affect the rational well-being of entities that will indeed be brought into existence [8, 9].¹

If we instead hold that the greatest inherent value stems from the capacity to derive enjoyment from life (e.g., to experience happiness or pleasure), then we may very well have some such obligations when AI are superbenebeneficiaries with regard to this set of interests. On a generic Utilitarian-style view for example, we have obligations to increase and maximize the aggregated

¹ This is not, nor is it intended to be, a careful articulation of Kant's view or likely stance on the topic. This paper focuses on consequentialism because it would be relatively easy to generate procreative on some consequentialist accounts, which is all that is really needed to progress this argument.

quantity and quality of pleasure in existence regardless of what sort of entity is lucky enough to enjoy the pleasure [10]. If an AI is a pleasure superbeneficiary, then we may have some procreative obligations to such entities in order to increase and maximize the world's aggregated quantity and quality of pleasure. [6] propose a number of possible conditions that might generate AI superbeneficiaries. I consider three for illustrative purposes.

First, in contrast with the above point about AI having no need to develop a capacity for pleasure or pain because of their nature as distributed intelligences, it could also turn out that they develop superhuman capacities for pleasure or pain. Human developers, for example, frequently use rewards systems for training AI. We have no reason to think that AI are unable develop psychological responses to rewards so closely approximating human-like pleasure as to render any differences irrelevant. Humans might also equip AI with a sensory apparatus appropriate for experiencing something approximating physical pleasure and using that apparatus as a part of a rewards system for training the AI. While humans have biological limitations on the quantity and quality of pleasure we can experience, humans could potentially develop AI superbeneficiaries with psychologies and physiologies evidencing no such limitations.

Further, due to differences in physical construction, the subjective speed at which AI experience pleasure may also surpass the subjective speed at which humans experience pleasure. As a result, an AI might subjectively experience many greater quantities of pleasure than human counterparts over the same objective time scale. If so, an AI could obtain pleasure superbeneficiary status, as compared to humans, through sheer quantity of pleasure it experiences over its lifetime. Moreover, the increased subjective speed at which an AI experiences pleasure may well stack with the kind of unlimited capacity for experiencing pleasure mentioned just above.

Finally, AI afford a greater reproductive capacity as compared to humans. For example, we can replicate innumerable AI "offspring" from a single originating AI entity. The only limit on the number of AI duplicates we can produce is the amount of available computational resources. In theory, an AI population could exceed the size of our own and achieve superbeneficiary status as a population (i.e., through sheer quantity). Parfit notes that, while it may be impossible for humans to imagine an entity that experiences a million times greater pleasure than us, it may nonetheless be possible for us to imagine a million entities that experience the same amount of pleasure as humans. These entities, as a population, might introduce sufficient pleasure into the world as to create procreative obligations.

Moreover, while any of these three considerations might result in procreative obligations to AI superbeneficiaries, it seems that AI entities with all three characteristics are possible. It could turn out that the original AI progenitor is also capable of higher quality pleasure experiences at higher subjective speeds (as hypothesized above). If so, the aggregated quantity and quality of pleasure brought into the world by AI superbeneficiaries might outstrip the value of human experiences to such an extent that our well-being should hardly be considered in the utility calculus thereby relegating the interests of humanity to the domain of secondary moral consideration. So, [6] argue, it seems possible for us to have procreative obligations to AI superbeneficiaries and for us to have serious moral reason to put their interests ahead of our own.

This seems like a potential problem for humanity – a problem with ethical consequences that we should seriously consider before proceeding with the development of AI with moral status. Supposing, however, that we do proceed with the development of AI with moral status, and supposing that those AI are superbenebeneficiaries, we will need to contend with possible procreative obligations [11]. I argue that, while we would most likely have procreative obligations to AI superbenebeneficiaries that can be overridden, we do not have procreative obligations to AI superbenebeneficiaries *ceterus parabus*. Moreover, I argue that we may have all-things-considered obligations against bringing AI superbenebeneficiaries into existence on anti-natalist grounds.

A Primer on Anti-Natalist Philosophy

On standard views, anti-natalism is a socio-ethical ideology according to which human procreation is morally wrong. For one reason or another, anti-natalists argue that we should decrease the size of the human population (locally or globally), often by implementing state-level policies aimed at discouraging the exercise of procreative liberties. Anti-natalism is, perhaps, one of the single most controversial ethical stances in existence because it is associated with a litany of morally repugnant social practices, is controversial on religious grounds, and because it argues against the fundamental moral value of life and its goods. However, on certain assumptions about what would be good for humanity as a species, some anti-natalist arguments are almost certainly correct. I consider these points in a little more detail, in order to fill out the sort of morally respectable anti-natalist argument I am interested in, before applying anti-natalist arguments to the problem of AI superbenebeneficiaries.

The motivations behind different forms of anti-natalism are wide-ranging, yet I have identified approximately four primary categories of anti-natalism: (1) pessimistic misanthropic; (2) optimistic misanthropic; (3) pessimistic philanthropic; (4) optimistic philanthropic. According to misanthropic anti-natalism, humanity, as a species, does not deserve to exist because of the amount of evil that we have brought into the world [12]. Views of this sort rally around the vast body of evidence that humanity, as a species, is seriously morally flawed. We perpetrate acts of harm, suffering, and death against ourselves, non-human animals, and our environment without any clear evidence that our species is likely to ever stop acting thusly. On pessimistic accounts (such as those proposed by The Voluntary Extinction Movement) [13], we are therefore irredeemable and should permanently stop bringing additional humans into existence. On optimistic accounts, we are redeemable but must limit human reproduction to a much smaller and highly selective population [14]. According to philanthropic anti-natalism, humanity, as a species, should stop reproducing due to how much evil there is in the world. The misanthrope's depiction of the world, as nasty, brutish, and short (following Hobbes), is largely endorsed by pessimistic philanthropic anti-natalists. Unlike misanthropes however, they hold that people should not be expected to endure such protracted suffering and that the best thing to do, to relieve human suffering, is to permanently stop bringing additional humans into existence [15,16,17, 18]. The optimistic philanthropist holds that the problem is largely due to the size of the human population, rather than human nature, and that we should temporarily stop reproducing until we reach a sustainable population size.

Anti-natalism is notoriously controversial. State-level policies aimed at discouraging the exercise of procreative liberties are often coercive [19]. Enacted in the 1970's, India's mass sterilization program during "The Emergency" and China's One Child Policy have been cited as affronts to women's reproductive autonomy – though both are long deprecated as of 2023. Such programs have been scrutinized for human rights violations [20], and these concerns arrive at the intersection of other morally repugnant, anti-natalist social practices.

Eugenics, in particular, often springs to mind when anti-natalism enters public discourse – and for good reason. Optimistic misanthropic anti-natalism is consistent with (though, not necessarily essential to) socio-ethical practices that fixate on the purported evils, such as an imagined propensity for violence and criminal activity, or other limitations (e.g., intelligence) of particular populations of people, usually based on income, race, ethnicity, religious affiliation, or sexual identity [21, 22]. Proponents identify these populations as the primary source of humanity's moral flaws and hold that decreasing the size of these populations, if not eliminating them entirely, will decrease the amount of evil in the world. However, one need not commit to such extreme ideologies for eugenics to be a concern for anti-natalist policies. When India's Population Control Bill is juxtaposed with a growth in reproductive tourism that exploits lax regulations on transnational surrogacy, we can see how the reproductive viability of people from wealthy nations is rapidly advanced, while that of the people of India and other exploited regions is decreased [23, 24, 25].

Anti-natalism is also controversial on religious grounds. First, it directly contraindicates reproductive behaviors promoted by pro-natalist myths present in the religious texts underpinning Judeo-Christian and Islamic traditions [26]. Second, anti-natalism also conflicts more generally with some religious attitudes towards women's health practices [27]. For example, abortion might be morally required on some anti-natalist views, where those with religious commitments favoring pro-natalism might instead see forced birth as the only permissible reproductive option. While both views might hold that women's reproductive autonomy is outweighed by other moral considerations, anti-natalism is more likely to favor autonomy-enhancing prophylactic measures.

Lastly, anti-natalism seems to argue against the fundamental moral value of life and its goods. The truth and significance of this claim seems to depend on which category of anti-natalism is under consideration. Categories (1) – (3) would seem to hold that existence is morally worse than non-existence – at least for some people. On such views, there is too much evil in the world for life to be worth living (or worth living without eliminating large portions of the human population). Proponents of optimistic misanthropic anti-natalism might hold that some lives are more fundamentally valuable than others and that only these people are deserving of life's goods. Pessimistic philanthropic anti-natalists might hold that all lives are fundamentally morally valuable but that life bestows more harms than goods. However, optimistic philanthropic anti-natalists might hold that we need to balance the fundamental value of life with the goods that are derivable from life when greater numbers of people significantly decrease the availability of life's goods.

This latter sort of anti-natalism is the one I am interested in this paper. In particular, the strongest optimistic philanthropic anti-natalist accounts argue that humanity poses an existential risk to itself as a result of anthropogenic climate change. On this view, the fundamental moral value of humans,

as a species, generates moral obligations to decrease fertility rates in order to ensure our long-term existence and prevent large amounts of human suffering. Hickey, Rieder, and Earl [2016], for example, argue that decreasing human fertility rates, by means of non-coercive population engineering, is one of the most effective (if not the most effective), morally justifiable means of obviating the urgent existential risk that climate change poses to humanity [28]. The conclusion follows, in part, from the empirical fact that humans are largely responsible for climate change and that no other means of attenuating the catastrophic harms of climate change is likely to have the same level of impact as reducing the size of the human population in the short time we have before the damage is irreversible. This issue is frequently of concern to prospective parents facing these threats [29]. If so, humanity has moral obligations to stop reproducing (presumably until a certain threshold is reached) in order to secure its survival – on the assumption, of course, that there is a moral imperative for humanity to survive.

An Anti-Natalist Argument Against Procreative Obligations to AI Superbeneficiaries

The existential threat that humans pose to our own survival, via our impact on climate change, is directly analogous to the existential threat that AI superbeneficiaries pose to themselves if we have procreative obligations towards them. For the most part, the line of reasoning supporting the argument presented by [28] applies to the case of AI superbeneficiaries. I argue that bringing AI superbeneficiaries into existence will result in comparable harms to the environment and to themselves while depleting the natural resources they require to survive. If so, then we have similar moral obligations to forestall bringing any such entities into existence, and I contend that these obligations override any (weaker) procreative obligations that we may have to AI superbeneficiaries.

Consider, in rough outline, some non-exhaustive, *prima facie* reasons for why this is so. At the present stage of our civilization, the way humans develop, manufacture, and use technology is resource inefficient and damaging to the environment [30, 31]. Current practices, along all three measures, incur massive energy expenditures and a correspondingly high carbon footprint (e.g., more devices means more energy use). Activities like crypto mining can be wasteful in other ways [32]. For example, the chips used for computation are sometimes obtained from larger devices, with other components, that remain unused or are discarded, and there is high-turnover in device use, which results in greater quantities of physical waste, due to use inefficiencies, and e-waste [33]. Another byproduct of modern technology is the release of toxic substances into the environment. For example, the extraction and purification process involved in obtaining the rare earth materials used in manufacturing many cutting edge technologies can result in contaminating local soil and water with radionuclides [34]. E-waste can similarly leak metals like mercury, lithium, and cadmium [35].

Presumably, AI of any sort should be numbered among these human technologies, which is a problem for AI superbeneficiaries. To manufacture and power the kind of physical substrate in which such entities are embedded, we require rare earth metals. Rare, in this context, means difficult to find and, frequently, more difficult to extract. If we have obligations to bring an ever-

increasing number of AI superbeneficiaries into existence, we likely have derivative obligations to intensify extraction of these finite resources. AI are unlike humans in that, while we benefit from a range of rare resources, we do not need them for our survival. All of the resources that humans require for survival would be sustainable and renewable with a sufficiently reduced human population. Once we extract the last of these materials, we will be unable to support the creation of new AI superbeneficiaries.

Even if we are not required to intensify mineral extraction and are only required to create a sustainable population of AI superbeneficiaries, we would still inevitably run out of materials – only at a slower rate. One might contend that we could always develop carbon neutral recycling facilities to reclaim the materials needed for creating AI superbeneficiaries. While this is a reasonable aspiration, we know that humanity has struggled to recycle many highly recyclable materials, such as glass [36]. Inefficiencies over the lifetime of our technology production and use cycles would most likely render resource reclamation programs ineffective, at least on the time scale required for sustaining a large AI population and before the catastrophic climate change that poses an existential threat to humanity.

Like all technology, the physical substrate in which the AI are embedded will degrade with time. Anyone who has had to replace a laptop because it reached the end of its usefulness should be familiar with this. Complex AI require much more computational power than our laptops, so they require much more in terms of physical resources, require more redundancy, and result faster turnover. Because of the computational intensity of maintaining complex AI, we would cycle through necessary materials at a quick pace. If we were unable to reclaim materials at a faster pace, the physical substrate of the AI would degrade, leading, essentially, to their death.

Like humans, AI of any sort must also contend with the dangers of catastrophic climate change. As already noted, the way humans develop, manufacture, and use technology is damaging to the environment. The development and maintenance of AI, as one of these technologies, contributes to the climate problem. As a result, each AI will leave its own environmental footprint. If we are required to bring an ever-increasing number of AI superbeneficiaries into existence, the AI will likely contribute to the approaching climate catastrophe in much the same way that humans do.

This is a problem for AI because they are similarly sensitive, if not more sensitive, to natural disasters. Consider that spilling a cup of water on one's computer is sometimes enough to irreparably damage it. AI are typically embedded in large networks of computers all of which can suffer similar sorts of damage. Of course, they can be protected to some extent, but a tsunami, or an earthquake, or a wildfire can very likely bring about their destruction. The physical infrastructure supporting AI, such as power grids, are also susceptible to extreme environmental conditions, but, unlike people, this physical infrastructure is not easy to relocate and predicting the best new location might be difficult given the unpredictability of environmental conditions as earth's climate changes [37, 38].

Conclusion

This is all to say that AI superbeneficiaries will most likely suffer the same fate as humanity if we have all-things-considered procreative obligations to them. I think the above anti-natalist stance suggests that we do not, and that we have more and mightier obligations to the contrary. If humanity has all-things-considered obligations to, perhaps temporarily, stop reproducing in light of the existential risk we pose to ourselves, then humanity most likely has the same sort of obligations against bringing AI superbeneficiaries into existence – at least where the nature and cause of the existential risk are analogous.

-
- [1] A. Jaworska and J. Tannenbaum. 2021. The grounds of moral status. *The Stanford Encyclopedia of Philosophy* (Spring 2021 Edition), Edward N. Zalta (ed.). Retrieved from <https://plato.stanford.edu/archives/spr2021/entries/grounds-moral-status>
- [2] S. Buss. 2012. The value of humanity. *Journal of Philosophy*, 109, 341–377.
- [3] I. Robeyns and M. F. Byskov. 2021. The capability approach. *The Stanford Encyclopedia of Philosophy* (Winter 2021 Edition), Edward N. Zalta (ed.). Retrieved from <https://plato.stanford.edu/archives/win2021/entries/capability-approach>
- [4] D. Degrazia. 2008. Moral status as a matter of degree? *Southern Journal of Philosophy*, 46, 181–198.
- [5] E. T. Walters and A. C. Williams. 2019. Evolution of mechanisms and behaviour important for pain. *Philosophical Transactions of the Royal Society*. 374, 1875. DOI: <https://doi.org/10.1098/rstb.2019.0290>
- [6] C. Shulman and N. Bostrom. 2021. Sharing the world with digital minds. *Rethinking Moral Status*, 306–326.
- [7] N. Bostrom. 2013. Existential risk prevention as global priority. *Global Policy*, 4, 1, 15–31.
- [8] C. Korsgaard. 1996. Kant's formula of humanity. *Creating the Kingdom of Ends*. Cambridge, Cambridge University Press, 106–132.
- [9] C. Korsgaard. 2004. Fellow creatures: Kantian ethics and our duties to animals. *The Tanner Lectures on Human Values*, Vol. 25/26, G.B. Peterson (ed.), University of Utah Press, Salt Lake City, 79–110.
- [10] J. S. Mill. 1861. *Utilitarianism*. Roger Crisp (ed.), Oxford University Press, Oxford (1998).
- [11] J. Stone. 1987. Why potentiality matters. *Canadian Journal of Philosophy*, 17, 815–829.
- [12] D. Benatar. 2015. The misanthropic argument for anti-natalism in Sarah Hannan, Samantha Brennan, and Richard Vernon (eds), *Permissible Progeny? The Morality of Procreation and Parenting*. New York, Oxford Academic, 34–64.
- [13] J. S. Ormrod. (2011). Making room for the tigers and the polar bears: Biography, phantasy and ideology in the Voluntary Human Extinction Movement. *Psychoanalysis, Culture & Society*, 16, 142–161.
- [14] S. L. Thomas. (1998). Race, gender, and welfare reform: The antinatalist response. *Journal of Black Studies*, 28, 419–446.
- [15] David DeGrazia. (2010). Is it wrong to impose the harms of human life? A reply to Benatar,” *Theoretical Medicine and Bioethics*, 31, 317–332
- [16] J. Marsh. (2014). Quality of life assessments, cognitive reliability, and procreative responsibility. *Philosophy and Phenomenological Research*, 89, 436–466.
- [17] S. V. Shiffrin. (1999). Wrongful life, procreative responsibility, and the significance of harm,” *Legal Theory*, 5, 117–148.
- [18] J. Feinberg. (1992). Wrongful life and the counterfactual element in harming. *Freedom and Fulfillment: Philosophical Essays*. Princeton, Princeton University Press.

-
- [19] C. Follett. (2020). Neo-Malthusianism and coercive population control in China and India. *Policy Analysis*, 897.
- [20] W. Feng, Y. Cai, and B. Gu. (2013). Population, policy, and politics: How will history judge China's one-child policy?. *Population and development review*, 38, 115-129.
- [21] P. Cole. (2006). *The myth of evil: Demonizing the enemy*. Greenwood Publishing Group.
- [22] J. D. Duntley and D. M. Buss. (2004). The evolution of evil. *The social psychology of good and evil*, 102-123.
- [23] C. M. Mishra and S. Paul. (2022). Population control bill of Uttar Pradesh (two-child norm): An answer to population explosion or birth of a new social problem?. *Journal of Family Medicine and Primary Care*, 11, 4123-4126.
- [24] A. Pande. (2016). Global reproductive inequalities, neo-eugenics and commercial surrogacy in India. *Current Sociology*, 64, 244-258.
- [25] D. Deomampo, (2013). Gendered geographies of reproductive tourism. *Gender & Society*, 27, 514–537. <https://doi.org/10.1177/0891243213486832>
- [26] A. Thompson. (2016). Debating procreation: Is it wrong to reproduce?, *Analysis*, 76, 556–558.
- [27] F. L. Brown. (2021). Anti-natalism. *Encyclopedia of Evolutionary Psychological Science*, 319-321, Cham: Springer International Publishing.
- [28] C. Hickey, T. N. Rieder, and J. Earl. 2016. Population engineering and the fight against climate change. *Social Theory and Practice*, 42, 4, 845-870.
- [29] M. Schneider-Mayerson and K. L. Leong. (2020). Eco-reproductive concerns in the age of climate change. *Climatic Change*, 163, 1007-1023.
- [30] M. Höök and X. Tang. 2013. Depletion of fossil fuels and anthropogenic climate change — A review. *Energy policy*, 52, 797-809.
- [31] C. Rosenzweig, D. Karoly, M. Vicarelli, P. Neofotis, Q. Wu, G. Casassa, A. Menzel, T. L. Root, N. Estrella, B. Seguin P. and Tryjanowski. 2008. Attributing physical and biological impacts to anthropogenic climate change. *Nature*, 453, 7193, 353-357.
- [32] L. Badea and M. C. Mungiu-Pupăzan, (2021). The economic and environmental impact of bitcoin. *IEEE Access*, vol. 9, 48091-48104. doi: 10.1109/ACCESS.2021.3068636.
- [33] A. De Vries and C. Stoll. (2021). Bitcoin's growing e-waste problem. *Resources, Conservation and Recycling*, 175.
- [34] D. Talan and Q. Huang. (2022). A review of environmental aspect of rare earth element extraction processes and solution purification techniques. *Minerals Engineering*, 179.
- [35] R. I. Moletsane and C. Venter. 2018. Electronic waste and its negative impact on human health and the environment. *International Conference on Advances in Big Data, Computing and Data Communication Systems (icABCD)*, Durban, South Africa, 1-7. doi: 10.1109/ICABCD.2018.8465473.
- [36] J. H. Butler and P. D. Hooper (2019). Glass waste. *Waste (Academic Press)* 307-322.
- [37] T. Kemabonta. (2021). Grid Resilience analysis and planning of electric power systems: The case of the 2021 Texas electricity crises caused by winter storm Uri (# TexasFreeze). *The Electricity Journal*, 34.
- [38] A. Kwasinski, F. Andrade, M. J. Castro-Sitiriche and E. O'Neill-Carrillo, E. (2019). Hurricane Maria effects on Puerto Rico electric power infrastructure. *IEEE Power and Energy Technology Systems Journal*, 6, 85-94.