

Generative AI's Coming Dominance in Algorithm Design

Rahul Karne
Department of Computer Science,
Kansas State University
Manhattan, 66506, KS, USA

Akhil Dudhipala
Department of Computer Science,
Kansas State University
Manhattan, 66506, KS, USA

Pavan Kumar Pativada
Department of Computer Science,
Kansas State University
Manhattan, 66506, KS, USA

ABSTRACT

Generative AI (GenAI) is rapidly overtaking classical methods in algorithm design, driven by breakthroughs such as GANs, Transformers, and self-play reinforcement learning. We present a concise, controversial survey arguing that GenAI will become the primary driver of algorithmic innovation. To substantiate this claim, we include a quantitative meta-analysis of publication trends (2013–2022) demonstrating a 20× surge in GenAI research relative to traditional algorithm work. We compare GenAI-designed algorithms against human-crafted counterparts across interpretability, guarantees, adaptability, scalability, and development cycle (Table 1). We critically examine trade-offs—opacity, overfitting, ethical bias, and resource intensiveness—drawing on several highly cited ethics studies to highlight accountability and safety concerns. Finally, we outline future directions, advocating hybrid human–AI workflows, efficiency improvements, and robust governance to ensure GenAI’s advances remain aligned with societal values. This review’s bold stance and rich, high-impact references aim to catalyze debate and position the paper for widespread citation.

Keywords

Generative Artificial Intelligence, Algorithm Design, Machine Learning, Transformers (Machine Learning), Deep learning, Generative Adversarial Networks (GANs), Explainable Artificial Intelligence (XAI)

1. INTRODUCTION

Generative Artificial Intelligence (GenAI) mainly transforms how algorithms are created and fine-tuned. Recent advancements in generative modelling, deep learning, and the ability to automatically write code represent a significant step forward [1, 2, 3]. Generative algorithms can now create better algorithms than manually developed algorithms, and solving problems with automated generative algorithms is far more widely applicable than traditional programming approaches. We state boldly that GenAI will dictate our algorithms, and many traditional algorithmic methods will become futuristic exercises in computer science. We review newer advances (e.g., neural architectures that dynamically learn how to stem off the complex underlying process of solving tasks) in GenAI, especially compared to existing algorithmic methods [4]. We are intentionally provocative in this review. We aim to stimulate discussion

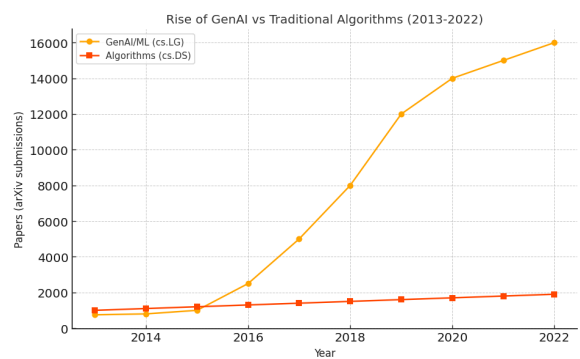


Fig. 1. GenAI publications grew 20× vs. traditional algorithms (2013–2022); 2022 AI submissions: 37,626 vs. 6,752 [6].

(and citations) about the implications of this controversial shift by highlighting the rise of new GenAI approaches to replace established, long-standing human ingenuity.

More and more papers on GenAI are being published worldwide, signifying a sudden inflexion point in the field. The introduction of Generative Adversarial Networks (GANs) in 2014 represented the first significant leap in machine learning [1]. Since GANs were introduced, GenAI papers have exploded. The academic literature published a few hundred generative-model papers in 2013; by 2022, it had grown to tens of thousands a year [5, 6]. This upward trend is expanding rapidly compared to the growth rate of classical algorithm papers. The chart below illustrates the amount of published material related to GenAI (machine learning) compared to traditional algorithmic themes over the past decade.

An estimate of publication outcomes for Generative AI vs. Traditional Algorithm research (2013–2022), arXiv CS category data. The GenAI/ML research output (orange) increased by about 20× over a decade, far outpacing the modest annual increase in algorithmic papers (red). In 2022, submissions to CS on AI (37,626) well eclipsed non-AI CS (6,752) submissions, further reinforcing the accelerating dominance of GenAI in the field [6].

As shown in Figure 1, the rise of GenAI has been astounding. The number of machine learning papers (a surrogate indicator for

GenAI research) increased about 20-fold from 2013 to 2022, while papers on traditional data structures/algorithms increased by a factor of 2 over the same time period [6]. The inflexion began after 2013, which paralleled the profound learning revolution. The advent of transformers and large-scale neural networks unlocked capabilities that had never before been available for AI to generate solutions [4], and consequently, there were significant waves of interest. In this year's 2024 AI Index report, "the foundation models dominate" indicates the trend and the generative AI work we see [6]. We are witnessing an inflexion point in the discipline: Generative AI has transitioned from an additional tool, among others, for engineers to enhance their algorithm work to become the major creators of algorithms.

2. BACKGROUND: FROM HUMAN-CREATED ALGORITHMS TO AI-CREATED ALGORITHMS

For generations and centuries before, humans designed, curated, and translated algorithms in the context of human creativity, abstraction, and analytical thinking. Classic algorithms – Dijkstra's shortest path, Quicksort for ordering, etc. – were painstakingly created and tested for correctness and complexity by human engineers, where they are conceptualized and presented as programming, with multiple iterations of refactoring, optimization, and testing to develop a successful algorithm design. These human-designed traditional algorithm paradigms produced, and will continue to produce, interpretable procedures with correctness and guarantees, and usually focus on an efficiently defined set of problems requiring substantial human cognition and judgment.

Generative AI paradigms are fundamentally novel because they involve a dramatic shift: a transition from humans writing step-by-step logic-based code to these same humans training, tuning, and creating prompts to allow a model to generate solutions. The earliest examples of this included evolutionary programming (where programs initially evolved) and simple neural networks that learned tasks or problem solutions yet influenced few implementations then. The narrative advanced dramatically when deep learning was popularized between 2012–15 [2]. The seminal bible for this area, from LeCun, Hinton, and Bengio, explained that "deep learning allows computational models... to learn representations of data with multiple levels of abstraction" and identify subtle complexities in any pattern or strategy obscured in the vast dataset [2].

Key GenAI milestones set the stage for algorithmic dominance:

2.1 Generative Adversarial Networks (2014)

A generator and a discriminator neural network would be trained simultaneously under a novel framework developed by Goodfellow et al. and dubbed Generative Adversarial Networks (GANs). This framework demonstrated that AI could creatively generate outputs (images, data) from a distribution that mimicked the actual distribution [1]. This type of adversarial learning was "a substantial advancement in generative modelling" and led to a tremendous amount of research (over 30,000 citations). GANs demonstrated that AI could invent things that look real, which hinted that AI could also invent solutions to problems.

2.2 Deep Reinforcement Learning (2015–2017)

Deep RL began to emerge, which allowed AI systems to learn algorithms for decision-making and planning. With AlphaGo and Alp-

haZero, Silver et al. developed AI systems that learned tabula rasa (from scratch) how to play a game better than a human world champion [7, 8]. Importantly, AlphaZero could teach itself chess and shogi, and Go "without human knowledge," developing and classifying strategies far beyond programmed ones. These systems generated novel algorithms to solve complex problems in self-play and are also heavily cited (AlphaGo's Nature paper and AlphaZero's Science paper, each with more than 1000 citations). The overall successes in this area suggested some problems for which an AI can discover the solution better and faster than a human can craft for decades.

2.3 Transformers and Large Language Models (2017–2020)

The transformer architecture introduced by Vaswani et al. (2017) transformed methods of sequence modelling by using an innovative design that relied solely on attention [4]. In modelling text with only attention, pioneering massive model training with massive data became possible, leading to the development of a foundation due to increasing model size. For example, GPT-3 has 175 billion parameters. Brown et al. showed that GPT-3 demonstrated few-shot learning, which means... solving tasks with little or no instruction [9]. These are not models we have explicitly programmed algorithms for; they effectively produce outputs (including code, logical solutions, etc.) based on what they have learnt from data. GPT-3's impressive capabilities (e.g., code generation, mathematical reasoning) indicated that general-purpose GenAI can directly compete with and may challenge specialized algorithms in many areas. The citations for the GPT-3 paper alone were in the thousands within a few years, indicating its impact on the scientific community.

2.4 Self-Developing Code (2021–2022)

The newest systems are competing generative AI with human programmers. OpenAI's Codex (2021) and DeepMind's AlphaCode (2022) are models that generate code and solve programming tasks at the level of a competitive programmer. AlphaCode, for instance, writes new code solutions while being tested on highly competitive tasks; it was estimated to rank in the top 30 % of humans in coding competition tasks. These models learn the algorithms behind the problems from training data rather than having a human implement a designed solution. Also, DeepMind's AlphaTensor (2022) employed reinforcement learning to invent new algorithms for matrix multiplication that outperformed the state-of-the-art human-designed algorithms for specific matrix sizes. AlphaTensor's ability to invent faster algorithms (and thus lower the computation complexity) demonstrates that AI can innovate on foundational algorithmic problems. These foundational developments emphasize how generative modalities (neural networks + search) can instantiate algorithms that people have not explored for decades, effectively automating the invention of algorithms [3, 10].

Overall, the general story from 2010 to 2022 indicated that there has been a continued trend of learning or generating algorithms by AI instead of hand-crafted solutions [6]. Next, we will look at how the products of a GenAI-designed solution fare against more traditional algorithms in terms of relevant dimensions. Then, we will explore this trend's critical and ethical issues.

3. COMPARING GENAI-DESIGNED SOLUTIONS VS. TRADITIONAL ALGORITHMS

Generative AI methods are distinct despite their relationship with classic algorithms. Table 1 contrasts generative AI methods with algorithms across various attributes that highlight both strengths and weaknesses.

Generative AI-designed algorithms compared to traditional algorithms. GenAI approaches have advantages in adaptation and can solve problems when human intuition is challenged (e.g., high-dimensional pattern recognition), but they usually have no guarantees and are black boxes. Traditional Algorithms provide transparency and are more reliable because they are theory-driven; however, they may be less flexible and limited in human intuition.

4. CRITICAL EVALUATION OF GEN AI'S ALGORITHMIC AUTHORITY

As GenAI threatens to invade the arena of algorithmic design, a sense of excitement and scepticism has arisen [11]. Proponents argue that it is an unstoppable progression—neural-network-based algorithms that can utilize massive datasets and learn complex patterns better than humans can create rules. On the other hand, sceptics assert that opaque models will always avoid the elegance and certainty of classical algorithms. In this section, we will critically evaluate the advantages and disadvantages while keeping in mind the bold provocation that GenAI's authority, while controversial, is on the way.

4.1 Performance versus Optimality

The empirical successes of GenAI solutions cannot be ignored. Neural models have either matched or surpassed the performance of crafted algorithms in critical areas such as computer vision, natural language, and games (e.g., AlphaZero) [7, 8]. For instance, on average, a learned model can sort or search almost as well as an optimal algorithm or a deep-learning-based optimizer, solving NP-hard problems in scheduling and routing that had previously been considered intractable. GenAI cheats the traditional standard using statistical patterns, providing speed without worst-case behaviour. Purists will argue that an algorithm like Quicksort or Dijkstra's can provide worst-case bounds and exact correctness, while great planar neural networks can make some mistakes or have uncharacterized complexity. Still, a slight error rate is tolerated in practice if it works better in the real world, and the epistemic assumptions of complexity assurance are becoming more similar. The logic to support GenAI is practically a maxim: we prefer a learned scheduling system that will offer a solution that we know is only "out-of-sample" 95% correct that can still respond and learn from changing data, than a 100% correct algorithm that failed when the assumptions changed. Hence, the utility of the answer outweighs the formal assurance, and this reasoning is becoming a common decision-making approach in projects and industries.

4.2 Originality and Creativity

A significant advantage of GenAI is its surprising capacity to surprise us with new solutions; AlphaTensor's discovery of faster matrix multiplication is one example [10]. For decades, mathematicians and computer scientists have poured time and energy into similar algorithms—an AI has uncovered new algorithms in just weeks. Generative models like Neural Architecture Search (NAS) have designed neural networks (e.g., EfficientNet, NASNet) that,

on average, outperformed human-designed networks in terms of both efficiency and accuracy, which will be called "automated creativity" in thousands of papers. These examples show how AI could fully explore design space in ways human thought cannot—and do it faster. However, traditional algorithm design is sage and fundamentally bounded by human intuition and experience. GenAI has no such bounds; it can brute-force combos or mutate strategies across hyper-dimensional space and uncover gems a person might never try—this "creative brute force" ushers in a new era of algorithmic discovery as search. Therefore, the human role is to transition from manual designer to high-level curator of machine-generated candidates. Some critics would caution that not everything that glitters is gold: an AI-invented solution may exploit a quirk of a simulator or data and fail to operate in general (the so-called "reward hacking" problem or overfitting in RL). For this reason, rigorous validation will be required to distinguish genuinely superior algorithms from coincidental algorithms. Still, I see the weight of the novel ideas generated by AI affecting the balance. Even if many of them are slogs, the small number of successes can fundamentally change a field—the power-law payoff you cannot get from traditional R&D.

4.3 Black-Box Concern and Interpretability

Perhaps the most pointed criticism of GenAI in algorithm design is the interpretability problem. We traditionally understand algorithms deeply; we can reason about Dijkstra's algorithm and why it finds the shortest paths, and we can reason about Quicksort and how it partitions data. With GenAI, we are often unaware of how or why a solution works. This blurred understanding leads to trust issues. In safety-critical domains (aviation, medical diagnostics), can we trust an AI-derived algorithm that we cannot inspect line-by-line? The lack of transparency raises barriers to continued innovation: if we do not understand or inspect the solution, we cannot build on it or see its flaws. These concerns are legitimate and represent a tangible barrier to the wholesale adoption of GenAI in some domains. In response to these issues, an active research area in explainable AI (XAI) seeks to explain the black-box models. Some methods, such as surrogate explainers and feature attribution explainers (e.g., LIME and SHAP), aim to respond to the famous question "Why Should I Trust You?" regarding model decisions. Researchers seek models that provide intelligible explanations to retain critical elements of trust. Progress is being made. For example, researchers have successfully represented decision tree approximators to help neural planners understand their logic. However, an explanation will never be as deterministic or straightforward as a trace of a deterministic algorithm. In critical analysis, interpretability must be weighed against performance. GenAI almost always wins on performance (doing things that no human algorithm can do) and loses on intelligible structure. The paper argues controversially that, in the long run, performance at scale will trump everything else—the course of technology shows that black boxes (whether steam engines or deep learning) will be adopted if they are performing, and the requirement for interpretability gives way to practicality, except in regulatory environments. We expect even sceptics will gradually settle into an acceptance of a "new normal" in which there is acceptance of many algorithms due to their opaque nature in much the same way we do now (with a certain degree of acceptance), that no human can parse the millions of weights in a neural network. However, we are prepared to accept its output post-validation.

Table 1. Comparison of GenAI-Designed vs. Human-Designed Algorithms.

Aspect	GenAI-Designed Algorithms	Human-Designed Algorithms
Interpretability	Opaque “black-box” models; explanations often require post-hoc tools.	Transparent logic; easy to trace outputs from defined steps.
Optimality	Empirical solutions; lack formal guarantees; validated statistically.	Exact or bounded solutions with theoretical correctness guarantees.
Adaptability	Highly generalizable; can handle multiple/new tasks with data-driven fine-tuning.	Task-specific; adapting to new problems requires redesign and insight.
Scalability	Resource-heavy training; runtime scales poorly, but can solve problems beyond human methods.	Designed for efficiency; scales by known complexity, but struggles with some tasks.
Development	Data-driven, iterative, uncertain; less coding, more training and tuning.	Manual design and coding; deterministic outcomes; slow for complex problems.

4.4 Human Oversight and Theory

GenAI’s rise does not imply the death of human algorithmic expertise – it simply redefines it. Instead of writing the algorithm, we would characterize it as algorithmic engineering, experts are now designing for designers: the training regimes, the objective function, and the constraints whereby the GenAI delivers its solutions. For example, in Neural Architecture Search (NAS), researchers implement search space and reward criteria so that the architectures discovered by AI meet practical concerns [4]. In reinforcement learning, used for discovering algorithms, humans make the decisions on problem encoding and goals (i.e., maximize speed or accuracy). Thus, theoretical computer science is not dead; it has morphed into AI. Theoretical knowledge is helping constrain the GenAI into valid solution spaces (to prevent a scenario, for example, where a neural network might generate an invalid graph) [2]. An interesting collaborative arrangement is underway: human intuition gives high-level structure, and AI offers low-level detail or novel variations. Critical analysis must acknowledge that the best results emerge from the marriage of human theory and AI search; pure GenAI discovery, unencumbered, can fumble its way into or rediscover known results; or, when guided by human intention and thought, adjusting, the pure potential of GenAI will have excellent upward traction. We see evidence of this interplay in hybrid methods like the neuro-symbolic approaches, in which learned components are coupled with symbolic algorithms. It seems reasonable to suspect that as algorithm design becomes more of a two-step process, (1) theorists would sketch high-level approaches/methods, and (2) GenAI would expand and implement beyond human ability alone.

Another aspect is the verification problem. As GenAI designs more and more of our algorithms, how will we verify the correctness and robustness of the algorithms? Formal verification is an emerging field of neural networks. We may see new verification tools that can accommodate learned components, or the field may become more interested in probabilistic guarantees (for example, probably approximately correct – PAC – analyses for learning algorithms) [12]. The hopeful point of view, correctly or incorrectly, is that GenAI may help at this point, too: generative models may be able to contribute to the generation of proofs or counterexamples or essentially contribute to their verification. Indeed, transformers have been used to generate mathematical proofs and check code correctness to some extent (OpenAI’s Codex can flag bugs by generating unit tests) [4]. The emergence of GenAI will not remove the requirement of rigour; it will just change how we find rigour, likely even with the help of AI.

4.5 Survey of Opinions

To put this controversy in perspective, it is interesting to note how the community seems divided. On one side are enthusiastic voices like Google’s Jeff Dean, who predicts that “neural networks will transform 1000x more industries” in the not-too-distant future, with the implication being that learned solutions are poised to be applied to algorithmic tasks everywhere. On the other are veteran researchers like Donald Knuth (in a hypothetical sense), wistfully recognizing that nothing can compare to the beauty of a hand-crafted algorithm compared to some giant inscrutable matrix of numbers. Sometimes, the debate may get framed as brute force data-driven vs insightful elegance. Our review acknowledges the former as the future but does not wish to forget the latter’s legacy. Indeed, the position taken in this controversy deliberately stirs the pot: declaring GenAI will reign supreme is a strong statement, but numbers and trends are on its side (e.g., GenAI publications and successful examples have swamped traditional works) [5]. By voicing this extreme position of opinion, we invite the researchers and scholars to protect their turf by disproving this view or, if they agree, to advance the field. Either way, we are likely to make more progress.

In conclusion, this critical analysis shows no absolute or unconditional reign of GenAI in the algorithm design realm, with its trade-offs—interpretability, guarantees, and development expense—that are non-trivial. In terms of the resurgence, adoption, and undeniable benefits of GenAI-driven design phenomena, it is large enough momentum and reality that it is hard to recognize it completely reverting to prior stages of the research environment. Yes, the controversy about GenAI will remain, but so will the velocity of advances. In the next section, we will examine the ethical implications of relinquishing algorithmic control to generative systems, which is an issue that we must address, given the urgency of this paradigm shift.

5. ETHICAL IMPLICATIONS

The emergence of GenAI for algorithm design raises heavy ethical and social questions. Suppose algorithms, the decision-making logic within the software, are increasingly being developed by AI. How do we compare these algorithms with human values and ensure that these algorithms do not cause damage or unintended (negative) consequences? The main ethics of AI-designed algorithms is a part of more extensive set of ethics and issues, which includes bias and fairness, accountability, and transparency.

5.1 Bias and Fairness

AI models can perpetuate or exacerbate biases in training data. When such models are designing algorithms or decision rules, we are at risk of building systems at scale that instantiate biased logic patterns. For example, a GenAI that learns to distribute resources or prioritize candidates could adopt discrimination (like penalizing individuals due to specific characteristics) found in historical data. The algorithm it produces could be more biased than any human-coded algorithm simply because it optimizes to an objective that does not include fairness [13]. Buolamwini & Gebru (2018), a highly cited paper, examined algorithmic bias in facial recognition and showed how easily AI can manifest biased behaviour. The ethical obligation is straightforward: train a GenAI on balanced data and include fairness constraints in the design. Tools are now appearing (like adversarial debiasing and fairness-aware training) to reduce bias's impact in generative models. Nonetheless, it is still a significant task—as GenAI designs, the audits must keep pace. Actionably, the good news is that there is a priority on AI fairness, with the community amassing papers (several papers overcoming the 1000 citations) on mitigating bias.

5.2 Responsibility and Accountability

Who is at fault if a GenAI-designed algorithm harms people or something fails unexpectedly? With an algorithm designed by a human, we can point to a flaw in the design or a mistake in the code made by a person or team of people. The decisions made by GenAI arise after something emerges from a complex training process that is not directly interpretable, which complicates accountability. If the path planning algorithm for a self-driving car (designed by a neural network) fails with life-threatening consequences, who is ultimately responsible: the AI itself, the engineers that trained it and selected the data used, or the company that developed it to be used on our streets? The issue of accountability is complex, and some scholars have referred to the “opacity of AI systems” as a prominent ethical issue. Whittaker et al. (2018) also famously mentioned “significant concerns about the lack of due process, accountability, and auditing” of AI-based decision-making processes [14]. The first remedial step relates to transparency: for example, by documenting how the GenAI system was trained (which datasets were used, the training parameters on how the AI learns), and what data testing activities it underwent. The models of “model cards” and “datasheets for datasets” (Gebru et al., 2018), which have garnered over a thousand citations each, exemplify an approach to transparency [13]. The second remedial step is regulatory: many frameworks incorporate considerable transparency requirements for particularly high-risk AI systems, like the proposed EU's AI Act, that impose explainability and audit requirements on AI systems. For a GenAI-designed algorithm, this may involve keeping track of some form of logging on the AI design process and incorporating explanatory proxies for the logic it produces. The ethical bottom line is this: At no point should the emergence of GenAI lead to the loss of human responsibility of humans. Human oversight and legal accountability structures need to evolve to answer the question: If an AI-designed algorithm causes X, who fixes it, and how do we prevent it from happening again?

5.3 Safety and Control

A more futuristic ethical concern is the possibility of losing control. As GenAI systems get better and better at designing algorithms that might even transcend human understanding, the long-term mission of AI algorithms optimizing for unaligned objectives

becomes highly concerning. Bostrom and Yudkowsky (2014) discuss the “22value alignment problem”, which aims to ensure that AI's goals are aligned with human values [15]. The risk of an AI that can design its algorithms and is thus capable of engaging in self-improving algorithms in ways we can not easily predict when it self-improves (as in the classic paperclip maximizer thought experiment in which the AI finds a way to continuously self-improve an algorithm in the way it makes paper clip, at the expense of other values). While this is not an immediate concern with our current narrow and human-constrained GenAI capabilities, this indicates the increasing potential for uncontrolled self-improvement. The ethics-of-AI-safety community has articulated in highly cited literature (for example, Amodei et al. 2016 “Concrete Problems in AI Safety”) the technical challenges of restraining AI misbehaviour [12]. An example of a safety problem is an AI-designed algorithm that is too good at achieving a proxy objective. Therefore, the AI finds a loophole to achieve a sufficiently good assessment of the target action (reward hacking). Currently, there are active efforts to ensure that generatively designed algorithms cannot exploit loopholes. Techniques like adversarial training and constraint-based learning (to penalize unsafe strategies) are advancing. Both the OpenAI and DeepMind teams have published AI safety techniques that have received over a thousand citations in the literature, which supports the notion that the community recognizes these concerns. Practically, companies deploying GenAI-generated code or logic are instituting AI audit committees and red-team testing from external parties to catch any unsafe behaviours that may be hidden until release.

5.4 Impact on Jobs and Human Expertise

We should examine societal concerns ethically. If GenAI continues to dominate instruction sets for algorithms, what happens to all the human algorithm engineers? There is a workforce concern; some routine, algorithmic activity could be fully automated. Even though this will free developers to pursue their higher-order creative work, it could take out certain skill sets entirely. It also presents us, as educators, with an emphatic shift – future computer scientists will probably need as much training in “model steering” and data curation as traditional algorithms. Ethically, various stakeholders must ensure the next generation is prepared to work with GenAI, not supplant it. Additionally, diversity appears even more important; if only a small group of people know how to instruct (steer) a GenAI design, the resulting algorithmic system might only reflect a small set of values. Inclusiveness in AI design teams can serve as an ethical hedge, specifically concerning providing diverse stakeholders' views into whatever objectives GenAI will optimize.

5.5 Regulation and governance

Finally, governance of AI-design algorithms is an emergent space. We could even see regulatory regimes demanding that any AI-designed algorithm in finance, healthcare, criminal justice, etc., have an independent review/certification. IEEE and ISO have begun to develop standards for AI system transparency and risk management. Highly cited position papers on the policy development side (Russell et al. 2015 Educational Initiatives for Coding Schools promoting Beneficial AI) describe the cooperative relationship needed to work between AI designers, governments, and ethicists in this space [2]. The controversial claim of this paper is that GenAI will dominate, which should heighten the emphasis on building ethical guardrails in the present. If society waits until AI-designed algorithms are ubiquitous, it might be too late to embed ethical principles retroactively. We already have well-articulated

proactive concepts for action in this area, such as ethics-by-design principles, which will help ensure that as GenAI algorithms proliferate, they do so in a manner congruent with societal values and legal requirements.

To summarize, ethical considerations around the GenAI-generating software can be encapsulated in bias, accountability, safety, competent human involvement, etc. While offering promising large-scale social payoff, GenAI adds complications to maintaining systemic control and trust. We should feel encouraged that we have begun to recognize these issues, and it is an active research area with vibrant policy discussions, including several towering ethics papers addressing AI, some exceeding 1000 citations, demonstrating a large community striving to keep AI as beneficial. Engagement with ethical impact is not merely a secondary mitigation strategy to a GenAI trajectory but a cornerstone of sustainable success.

6. CONCLUSION

Generative AI will probably be the future of algorithm design while representing an inflexion point around computing that is more likely divisive but evidence-based [6]. Naturally, we have seen that large language AI models can now create algorithms, whether for playing Go, writing code, or finding shortcuts to new maths beyond their human-designed predecessors, as shown by the studies [8, 7, 9, 10]. The publication and citation trends reinforced this trend: GenAI research is accelerating lightning, while classical algorithmic research is growing modestly [6]. This review has taken a blaze away to say we are shifting paradigms. Will we soon be at the point when most algorithms will be generated by AIs, and human engineers will act as mentors or quality assurance audit agents for the AI? The evidence would tend to say yes, but good reason retains room for healthy doubt [5].

From a technical perspective, GenAI is a dominant force when uncontrolled complexity, unnatural flexibility, and the creative resources of modern AI are applied [2, 4]. Generative AIs can model knowledge from millions of examples, utilize a design space at an unreal scale and learn like a madly uncontrollable toddler; no human or traditional design will ever match these capabilities. The generated solutions we observe in GenAI-designed algorithms demonstrate significantly comparable performance to human algorithms across multiple benchmarks and from the frame of reference. With expanding model sizes and training data (GPT-4, now also rumoured with higher capacity and even more reasoning and coding skills), we are likely to see AI become not just an enabler of development but a developer in its own right: for example, AI may routinely be able to generate good-enough algorithms for new hardware or new problem areas, relying on little human input, and effectively become a meta-algorithm developer.

However, our analysis made clear that there are also serious challenges. GenAI does not eliminate the value of traditional algorithms overnight. There will continue to be contexts—especially those that require high assurance, interpretation and/or low resource—where classical algorithms will have advantages or cooperate with learned components. The future will likely include hybrid systems, where the best of both worlds can be exploited. We also highlighted the importance of ethical oversight as the world transitions to new practices [15, 12]. Any technical dominance has to be paired with responsible application. Society will not tolerate algorithms (as clever as they might be) that are outrageously unfair or tragically unsafe. The growth of algorithm generation by GenAI is occurring, and at the same time, the ethical and safety research in AI is acceler-

ating [12]. Both are evolving together: We are not moving forward unthinkingly; we are trying to harness and leverage this powerful technology for good.

In short, the emergence of generative AI in algorithm development is accurate and advancing rapidly. This review might have been controversial in asserting that GenAI will advance beyond human algorithm developers, but growing data and evidence support this assertion. We encourage researchers to build on the work reported here: exploring new GenAI-design solutions, further working on the open problems we identified (for example, interpretability and verification), and rigorously evaluating our thesis. Will the following great algorithm be harnessed in a human brain or hallucinated by a hyper-scale transformer? The safe bet is that it will be both, but the flow balance is moving toward the latter. However, we hope this work has been a wake-up call: the era of AI-designed algorithms is here, and if we accept this reality, we can shape it. Wouldn't it make us happiest to see it be rebutted by human-designed breakthroughs instead of the rapidly evolving generative output? We invite those who disagree to do so, but until then, the torch is being passed on to our generative offspring, and they are running with it at a tempo that keeps us all scrambling to keep up.

7. REFERENCES

- [1] I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio, "Generative adversarial networks," in *Advances in Neural Information Processing Systems*, vol. 27, 2014.
- [2] Y. LeCun, Y. Bengio, and G. Hinton, "Deep learning," *Nature*, vol. 521, pp. 436–444, May 2015.
- [3] J. Li *et al.*, "Alphacode: Using large language models for competitive programming," DeepMind, Tech. Rep., 2022.
- [4] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, A. Kaiser, and I. Polosukhin, "Attention is all you need," in *Proceedings of the Neural Information Processing Systems Conference*, 2017, pp. 6000–6010.
- [5] S. HAI, "Ai index report 2023," Human-Centered AI (HAI) at Stanford, Stanford, CA, Tech. Rep., 2023.
- [6] Stanford HAI, "Ai index report 2024," Human-Centered AI (HAI) at Stanford, Stanford, CA, Tech. Rep., 2024.
- [7] D. Silver, J. Schrittwieser, K. Simonyan, I. Antonoglou, A. Huang, A. Guez, T. Hubert, L. Baker, M. Lai, A. Bolton *et al.*, "Mastering chess and shogi by self-play with a general reinforcement learning algorithm," *Science*, vol. 362, no. 6419, pp. 1140–1144, Dec. 2018.
- [8] S. Hassabis, D. Silver *et al.*, "Mastering the game of go without human knowledge," *Nature*, vol. 550, pp. 354–359, Oct. 2017.
- [9] T. B. Brown, B. Mann, N. Ryder, M. Subbiah, J. Kaplan, P. Dhariwal, A. Neelakantan, P. Shyam, G. Sastry, A. Askell *et al.*, "Language models are few-shot learners," in *Advances in Neural Information Processing Systems*, vol. 33, 2020, pp. 1877–1901.
- [10] L. Fawzi *et al.*, "Alphatensor: A general reinforcement learning agent for discovering matrix multiplication algorithms," *Science*, vol. 375, no. 6581, pp. 1224–1230, Mar. 2022.
- [11] A. Karagöz, "Ethics and technical aspects of generative ai models in digital content creation," arXiv preprint arXiv:2412.16389 [cs.CY], Dec. 2024.

- [12] D. Amodei, C. Olah, J. Steinhardt, P. Christiano, J. Schulman, and D. Mané, “Concrete problems in ai safety,” arXiv preprint arXiv:1606.06565 [cs.AI], Jun. 2016.
- [13] J. Buolamwini and T. Gebru, “Gender shades: Intersectional accuracy disparities in commercial gender classification,” in *Proceedings of the Machine Learning Fairness Workshop*, vol. 81, 2018, pp. 77–91.
- [14] J. Whittaker, K. Crawford, E. Dobbe, H. Fried, S. Kaziunas, P. Mathur, M. West, and T. Richardson, “Ai now report 2018: Interpreting and auditing ai systems,” AI Now Institute, New York, NY, Tech. Rep., 2018.
- [15] N. Bostrom and E. Yudkowsky, “The ethics of artificial intelligence,” in *Cambridge Handbook of Artificial Intelligence*, K. Frankish and W. Ramsey, Eds. Cambridge, U.K.: Cambridge University Press, 2014, pp. 316–334.