

The Rise of Visual AI: Opportunities, Misuses, and the Need for Ethical Oversight

Zakiya Sayyed¹, Mehak Fatima², Hasan Phudinawala³

^{1,2}P.G Students, Department of Data Science

³Coordinator, Department of Data Science

Royal College of Arts, Science and Commerce (Autonomous), Mira Road (East)

Abstract: *Visual Artificial Intelligence (Visual AI) systems apt at seeing, understanding, and creating visual content have quickly evolved a variety of industries, both in medical utilize diagnostics and creative media. Nevertheless, these very technologies pose high moral and social questions, e.g., in the form of deepfake fake news, automatic discrimination, and mass surveillance. The present paper is a systematic review and made up of recent articles and case studies (2023-2025) in which the opportunities, misuses, and gaps in governance of Visual AI are discussed. By being compared to one another, it determines the drawbacks of the existing detection mechanisms, the issues of the generalization, and the inadequacy of legal and ethical controls. An example of failed administration on a real-life situation can be seen in a case study of the implementation of facial-recognition in the city of New Orleans. The discussion summary is that innovation and technical progress cannot be ensured to be beneficial to the population in terms of ethical excellence and transparency, legally overseen, and not irresponsible to anyone.*

Keywords: Visual Artificial intelligence, Deepfakes, Ethical AI, AI Governance, Surveillance, Algorithmic Bias, Responsible Innovation

I. INTRODUCTION

Visual Artificial Intelligence (Visual AI), the term denotes one of the most impactful subfields of Artificial Intelligence (AI) that changes the technologies of machine perception and image generation within the last decade. Visual AI is dominating almost all spheres of human activity from autonomous vehicles and the medical imaging industry to surveillance work and the creation of digital media. With the introduction of deep learning models, both Convolutional Neural Networks (CNNs) and Generative Adversarial Networks (GANs), the world has undergone a new swift visual AI technologies growth as machines are now capable of synthesizing, manipulating, and analyzing visual data with new levels of perception and realism like never before [1]. Croitoru et al. state that due to this development, applications of this nature have grown exponentially, including face recognition, video-manipulation, object-detecting, and creative-media-generation [2]. Nevertheless, with more functional systems, risks of abuse, breach of ethics and damage to society have become even greater [3]. The digital creativity revolution posed by the democratization of generative AI tools, including deepfake generators and text-to-image diffusion models has ruined the belief in visual evidence at the same time. According to the findings presented by Khan et al., the detection systems, but with their high level of technology, are susceptible to adversarial attacks and manipulations [4]. Antonitou, Bale, and Ugwu (2024) were able to prove that biases of the datasets are oftentimes recreated by AI-based detectors to mirror the array of gender and racial inequality [5]. The two- pronged aspect of Visual AI, which offers both innovative and inception as well as deceptionary possibilities such as advancement, underscores the necessity of a robust ethical regulation, transparency as well as accountability models. International cases of misinformation and identity theft caused by the use of deepfakes and unlawful surveillance have served to make these discussions urgent, with examples like Project NOLA (2025) reporting the use of facial recognition technology despite it being outlawed by law, and it was disproportionately used against marginalized communities [6]. In addition to deepfakes, the role of Visual AI in law enforcement, healthcare and governance raises some serious ethical and policy challenges. Intel use Automated surveillance and facial recognition is an area of systems engineering where further efficiency with minimal regulation is assured, leading to invasion of privacy and false



identifications [7]. Visual AI can be useful in the field of medicine in early disease detection and accurate diagnostics, although the aspect of algorithmic unpredictability concerns clinical accountability and the lack of objectivity in patient outcomes [8]. Moreover, Liu, Tao, and Zhou (2024) presume that although multimodal systems- the ones combining both visual, audio and textual information- can better capture the accuracy of the model, they also impact computational requirements and energy expenses with more concern to sustainability [9]. Therefore, Visual AI is a techno-economically symbiotic choice, as it is both a moral and a social intersectionology where all the responsibility and the innovation must be balanced.

Researchers continue to point out the moral paradox of AI governance: AI systems can be hard to control because the very qualities that have made them successful such as freedom, scale, and malleability are the same factors that make AI systems challenging to control [10]. Gong and Li (2024) indicate that AI-driven tools of detection lose their reliability significantly in the conditions of the real life, which makes their applicability problematic in both legal and journalistic pieces of evidence [11]. The evaluation metrics used today will not compare among studies as it is characterized by Singh Chauhan et al. (2024) and Vyas et al. (2024), and thus, coming up with exaggerated claims about accuracy and damaging the concept of transparency [12][13]. Such inconsistency is not only weakening trust levels of people towards the AI systems but also revealing the insufficiency of the existing regulations systems to match the technological innovation. In addition, the visual AI squares the difference between human and artificial perception. Its mimic power of reality has become an object of argumentation on the question of epistemic trust lien-- the facade to differentiate truth and lie in a visualistic photobombing virtual world. This epistemic insecurity is an immediate menace to journalism, elections, and social unity in the democratic society. Even with an explicit warning about deepfaked media, the performance of human subjects is comparatively poorer than chance in spotting attempt manipulation [14] even though the results of the study showed that the difference between human control and non-conclusive media presence hovers around, according to the findings of Computers in Human Behavior Reports (2024) [14]. This result highlights the aspect of the limitations of human thoughts in fighting visual misinformation and indicates the need to involve collaborative frameworks by using both human judgment and automated recognition. Visual AI is thus less about technological sophistication than with regard to institutional responsibility, public policy and interdisciplinary governance. Regulation transparency requirements and audits of biases Studies like the EU Artificial Intelligence Act (2024) and the US Blueprint of an AI Bill of Rights (2023) are key milestones on the way to codification of transparency requirements and bias audits. Nevertheless, all these frameworks are mostly reactive in nature yet they resolve the harms after they occur unlike in their proactive approach [15]. Such a gap between the speed of AI innovation and the slowness of the regulatory mechanism, poses essentially a regulatory lag, and establishes an environment in which unethical activities, including overreach by surveillance systems and fake news via synthetic media, can proliferate without much payment [16].

It is against this background that the current study will attempt to summarize available literature on the emerging popularity, opportunities, and ethical issues around Visual AI. In particular, the following are the objectives of this research: 1. To analyze the technological development of Visual AI and determine its most significant levels of implementation in the domains. 2. To critically assess the available literature on the abuse of Visual AI, specifically, deepfakes, surveillance bias, and privacy issues. 3. To examine the ethical and governing philosophy presented so far in mitigating the dangers of Visual AI. 4. To establish the gaps in the research and offer the perspectives on the further interdisciplinary interaction between the technologists, ethicists, and policymakers.

II. LITERATURE REVIEW

Visual Artificial Intelligence (Visual AI) is currently one of the fastest growing, as well as the most rapidly discussed spheres of artificial intelligence research. It has received a wide variety of different applications in automated image recognition and content generation, surveillance, and deepfake creation, transforming both the technological and ethical landscape. There has been a two-pronged trend in the literature when, on the one hand, Visual AI is seen as a source of innovation and economic benefit and, on the other hand, individuals are exposed to false information, prejudice, and a worrisome amount of privacy to compromises. In a bid to address these dimensions in detail, this section summarizes twelve peer-reviewed works (2023-2025) that represent in totality the development, opportunities, and ethical consequences of the Visual AI. Table 1 summarizes the focus of each study, their findings, ethical implications as well



as any gaps that were identified in the research. The choice of works concerning deepfake detection, fairness, transparency and governance models is aimed to demonstrate how appropriate AI implementation can be complicated.

2.1 Technical Foundations and Area of use.

The engineering backgrounds of Visual AI and deepfake studies have greatly advanced over the previous couple of years with various studies providing various thoughts on generation, identification, and endurance. Deepfake Media Generation and Detection in the Generative AI Era [1] offers a high level of taxonomy regarding the way in which a deepfake works, such as GAN-based, autoencoder, and diffusion model, and analyses the boundaries to generalization of the current detection tools. This paper points to the reality that most detectors are capable of succeeding on controlled data, but cannot be generalized to unseen manipulations, or noise introduced by the environment.

Continuing on this challenge, Unmasking Synthetic Realities in Generative AI [2] discusses the forcefulness against adversarial attacks and how a significant detector decrease can happen at a number of few perturbations or artifacts of compression. The authors emphasize that not only a technical issue but also an ethical necessity, adversarial robustness is the determinant of false information, which may slip with untrustworthy systems. Adding to these findings, A Survey on Multimedia-Enabled Deepfake Detection [3] highlights the advantages of multimodal fusion - the combination of visual, audio, and text information in the study - to be reliable in the actual testing setting.

Continuing on this Evolving from Single-Mode of Multi-Model Facial Deepfake Detection [4], which analyzes the hybrid architectures of convolutional and transformer-based systems to improve the capacity to extract spatial-temporal features. Although these enhancements in performance are observed, the study is much wary of multimodal approaches which require huge datasets which are mostly not available and are balanced. Lastly, empirical results on Deepfake-Eval-2024 Benchmark [5] indicate that there is a significant decline in detection performance in uncontrolled and real-world scenarios, which proves the existence of a divide between lab and field relevance. Collectively, these studies [1]-[5], demonstrate the technical challenge of Visual AI studies and the immediate need to develop accessible and acceptable models of the core in all three aspects: accuracy, generalization, and interpretability.

2.2 Misuse, Bias, and Case Studies

Visual AI technologies become more and more abused to serve deceitful and unethical intentions, despite the opportunities to be creative and innovative in the message. According to Human Performance in Detecting Deepfakes [6], even trained humans show only slight improvement compared to the possibility of pure chance detection of synthetic media, which proves that technological aids and media literacy trainings are indeed necessary. Equally, Deepfake Detection and Classification of Images from Video responded with the studies that see a significant issue in the lack of diversity in datasets, and underrepresented population groups cannot be mitigated by the models because of the problem of fairness and bias [7]. The prominent real-life example of abuse and lack of control relates to New Orleans Facial Recognition Surveillance Case [8], where surveillance networks created by privately funded companies used AI-supported facial recognition without any obvious control guidelines. The case brought in the issue of racial discrimination, risk of false identifications and gaps in accountabilities in the nature of the systems of public and private surveillance. The same case of governance has been experienced in European union and other parts of the world where widespread public bias in terms of biometric surveillance tend to edge out the law and ethics. These cases help identify how the abuse of the Communist AI is not limited to the realm of deepfake development but permeates into the governance system with urgent ethical issues.

2.3 Mitigation, Governance and Oversight.

The process of overcoming it is unified by combining technical innovations with regulation and ethical regulation on the major risks of Visual AI. Scientific studies in the areas of adversarial resilience, digital watermarking, and explainable artificial intelligence have improved the technological protection against abuse. However, as demonstrated in [1]-[5], largely technical fixes do not give such sufficient assurances of responsible deployment. A significant move, albeit not the final one, towards accountability, transparency, and human rights of AI systems can be seen in the introduction of global governance paradigms, including the AI Ethics Guidelines (2023) of UNESCO or the EU AI Act (2024) [12]. These



frameworks promote risk based categorization, computer program auditability, and obligatory records to enhance the belief of the populace and guard against the misfortune. However, implementation requirements are complicated by the differences in jurisdiction and the phenomenal pace of the development of generative technologies. Visual AI should be placed under sustainable management thus, an integrated model is necessary incorporating technical validation, legal regulations and participatory ethics- whereby the innovation is prone to the societal values and democracy.

recent research is still falling in line with the growing complexity and social effect of Visual AI systems. In addition to technical detection and generation, dispensers started focusing on the aspects of explainability, shape surface of models and sustainability in AI-based media synthesis. Other publications have delved into the principles of combining ethical auditing models to determine the fairness of algorithms, especially with facial manipulation recognition of which demographic bias is still a risk factor. Multimodal systems are also emerging that are being optimized to handle contextual and behavioral signals, creating enhancements of contextual verification of authenticity through social media platforms. Even under such developments, the literature continually reports the existence of a gap between the laboratory testing and the practical application, clarifying the significance of strong datasets, the cross-domain learning and consistency in benchmarking. Taken together, this literature presents the view that efficient governance of Visual AI should be linked with technical expertise and ethical responsibility.

The development of Artificial Intelligence (AI) and Visual AI technologies has progressed quite fast and has made a significant change in the areas of healthcare, surveillance, and decision-making systems. Recent studies emphasize the ability of AI to analyze visual data with very high precision and hence help in medical diagnostics, security control, and governance through data. Nevertheless, in time with these developments, researchers have expressed worries over the problem of algorithmic bias, privacy encroachment, the absence of transparency, and the ethical nature of automated decision-making. It has been demonstrated in the literature that on the one hand, Visual AI can improve the efficiency and capacities of humans, but on the other hand, it also presents the risk of discrimination, misuse of data, and loss of human responsibility. In order to become responsible in innovation, researchers highlight that it is important to have a fair representation of data, explainable AI models, and robust ethical governance structures. The table below provides an overview of the major researches that analyze the opportunities, difficulties, and ethical issues of AI and Visual AI use within different industries.

Table 1: Literature Review on Visual AI -Opportunities, Misuses, and Ethical Oversight

Author(s) & Year	Paper Title	Purpose / Focus	Significant Results of the Paper	Ethical / Sociological Detention
Nguyen & Tran (2024) [1]	<i>Deepfake Media Generation and Detection in the Generative AI Era</i>	Keep up with the available variants of deepfakes generation and assess the potential issue of detection.	Speculations on creation and detection algorithms of the deepfakes; explains the deficiency of threat detectors in the area of generalization.	Recommendations of public datasets and conscientious use; comments that the applicability in unobserved circumstances will be ineffective.
Khan & Li (2025) [2]	<i>Unmasking Synthetic Realities in Generative AI</i>	Resistance and defense systems	It specializes in resistance and defense mechanisms on opponent detection models.	Soundeths danger in the system of adversary concerning design; is interested in acting rightly when forming systems.
Zhou & Patel (2024) [3]	<i>A Survey on Multimedia-Enabled Deepfake Detection</i>	To address the models based on the multi-modal combination (audio, visual,	Positively reinforces the fusion of imagery, aural and text fusion that are more multiprequent to	There is cross-modal ethical authentication and greater variance in larger dataset in the protagonists.



		text) in the detecting system.	enhance the accuracy of detection.	
Croitoru et al. (2024) [4]	<i>A Comprehensive Survey on Deepfake Generation and Detection Techniques</i>	To be able to automatically determine the level of performance.	Rather, discussions on GANs research and diffusion-based detectors.	Increases openness and consistent cuteness to become fair.
Liu, Tao & Zhou (2024) [5]	<i>Hybrid Multimodal Systems for Visual AI Detection</i>	Liu, and both offered Hybrid Multimodal systems which considers text, audio and vision to detect.	Performance trade-offs are reviewed by looking at vision-audio-text hybridist.	Raises concerns regarding the energy utilization and environment sustainability.
Bansal et al. (2024) [6]	<i>DeepFake-Eval-2024: Benchmarking Deepfake Detectors in the Wild</i>	To understand deepfake detectors in the wild.	Practically, its accuracy has been found to be low when compared to that of the laboratory.	Whalid warns on the dependence on standardization overdependence; reports that there is a discrepancy between lab-real-world performance.
Bale, Ochei & Ugwu (2024) [7]	<i>Comparative Study of Classical vs. Deep Learning-based Deepfake Detection</i>	To compare the traditional and modern approaches.	Among classical models, shows with deep learning have shown superior findings but according to deeply deep research they have been found to be biased in terms of demographics.	Highlights demand that the representative information should reduce systemic bias.
Chauhan et al. (2024) [8]	<i>Fairness, Robustness, and Explainability in Deepfake Detection</i>	To search ethical matters such as fairness in the detectors of AI capabilities.	Presents deception of non-fairness and non-explainability of detecting algorithms.	Demands to examination of AI which is morally influential.
Gong & Li (2024) [9]	<i>Reliability of Deepfake Detection under Adverse Conditions</i>	To measure the operation of the machine detecting models that are obtained under the conditions of contamination and stress.	Finalizes compression, entails reduction in the performance of detection through noise.	Laptuses are opposed to these forms of blind use of automated verification.
Vyas et al. (2024) [10]	<i>Evaluation Metrics of Deepfake Detection: A Meta-Review</i>	In a bid to assess consistency and disclosure in regards to	Reveals discrepancies in the stated findings and activities.	Proposes uniform performance appraisal processes.



		reporting metric of detection.		
Chen, Yang & Yuan (2025) [11]	<i>Transformer-Based Detection of Next Generation Deepfakes</i>	To prove that transformers easily and reliably out performs CNNs in identifying next-gen deepfakes.	Transformer models are also more accurate detectors as compared to CNNs.	Suggests explicable nascent AI towards human wellbeing.
European Commission (2024) [12]	<i>AI Act 2024: Towards Responsible and Transparent AI Governance</i>	To recommend AI governance through responsible and transparent AI classification and regulation through risk based classification and regulation.	Proposes a classification of AI and compliance mechanisms, depending on the risk.	Establishes a code of any ethical responsibility and fairness in the use of AI.

IV. ANALYSIS AND CONCLUSIONS

The development of the Visual Artificial Intelligence (Visual AI) has triggered the paradigm shift in various fields, including the Digital content creation and entertainment to security, forensics, and journalism. As existing literature confirms, Visual AI in its innovative approach to creative and analytic solutions to problems on the one hand brings significant changes, on the other hand, it implies immense ethical, societal, and governance problems. The reviewed papers [1]-[12] experience were analyzed and found several pattern patterns such as acceleration of the speed of algorithms, constrained within the validity in the real world application, the increasing use of technology misuse fully, and the lack of system structures to handle such applications. Technically, the advances made in the recommendation of generative adversarial networks (GANs), autoencoders, and diffusion-based models have drastically advanced the realness of created media [1]. The advances have made it possible to come up with synthetic images, video and voice tapes that cannot be distinguished by the real ones. On the one hand, this development has triggered innovations in the spheres of entertainment, marketing, and design, yet on the other hand, due to the advancement, it has become even more challenging to spot any manipulated content. In a set of different researches [2][4][9], deepfake detectors learn with the controlled datasets also have the inability to experience the generalization at the conditions of facing a new manipulative tool and compression stages.

This technical issue exposes a key paradox of Visual AI that the same power of creating gives to creativity allows for deception and misinformation. It can also be hinted out by the review that combining multimodal detection systems, such as visual, auditory, and textual detection methods, can promise us better levels of reliability [3][5]. These systems replicate the process of the human brain by comparing several signals of different senses in the identification of inconsistencies. Nevertheless, these hybrid architectures require very high levels of computational capabilities and data sets that are big but not necessarily homogeneous. The sustainability and environmental questions of multimodal learning tend to be ignored in the context of the AI ethics debate as Liu, Tao, and Zhou [5] point out that developing multimodal learning is quite energy-consuming. On this way the scalability of such detection frameworks should be balanced with their ecological footprint and availability in low resource areas. The theme that stands out critically in the literature is the lack of bias and inequality in Visual AI technology. Research [6] and [7] has shown that in many cases, deepfake detectors do especially poorly in the cases of groups not represented well by the training process because their datasets are biased. The ethical implications of this are devastating because, in case of prejudiced detection system, there is a possibility of



exposing a marginalized group of people to danger, wrong information, or false identifications in irregular manner. In addition, accountable and explicability, key components of trustful AI are not properly integrated into the majority of existing detection proposals [8].

Visual AI is an emerging technology that has gone beyond extraordinary and poses challenges like never before. On the one hand, generative modelling, vision-language integration and transformer architecture innovations have increased human creativity and abilities. On the other hand, the same technologies allow spreading of informed, abuse of surveillance, and undermining digital trust. The reviewed literature highlights that, neither technological sophistication can guarantee ethical use, but with proper governance, transparency, and fairness, the technology can be used. A balance here can be attained through coming up with explainable, bias-sensitive detection systems and putting into practice globally consistent AI ethics regulations in future research. An ecological approach between developers, regulators, and end-users is imperative to agree on the responsibility in AI-inspired societies. Ethical regulation of Visual AI, then, is not just a measure of control in form of regulation, it is also a value and social requirement of preserving digital integrity in the era of smart images.

REFERENCES

- [1] Nguyen, T., & Tran, P. (2024). *Deepfake Media Generation and Detection in the Generative AI Era*. IEEE Access.
- [2] Khan, R., & Li, J. (2025). *Unmasking Synthetic Realities in Generative AI*. Journal of Machine Vision and Applications, 38(2), 45–59.
- [3] Zhou, S., & Patel, M. (2024). *A Survey on Multimedia-Enabled Deepfake Detection*. ACM Computing Surveys, 56(1), 1–25.
- [4] Croitoru, I., et al. (2024). *A Comprehensive Survey on Deepfake Generation and Detection Techniques*. Pattern Recognition Letters, 178, 108–127.
- [5] Liu, Y., Tao, Z., & Zhou, Q. (2024). *Hybrid Multimodal Systems for Visual AI Detection*. IEEE Transactions on Multimedia, 26(3), 2301–2315.
- [6] Bansal, R., Ochei, K., & Ugwu, C. (2024). *DeepFake-Eval-2024: Benchmarking Deepfake Detectors in the Wild*. arXiv preprint arXiv:2403.08976.
- [7] Bale, M., Ochei, K., & Ugwu, C. (2024). *Comparative Study of Classical vs. Deep Learning-based Deepfake Detection*. Journal of Information Security and Applications, 79, 103622.
- [8] Chauhan, S., Singh, R., & Li, X. (2024). *Fairness, Robustness, and Explainability in Deepfake Detection*. Artificial Intelligence Review, 57(6), 3491–3510.
- [9] Gong, W., & Li, H. (2024). *Reliability of Deepfake Detection under Adverse Conditions*. Multimedia Tools and Applications, 83(12), 28765–28789.
- [10] Vyas, P., et al. (2024). *Evaluation Metrics of Deepfake Detection: A Meta-Review*. International Journal of Computer Vision, 132(7), 1590–1611.
- [11] Chen, J., Yang, Q., & Yuan, S. (2025). *Transformer-Based Detection of Next Generation Deepfakes*. IEEE Transactions on Neural Networks and Learning Systems, 36(4), 1023–1040.
- [12] European Commission. (2024). *AI Act 2024: Towards Responsible and Transparent AI Governance*. Official Journal of the European Union.
- [13] UNESCO. (2023). *Recommendation on the ethics of artificial intelligence*. United Nations Educational, Scientific and Cultural Organization.
- [14] The City of New Orleans Office of the Inspector General. (2023). *Facial recognition technology and its use in law enforcement: A policy review*.
- [15] European Data Protection Board (EDPB). (2024). *Guidelines on facial recognition in public spaces*.
- [16] Harwell, D. (2023, August 10). *AI surveillance and racial bias: The New Orleans case*. The Washington Post.

