

A Unified Framework for Multimodal Deepfake Detection: Understanding Video, Audio, Image, and Text Manipulations

Gayatri Devgaonkar, Anisha Dudhane, Gauri Haralkar, Pratiksha Dhobale, Dr. Suresh Mali

*Department of Computer Engineering,

Dr. D. Y. Patil College of Engineering and Innovation, Varale, Pune, India

gdevgaonkar9022@gmail.com, anishadudhane0@gmail.com,

haralkargauri30@gmail.com, dhobalepratiksha46@gmail.com, guide.email@college.edu

Abstract: *In recent years, artificial intelligence (AI) has enabled the creation of deepfakes—highly realistic fake videos, audios, and images that are difficult to distinguish from real ones. Using techniques such as Generative Adversarial Networks (GANs) and other advanced deep learning models, deepfakes can convincingly imitate people's faces, voices, and even writing styles. While these technologies demonstrate the creative power of AI, their misuse has led to serious concerns related to privacy, misinformation, fraud, and identity theft. This paper explores the growing problem of AI-generated manipulation and focuses on the importance of deepfake detection and mitigation systems. It reviews how deepfakes are created, their impact on individuals and organizations, and the existing tools for detection. To enhance detection accuracy, the proposed study integrates the YOLO11 (You Only Look Once, version 11) algorithm—a state-of-the-art object detection model known for its real-time performance and precision in identifying visual manipulations within images and videos. YOLO11's ability to detect subtle inconsistencies, such as abnormal facial movements and mismatched lighting, makes it highly effective for identifying forged visual media. By analyzing current detection tools and implementing YOLO11-based visual analysis along with AI-driven text and audio examination, this research emphasizes the need for stronger and more adaptable technologies. Additionally, it highlights the importance of promoting public awareness and enforcing ethical AI policies. The findings of this paper aim to contribute to developing safer and more trustworthy digital environments where information authenticity can be verified and manipulation through AI-generated media can be effectively reduced.*

Keywords: Deep fake, Artificial intelligence (AI), Machine learning, Neural networks, mitigation, manipulation

