

Authentication Challenges in Generative AI Era: A Mapping Review

1st Wesley dos Reis Bezerra
PPGCC/CTC

Federal University of Santa Catarina
Florianópolis, Brazil
wesleybez@gmail.com
0000-0002-6098-7172

2nd Laís Machado Bezerra
SENAISC

SC, Brazil
lais.machado1@gmail.com
0000-0003-3331-7277

3rd Carlos Becker Westphall
PPGCC/CTC

Federal University of Santa Catarina
Florianópolis, Brazil
carlosbwestphall@gmail.com
0000-0002-5391-7942

Abstract—Authentication and authenticity have been a security challenge since the beginning of information sharing, especially in the context of digital information. With the advancement of generative artificial intelligence, these challenges have evolved, demanding a more up-to-date analysis of their impacts on society and system security. This work presents a mapping review that analyzed 88 documents from the IEEE Explorer, Scopus, and ACM databases, promoting an analysis of the resulting portfolio through six guiding questions focusing on the most relevant work, challenges, attack surfaces, threats, proposed solutions, and gaps. Finally, the portfolio articles are analyzed through this guiding research lens and also receive individualized analysis. The results consistently outline the challenges, gaps, and threats related to images, text, audio, and video, thereby supporting new research in the areas of authentication and generative artificial intelligence.

Index Terms—generative artificial intelligence, genai, authentication

I. INTRODUCTION

Digital resources are important assets across different areas of knowledge, and ensuring their authenticity is a recurring challenge across these fields [1], [2]. Authenticity is an important factor across different media; whether textual, audio, video, photographs, illustrations, or even works of art, the challenges of proving a work's authenticity and authorship have become a major problem [3].

Using generative artificial intelligence (GenAI) tools, malicious users have created fake profiles and violated the authenticity of works and messages worldwide [4], [5]. Whether these violations are intended for profit or to damage others' image, this type of inauthentic resource has been used widely and harmfully. Some adversaries have used GenAI capabilities to deceive less informed people or reinforce prejudices by spreading false messages through social media and using bots.

This work sought to identify the main challenges posed to the authentication and authenticity of digital artifacts by the use of generative artificial intelligence. Furthermore, a systematic literature review was applied, briefly mentioned in the text, which enabled the analysis of a portfolio of works that showed that some aspects of the challenges are associated not with technological issues but with social issues. Among these social issues, the reader's internal bias toward fake news (their

prejudices) or even the lack of knowledge among other media users, such as video users, who are deceived by deepfakes.

The rest of the work is organized as follows: Section II presents an overview of the methodology used for the systematic literature review. Following, Section III provides a brief discussion of the challenges encountered through the systematic review. The paper concludes with the Section IV, which presents the conclusions and vision for the future development of new work in the area.

II. METHODOLOGY

A. Objectives and Research Question

This work focused on security aspects of GenAI, providing a broader view of the challenges and how they materialize in the real world. Some authentication and authenticity challenges have currently emerged in the GenAI field. However, it is important to have a general overview of the applicability of these challenges to existing cyber-physical systems. This general overview allows for a plan for future work to address the challenges encountered.

From this perspective, a broad research question was developed that allows for a focus on applications and practical aspects of GenAI in breaching authentication/authenticity security in cyber-physical systems. This question can be seen below.

"What are the security challenges for authentication brought about by the advent of Generative Artificial Intelligence and their impact on threat generation that can be found in the most relevant work today?"

The hypothesis of creating a document that would support future research in the area was the starting point, and therefore, a more comprehensive question was formulated. The hope is that this document will support the development of new technologies that prepare cyber-physical systems for the challenges posed by advances in Generative Artificial Intelligence. However, to arrive at an adequate answer to this question, we must segment the purposes of this research and seek support from the six guiding questions listed below.

- Q1** What are the most relevant works in the area of authentication and GenAI?
- Q2** What are the security challenges encountered in these works?
- Q3** What are the attack surfaces most targeted by the attacks?
- Q4** What are the main threats listed in the works found?
- Q5** What are the main proposed solutions?
- Q6** What are the gaps found in the selected studies?

Question **Q1** provides an overview of the most relevant research work within the research area, along with an outlook on advances, challenges, and future work to be explored within the framework of these documents. **Q2** focuses specifically on security problems or challenges without limiting areas of activity or the type of security addressed. **Q3** and **Q4** seek to understand which types of systems or resources are the main target of adversaries and the approach to exploiting vulnerabilities in them, respectively. Furthermore, question **Q5** provides an overview of existing solutions within their specific areas and their characteristics. Finally, **Q6** outlines an overview of existing research opportunities within the area and the resolution of still unresolved challenges.

To select the databases, we chose three databases that together represent a large portion of the most well-known conferences and journals in the field of computing. The selected databases were IEEEExplorer, Scopus, and the ACM Digital Library, all of which allow the use of a specific search language and provide tools for advanced searches and filters. Furthermore, these databases allow the export of search content in a standard format, such as BibTeX or CSV.

TABLE I
ARTICLES BY DATABASE

Database	Quantity	Percentage
IEEEExplorer	8 documents	9.5%
Scopus	15 documents	17.6%
ACM	62 documents	72.9%
Total	85 documents	100%

As shown in Table I, a total of 85 documents were retrieved through the search. For the IEEEExplorer database, a total of eight documents (9.5%) were retrieved, while for Scopus, 15 documents (17.6%) were obtained, and ACM had the largest number of documents, with 62 documents (72.9%). Although the sample size is not large, the documents obtained from the three databases represent a concise and non-segmented set of articles.

B. Analysis of Guiding Questions

The final portfolio will be analyzed in accordance with the guiding questions outlined in this study, which are listed in the Methodology section (Section II).

1) *Q1. What are the most relevant works in the area of authentication and GenAI?*: Starting with **Q1**, we can affirm that the selected works are the most relevant within the set of works available at the time of the search. However, because

this is an area undergoing significant research expansion, the dynamics among researchers, authors, co-authors, and funders tend to change. An overview of the most important works is presented in Table II, which summarizes our portfolio data.

2) *Q2. What security challenges have you encountered in these projects?*: Regarding question **Q2**, which addresses the challenges encountered, we can categorize them into: *deep-fake*, lack of regulation, bias and lack of *fairness*, Information Bias Anchoring (IBA), fake news, swatting attacks, other risks involving false emergency alarms, among others. As shown, there are many challenges, and their solutions depend on different approaches. For a better understanding, the portfolio projects and their specific challenges are listed below:

- [6] highlights the main challenges of combating information bias, which is related to both technological aspects that hinder the identification of fake news and psychological aspects that make users inclined to believe fake news.
- [7] for the authors, the challenges are linked to AI regulation, mainly regarding aspects such as data protection and non-discrimination.
- [8] presents the main challenge of preventing the use of audio files published by users to generate deepfakes.
- [9] presents the stress caused to citizens and the unavailability of services due to the use of deepfakes audio in phone calls.
- [10] presents problems arising from image forgery, which pose challenges in its identification. Since forgery identification solutions are specific to each scenario, there is currently no general solution to the problem.
- [11] highlights the problem of ensuring video authenticity even with technological advances in GenAI image generation. As models have advanced, passively detecting altered or AI-generated videos has become increasingly difficult.
- [12] presents privacy as one of the main challenges and details this issue in four dimensions: data, technology, people, and processes.
- [13] explores the use of videos as evidence by police and law enforcement in the era of digital media synthesis.
- [14] identifies verification of audio authenticity based on voice characteristics. Both sighted and blind people evaluated these.
- [15] raises the challenge of identifying hallucinations generated by automatic AI correction modules currently found in photographic devices.
- for [16], the main problem is the inability to identify fakes in photographs, that is, the inability to easily identify a computer-generated image.
- [18] addresses the identification of *fakes* generated by Large Language Models (LLM) used for manipulation and other types of violations on the internet.
- for the authors of [17] the main challenges are the identification of human-made artworks and artworks made by GenAI.

TABLE II
PORTFOLIO - MOST REPRESENTATIVE ARTICLES

#	Ref.	Article
1	[6]	A Perfect Storm: Social Media News, Psychological Biases, and AI
2	[7]	Out of Context: Investigating the Bias and Fairness Concerns of “Artificial Intelligence as a Service”
3	[8]	AntiFake: Using Adversarial Audio to Prevent Unauthorized Speech Synthesis
4	[9]	Not My Voice! A Taxonomy of Ethical and Safety Harms of Speech Generators
5	[10]	A Comprehensive Survey on Methods for Image Integrity
6	[11]	ProvCam: A Camera Module with Self-Contained TCB for Producing Verifiable Videos
7	[12]	Towards a Privacy and Security-Aware Framework for Ethical AI: Guiding the Development and Assessment of AI Systems
8	[13]	Ways of seeing: The power and limitation of video evidence across law and policy
9	[14]	Uncovering Human Traits in Determining Real and Spoofed Audio: Insights from Blind and Sighted Individuals
10	[15]	Advocating Pixel-Level Authentication of Camera-Captured Images
11	[16]	CIFAKE: Image Classification and Explainable Identification of AI-Generated Synthetic Images
12	[17]	Identifying AI-Generated Art with Deep Learning
13	[18]	Deepfakes, Misinformation, and Disinformation in the Era of Frontier AI, Generative AI, and Large AI Models

3) *Q3. What are the attack surfaces most targeted by these attacks?:* Most attacks occur on the web, specifically on social media. The spread of deepfakes and fake news is a growing problem and has become increasingly difficult to identify for users without technical expertise. Due to the increased image, video, and audio generation capabilities provided by AI tools, many people have difficulty identifying whether AI produced a piece of media or not. Below is the analysis for each publication:

- [6] presents social media tools as a platform for the dissemination of fake news and the anchoring of bias.
- [7] for the authors, the threats present themselves as more general challenges to fairness. We can say that the topic addresses three areas: the provider, the user, and legislation.
- [8] presents surfaces that use voice as information, such as voice authenticators, phone calls, audio messages, and public audio publications.
- [9] attacks are carried out through phone calls and occur mainly in schools and citizen service centers.
- [10] the use of generated images can affect communication between people, the safety of children and adolescents, system authentication, political systems, among others.
- [11] for the authors, the dissemination of altered or AI-generated videos on social media has been a problem. This problem is further exacerbated by the complexity of identifying the alteration of these videos.
- for [12], according to the framework proposal, computational resources, data, and system (technology) are the chosen attack surfaces.
- for [13], the legal system and the police are directly impacted, since authenticity challenges directly impact these systems.
- [14] this study does not directly focus on attacks but comments on the vulnerability of people who must base their identification of emotions and individuals solely on voice: blind people.
- [15]’s work highlights the authenticity challenges faced by modern image acquisition devices and AI post-

processing.

- [16] primarily focuses on image-based media, such as art, photography, and any other medium that disseminates these types of digital artifacts.
- [18] highlights information publishing platforms as one of the most affected, especially social media.
- [17]’s most affected entities are companies that price, evaluate, or trade works that humans may not have created.

4) *Q4. What are the main threats listed in the papers found?:* Although there is a large set of relevant threats when it comes to security and Generative AI, here are the main threats existing in our portfolio and selected by the authors. They will be presented and briefly commented on in their specific papers, listed below:

- [6] highlights the spread of fake news, misinterpretation, the uncertainty of significant, frequent, and anonymous information sets, and internal bias.
- [7] presents an argument that can be summarized as the lack of legislation and monitoring to ensure accountability for the lack of fairness in the use of AI as a service.
- [8] discusses the problems posed by the use of deepfake for voice synthesis. Threats include conducting financial scams, compromising security and privacy, spreading hate speech and disinformation, and violating voice authentication.
- [9] presents risks such as swatting attacks, which use synthesized voices to spread false alarms or call emergency services (e.g., fire departments), causing some service units to be unavailable due to their use in false alarms.
- [10] highlights some of the main threats, including encouraging child prostitution, political misconduct, falsifying medical images, and violating system authentication.
- [11], this paper presents image damage and misinformation as the main problems arising from the difficulty in identifying the generation and alteration of videos made by GenAI.
- [12] lacks guarantees of the process’s operation, access to sensitive information, and transparency in the process,

leading to privacy violations. In general, these are problems related to privacy.

- [13] the lack of safeguards for this type of situation, the compromise of evidence, and the lack of preparation of legal systems to address this type of situation.
- [14] the paper discusses the risks of impersonation, audio replay, speech synthesis, and voice conversion. These situations are analyzed by individuals in a study presented in the paper.
- for [15] altering the information contained in the image can lead to misidentification of some elements of the image. This can lead to decision-making based on information inserted by an AI module's hallucination.
- [16] the main threat lies in the failure to identify a computer-generated image. From this point on, several ethical, professional, and legal issues arise in the use of this type of technology.
- [18] cites the effects on democracy and public opinion, impacts on privacy and personal security, consequences for media and journalism, erosion of public trust, and ethical and legal dilemmas as threats.
- [17] incorrect identification can lead to problems in assessing the value of a work, for example. Furthermore, credibility issues can arise in cases where works are misjudged.

5) *Q5. What are the main proposed solutions?:* Each problem demands a different specific solution approach. However, in general, any solution must involve a technological, legal, or cultural approach. The technological approach is critical in identifying digitally generated artifacts, but only education and legal penalties can resolve the situation in the long term. More informed people, with a more open outlook on life and receptive to different opinions, are less susceptible to being deceived and to replicating fake news. In addition, there is a need for legal accountability for the generation and dissemination of fakes by the participating parties. A more detailed discussion per portfolio article follows.

- [6] brings transparency and the reduction of internal bias as objectives for solutions. As a tangible technological solution, it brings the use of eXplainable Artificial Intelligence (XAI) in generative AI.
- [7] presents regulation, surveillance, transparency, and accountability as part of the solution to addressing fairness issues in AI platforms.
- [8] proposes using the solution presented in the article, which processes audio before it is made public. This processing aims to embed information that prevents deep-fakes from being overly precise and difficult to identify.
- [9] presents a framework that allows for a better understanding of the data these threats can generate and follows this classification into three dimensions: data types, specific data, and motivation for the harm.
- [10] presents two types of forgery detection: active detection and passive detection. Active detection can be used through watermarking, digital signatures, and

encryption techniques. Passive detection, on the other hand, consists of applying detection techniques through feature extraction to identify image tampering, identify nonconformities, and distinguish between natural and computer-generated images.

- [11] proposes a solution based on secure hardware that utilizes cryptographic techniques to ensure the origin of video generation and the integrity of the image sequence that composes it.
- [12] the creation of a framework for the ethical use of AI with security and privacy. This can be expressed through the recommendations proposed by the authors.
- [13] improvements in countries' legal systems and the adoption of cutting-edge technologies that can address these types of situations where video authenticity issues may arise.
- [14] proposes the possibility of guaranteeing the origin of audio files and the use of digital watermarks that are perceptible to humans.
- [15] demonstrates a possible solution through the use of spatial masking for altered pixels, the use of metadata, and the pixel-by-pixel identification of the changes made.
- [16] proposes the solution by using its image dataset in conjunction with XAI and a proposed method for identifying fakes.
- [18] presents a framework that divides responsibilities and promotes collaboration between the parties involved, incorporating technology and legal aspects to solve this problem.
- [17] the authors propose the use of a dataset proposed by them, and the adoption of XAI and Gradient Class Activation Mapping (Grad-CAM) in the identification of AI-generated art.

6) *Q6. What gaps were found by the selected studies?:* Due to the recent adoption of GenAI by the general public, new forms of security breaches are emerging every day. Added to this are advances in models and the quality of generated artifacts, which are becoming increasingly closer to reality. In this scenario, many research opportunities are needed to address the gaps that still exist in the literature. Below, we will list some gaps found in each selected article.

- [6] proposes the use of XAI as part of the content generation solution. This would be a way to identify whether the generated content is sufficiently transparent and free of bias.
- [7] presents the lack of regulation and transparency as important gaps in the topic of fairness.
- [8] applying pre-processing to audio before publishing remains complex for the end user. The lack of different samples of adequate length also poses challenges for research, as well as ethical issues and bias in approaching the topic.
- [9] highlights the need for: distinguishing between testable and untestable data, studying model outputs, and understanding the potential harm caused by exposure,

ways to mitigate these types of harm, and legal support for accountability through policies and civil society interventions.

- for the authors of [10], research opportunities exist primarily in the new areas of image modification detection, namely image anomaly detection and deep anomaly detection. The authors also classify the use of deep learning in anomaly detection as one of the main trends.
- in [11], the authors outline improvement proposals for the application of their solution to real devices such as drones, robots, cameras, cell phones, among others.
- also for [12], the need to evolve the framework in line with the evolution of algorithms and the application of AI in various areas.
- in [13], there are opportunities for advancement in countries' legal systems, in the development of new technologies for identifying deepfakes and authenticating videos to be used as evidence.
- for [14], the advancement of audio synthesis has generated results very close to reality, allowing even individuals more sensitive to the different characteristics of voice to be deceived. Research into source identification and human-identifiable watermarks is necessary.
- [15] highlights research opportunities in how to ensure that the mask and metadata generated during acquisition are not altered, as well as how to store such information.
- [16] argues how gaps represent opportunities for using different classification techniques within their dataset, the use of other XAI approaches, and the continued evolution of the dataset to keep pace with new advances in image generation.
- [18] proposes continued improvement in identification methods to keep pace with GenAI developments, adaptability to address new forms of disinformation, adoption and implementation of standards, and education of information consumers. The work also complements the list of six open research opportunities: technological advancements, global cooperation, public engagement, behavioral insights, economic models, and technological innovations.
- already [17] proposes an ensemble approach that combines multiple models for a more robust output.

III. DISCUSSION

This work adopted a literature review methodology to map the challenges of using generative artificial intelligence for the authentication and authenticity of digital artifacts. Thus, challenges related to the use of AI-generated texts to spread disinformation (i.e., fake news) were analyzed; the generation of audio and video for impersonation was evaluated, including the creation of deepfakes across various media. Furthermore, it was observed that public services are precarious due to AI-generated fake calls, falsification of works, and legal implications for both AI service providers and users (i.e., police), reflecting a need for improvement in the current legislation of various countries.

Therefore, intelligence, like any tool, is subject to the adoption criteria of its users. We perceive that the more power a tool possesses, the more prepared and well-intentioned its user must be, whether consumer or content producer. Content producers must be aware of their responsibilities in the information dissemination process and clearly present generated media as untrue or text as fictional. Furthermore, regarding production, the tools used to generate inauthentic media, for which they are paid, should also share responsibility for creating a world better prepared to address this type of artifact and be partially accountable for crimes committed with those tools. Finally, consumers must be increasingly attentive to media and texts that do not reflect reality, and be aware of the need to verify facts before disseminating them.

However, despite many challenges, artificial intelligence is already being applied to solve various types of issues that previously required significant effort. We can cite the dynamic modeling of threats [19], [20], in partnership with machine learning in the detection and prevention of intrusion [21] and in the identification of the authorship of texts through techniques such as stylometry [22], avoiding security breaches in social media messages. Beyond the challenges, generative artificial intelligence offers diverse opportunities for research and development across various areas, especially computing, as previously mentioned.

IV. CONCLUSION AND FUTURE WORK

We can affirm that this work successfully identified some of the main authentication challenges encountered in this era of generative artificial intelligence. This work lists articles from a portfolio found through a systematic literature review of important portals for the scientific community in computing. The works were examined through six questions that help refine the search for challenges in this area and better understand them. Additionally, a discussion was held on the main results and how we can use them to advance research in this area of knowledge.

However, improvements in this mapping review and further developments were observed. It was noted that greater detail is needed for the sub-areas cited as presenting challenges. For example, authorship and authenticity identification (impersonation identification) is an important area that should be further explored across its different facets: content, writing style, speech characteristics, and video characteristics. In this way, the challenges of this area of authentication and GenAI would be better understood.

REFERENCES

- [1] W. Ibrar, D. Mahmood, A. S. Al-Shamayleh, G. Ahmed, S. Z. Alharthi, and A. Akhuzada, "Generative ai: a double-edged sword in the cyber threat landscape," *Artificial Intelligence Review*, vol. 58, no. 9, p. 285, 2025.
- [2] E. Ferrara, "Genai against humanity: Nefarious applications of generative artificial intelligence and large language models," *Journal of Computational Social Science*, vol. 7, no. 1, pp. 549–569, 2024.
- [3] J. Collomosse and A. Parsons, "To authenticity, and beyond! building safe and fair generative ai upon the three pillars of provenance," *IEEE Computer Graphics and Applications*, vol. 44, no. 03, pp. 82–90, 2024.

- [4] M. Casu, L. Guarniera, P. Caponnetto, and S. Battiato, "Genai mirage: The impostor bias and the deepfake detection challenge in the era of artificial illusions," *Forensic Science International: Digital Investigation*, vol. 50, p. 301795, 2024.
- [5] L. Sandrini and R. Somogyi, "Generative ai and deceptive news consumption," *Economics Letters*, vol. 232, p. 111317, 2023.
- [6] P. Datta, M. Whitmore, and J. K. Nwankpa, "A perfect storm: social media news, psychological biases, and ai," *Digital Threats: Research and Practice*, vol. 2, no. 2, pp. 1–21, 2021.
- [7] K. Lewicki, M. S. A. Lee, J. Cobbe, and J. Singh, "Out of context: Investigating the bias and fairness concerns of "artificial intelligence as a service"," in *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*, 2023, pp. 1–17.
- [8] Z. Yu, S. Zhai, and N. Zhang, "Antifake: Using adversarial audio to prevent unauthorized speech synthesis," in *Proceedings of the 2023 ACM SIGSAC Conference on Computer and Communications Security*, 2023, pp. 460–474.
- [9] W. Hutiri, O. Papakyriakopoulos, and A. Xiang, "Not my voice! a taxonomy of ethical and safety harms of speech generators," in *Proceedings of the 2024 ACM Conference on Fairness, Accountability, and Transparency*, 2024, pp. 359–376.
- [10] P. Capasso, G. Cattaneo, and M. De Marsico, "A comprehensive survey on methods for image integrity," *ACM Transactions on Multimedia Computing, Communications and Applications*, vol. 20, no. 11, pp. 1–34, 2024.
- [11] Y. Liu, Z. Yao, M. Chen, A. Amiri Sani, S. Agarwal, and G. Tsudik, "Provcam: A camera module with self-contained tcb for producing verifiable videos," in *Proceedings of the 30th Annual International Conference on Mobile Computing and Networking*, 2024, pp. 588–602.
- [12] D. Korobenko, A. Nikiforova, and R. Sharma, "Towards a privacy and security-aware framework for ethical ai: Guiding the development and assessment of ai systems," in *Proceedings of the 25th Annual International Conference on Digital Government Research*, 2024, pp. 740–753.
- [13] S. Ristovska, "Ways of seeing: The power and limitation of video evidence across law and policy," *First Monday*, 2023.
- [14] C. Han, P. Mitra, and S. M. Billah, "Uncovering human traits in determining real and spoofed audio: Insights from blind and sighted individuals," in *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*, 2024, pp. 1–14.
- [15] A. Punnappurath, L. Zhao, A. Abdelhamed, and M. S. Brown, "Advocating pixel-level authentication of camera-captured images," *IEEE Access*, 2024.
- [16] J. J. Bird and A. Lotfi, "Cifake: Image classification and explainable identification of ai-generated synthetic images," *IEEE Access*, vol. 12, pp. 15 642–15 650, 2024.
- [17] T. Bianco, G. Castellano, R. Scaringi, G. Vessio *et al.*, "Identifying ai-generated art with deep learning," in *CREAI@ AI* IA*, 2023, pp. 16–25.
- [18] M. R. Shoaib, Z. Wang, M. T. Ahvanooey, and J. Zhao, "Deepfakes, misinformation, and disinformation in the era of frontier ai, generative ai, and large ai models," in *2023 International Conference on Computer and Applications (ICCA)*. IEEE, 2023, pp. 1–7.
- [19] J. A. Da Silva, C. B. Westphall, and W. D. R. Bezerra, "Intrusion detection of data manipulation attack in iot network based on fog computing," in *2026 40th International Conference on Information Networking (ICOIN)*. IEEE, 2026, pp. 894–899.
- [20] W. dos Reis Bezerra, C. A. de Souza, C. M. Westphall, and C. B. Westphall, "Hyvaw: A hybrid vulnerability analysis workflow threat model methodology for complex systems based on distributed components," *Journal of Open Innovation: Technology, Market, and Complexity*, p. 100654, 2025.
- [21] W. R. Bezerra, E. A. Galhardo, L. M. Bezerra, and C. B. Westphall, "Threats and ai trends in threat modeling for 5g/6g," in *2026 40th International Conference on Information Networking (ICOIN)*. IEEE, 2026, pp. 645–649.
- [22] C. G. Butzke, M. L. Brocardo, W. R. Bezerra, and C. B. Westphall, "Identifying fakes in social media text messages using stylometry," in *2026 40th International Conference on Information Networking (ICOIN)*. IEEE, 2026, pp. 581–586.