

The Adverse Effects of Utilisation of AI-Generated Content by Older Adults: A Scoping Analysis of Distribution Channels, Vulnerabilities, and Countermeasures

Yuhan Tang

*Beijing University of Chemical Technology, Guiyang, China
1679848399@qq.com*

Abstract. The application of generative AI technology in the creation of videos and audio contents increases its spread rate but also increases chances that fabricated information will mislead public opinions when not rigorously verified; Simultaneously it makes illegal activities easier as identity is stolen or abused for different purposes within massive scale without being easily detected. Driven by gaps in digital literacy, verification habits and social-psychological needs, older people are more likely to fall victim to fake media on short-video apps, social networks and other peer-to-peer channels. At the same time, this phenomenon also leads to financial losses for parents; it harms the mental state or physical safety of children; and disrupts the consistency of daily routines and development activities. The current mitigation tactics are mainly clustered at four levels: personal digital literacy education, family-based intergenerational assistance, platform-content regulation, and regulatory collaboration. However, it is still necessary to consider the extent of its influence in practice and how much evidence there is to back it up currently. Based on the research results, this article constructs a unified "technology-cognition-behaviour" framework and puts forward concrete recommendations for families in China, internet platforms, and the general public.

Keywords: AI-generated content, deepfakes, voice cloning, misinformation, older adults

1. Introduction

The development of generative AI has greatly reduced the barriers to producing high-quality audio-visual works; with a small amount of input, an ordinary person can generate realistic video material featuring individuals, news reports, or clarifications [1]. At the same time, some such innovations have also been exploited by evil people to manufacture false information, impersonate others for illegal gains, and autonomously generate enormous amounts of content with purposes. As a result, there has been an increase in the authenticity problems and the verification difficulties of digital information [2]; Therefore, there is a general increase in the extent to which they have occurred over time.

There are three common ways that the misuse of AI manifests in video communications. First of all, it is more realistic, that is, the improvement of facial expressions, lip-syncing and sound effects

can make people more convincing [3], Second, higher levels of interaction experience. Finally, there is authority impersonation, which involves imitating professionals, broadcasters, spiritual leaders, and official accounts [4]. In addition, the scale of distribution can be realised through large-scale content production and synchronous social networks; therefore, it is relatively simple for false information to spread via feeds and reach all corners of society [5].

Seniors are more prone to being influenced by this kind of material in the online media environment not only due to their "diminished judgmental abilities" but also because multiple factors affect them collectively [6]. Firstly, source validation often involves some extra work and the use of various information channels; On the other hand, social network interactions and credibility signals may weaken people's motivation to verify [7]. In addition, older people are also more susceptible to affective rhetoric because of loneliness, the desire for company, and concerns about their health; combined with being easily swayed, they often believe in emotionally charged information that offers an immediate solution [8].

This study systematically summarises the main manifestations and transmission channels of AI-fabricated media misconduct; Factors that cause older people's weak ability to verify information; The monetary losses, mental anguish, physical health problems, psychological maladjustment, etc., caused by being deceived; Existing proof robustness and gaps in current prevention strategies. Based on this, the article presents a unified "Technology-Cognition-Behaviour" structure and offers concrete improvement suggestions.

2. Conceptual definitions and theoretical framework

In this study, "AI-generated media abuse" refers to the activity of using generative models to create false data in various forms, such as videos, images, and audio; or to impersonate others; and then disseminate these materials improperly. It includes, but is not limited to, deepfake/face-swapping videos [9], voice mimicry reproducing friends or authority figures; simulated hosts or experts explaining events [4], and a large amount of promotion and opinion material [5].

The term "senior citizens" has different definitions in various studies based on varying age criteria [10]. A 60+/uniform standard has been set for all comparisons in this survey. Nevertheless, the existing studies mainly focus on the scenarios with higher relevance to media integration; that is, in the scenario of short video applications, social network updates, collective message sharing and portable/voice-based interaction.

3. Methods

This research uses the scoping review method to map out an evidence framework on the impact of AI abuse on older adults and identify investigation gaps and directions; it is carried out in four steps: "identification - screening - extraction - synthesis" [11].

Keywords were divided into three categories: AI-generated media (including "deepfake technology", "authenticity indicators", "algorithmic amplification"), associated hazards (including "financial fraud against older people", "disinformation and psychological problems", "reduced trust in institutions"); The second group consists of terms related to the elderly, including "elderly people", "loneliness", "susceptibility to misinformation".

The inclusion criteria were: the research subjects were older people or age-disaggregated data was provided; the participants interacted with AI-generated content, which included abuse and risks of misinformation or deception; there were no restrictions on the research method (cross-sectional study, experiment, qualitative analysis, or intervention). Exclusion criteria entailed: First, technical

detection analysis without population-level analysis; Second, the sample did not include older adults or lacked age classification; Third, there were standard creative uses unrelated to abuse hazards.

The obtained dataset covered the following information: Country/Region; Sample attribute and age range; Platforms and distribution environment, form of maltreatment; Propagation method; Result metric (including confidence level, victimisation rate, mental impact, wellness practice), etc. Research framework and the main research achievements. Maltreatment types include physical mistreatment; Spread through various networks, including digital media and virtual games, which young people are more prone to (online); At-risk groups mainly consist of children under ten years old, marginalized children, and disabled children, etc. The result is to hinder individual progress and fail the family's expectations in quality assessment; At the same time, this will have an adverse impact on social development outside the scope of daily life at home.

Given that there is considerable variation in the research methods across different studies, this article primarily employs qualitative synthesis to compare the evidence rigour of these studies. In terms of the upper limit of causal deduction and the future improvement direction, they will be elaborated on in detail later.

4. Results

4.1. Forms of AI-generated media abuse

Based on the cited literature, it has been found that there is a relatively distinct mainstream tendency in terms of the current forms of abuse. In summary, there are four specific risk areas that can be identified, as shown in Figure 1; each has its own technical route, method of transmission, and level of danger.

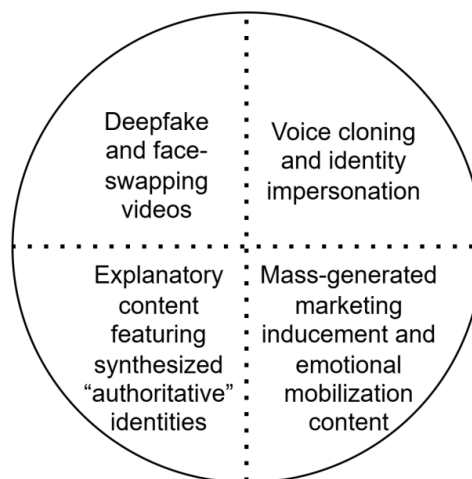


Figure 1. Four forms of AI-generated media abuse

Deepfakes and face-swapping movies (1). Exploitation in this category primarily involves inserting or replacing the victim's facial identity into a source clip, or merging such graphic modifications with audio recordings to create the individual's voice, which is then utilised to imitate the personality of the original person in video content [12]. With the continuous progress in the construction of these frameworks, more profound research has been carried out on facial expression simulation technology; Such as lip-synchronisation accuracy and Voice alignment, etc., Which can

create a very realistic visual or auditory effect. It will make it much more difficult for ordinary people to subconsciously assess the authenticity of the output. Deepfake videos that occur in real life not only spread false words of famous people or spread incorrect information, but also take the form of what seems to be typical expressions, such as "clarification" and "announcement", which increases their likelihood of being believed.

Acoustic Replication and Person Simulation (2) Acoustic Replication and Person Simulation: Audio synthesis technology can replicate the unique voice characteristics of a target person based on a small amount of audio data; Then, it performs identity spoofing during phone calls, recorded sound files, and real-time communication scenarios. Based on the current literature, the achievement of voice replication techniques has reached a relatively high level in terms of authenticity and functionality; However, the verification tools for fake speech have difficulties in both false positive rates and strong dependence on context, making it difficult to be efficiently applied in actual conversational scenarios [13]. Typically, some people use more advanced techniques to simulate relatives, friends, or staff members of organisations; At the same time, typical discourse patterns such as "crises" and "urgent needs" are employed to create psychological pressure that aims to prompt a response from the individual at risk immediately.

(3) Explanation materials with combined "credibility" personae: This category includes amalgamated digital representations of experts, journalists, clergy, etc., who hold symbolic positions of authority in society, expressed as forms such as "scientific outreach," "expert analysis," or "critical warnings" [4]. Based on previous research, when audiences encounter explanatory data with clear professional indicators, they often believe it to be credible information and neglect to further verify the credibility of the originator or the basis for this information. Falsified authoritative information about medical data, financial advice, and explanations of social risks, among others, is still quite common. Typical deception generally does not cause such a misleading effect; it is usually some curated representations of reality, an exaggeration of the threat factor or a simplification of complex issues.

(4) Mass-produced promotional materials that trigger incentives and appeal to emotions: This particular abuse takes advantage of the low cost and high efficiency of generative architectures to produce a large number of consistent internal content, which is then spread through account networks or automatic publishing methods [14]. In general, the output focuses on short-term transactions, emotional interactions, and collective conflicts, among other things; at the same time, it is emphasized that there is an immediate call for response, such as liking, sharing, purchasing, and donating. Given that platform recommendation systems have an impact, content that is more entertaining and emotionally resonant is more often promoted; therefore, it spreads further.

Broadly speaking, the types of AI-generated content that often present difficulties for older people fall into two categories. One type includes high-authority cues, which are divided into three subtypes: obvious identity recognitions, technical terms used, and conventional delivery forms [15]. The other type is powerful emotional cues, which cover fear, anger caused by disease, and exclusive needs that must be met in personal time. When these parts are combined with the former parts, it can significantly increase older people's sense of security and shorten the cycle required for data verification or interaction with peers [8]; In turn, this increases vulnerability.

4.2. Dissemination and reach mechanisms

Within which, the included research suggests that the distribution and effects of AI-synthesized content are not results from a single path; instead, they are formed by the combined influence of "platform-suggested expansion + social-circle-sharing + expert-obfuscation".

In terms of the platform tier, short-form video and social networking sites generally use feed-recommendation engines to deliver content [16]. By constantly enhancing the target exposure's tuning by incorporating, into our recomposition algorithm that pays attention to behavioral data on users in each group, more appealing videos will be shown in front of people. Synthetic media can quickly appear in recommendation feeds, offer lower creation costs and stronger expression; therefore, it may break through the initial audience limit.

At the social level, collective message and information distribution among peers have become important ways of dissemination outside of algorithmic recommendations. Compared with the "systematic approval" provided by platform engines, such transmission occurs in close relationships and is based on people's long-term interaction experience. According to some studies, since the data is distributed among relatives, neighbours and locals in a close-knit relationship within the community, there tends to be an initial trust phenomenon among listeners that no further verification is needed [7]. For older people, this supplementary verification based on personal relationships is more prone to maintain false information in their daily life's information environment and spread extensively within relatively separate social networks.

Next, we will focus on how to improve recognitions of special materials using this prediction method; At the same time, consider enhancing model generalisation performance for diverse types of images. From the investigation, we find that materials with strong emotional resonance, clear narrative structures, and direct calls to action do indeed increase likes, comments, and posts; Algorithms then categorise this content as "high-quality" and promote it more widely [17]. The primary assessment criterion is the strength of consumers' feedback on issues rather than the factuality of information; It measures the resilience of individuals in participation and interaction. As a result, some "inferred content" that is not entirely deceptive will still occupy visible positions in streams because it has the ability to evoke emotions; Therefore, better but less captivating material may be marginalised.

In terms of peer-network data circulation, the standard for determining whether it is genuine has often been simplified as follows: Is the originator reliable? and thus become "Is the individual reliable?" ", which is especially obvious when AI-generated materials masquerade as minors, friends or figures of power. At this time, older people's experience of operating information technology is mainly based on trust in interpersonal relationships and pre-digital role identification, lacking empirical verification or digital proof. If some materials use the format of "crisis storytelling", such as "a mishap occurred", "immediate funds are needed" and "details not yet available", this will also reduce the chance of Confirmation and Interaction, making it much more dangerous [18].

In short, the platform's recommendation mechanism creates an environment where information is highly visible to all users; At the same time, the diffusion of social connections strengthens the perception of reliability based on personal identity recognition; Meanwhile, the imitation of experts combined with crisis storytelling accelerates cognition and emotional response. In summary, due to the influence of these three factors mentioned above, artificial intelligence-generated data may serve a function that is difficult to guard against when conveyed to older people and can become one of the means through which fraud is carried out and false information spreads.

4.3. Vulnerabilities of older adults in information verification

As shown in table 1, the problems raised by the elderly at the validation stage in these cited researches are categorised into three major categories across all these references: A relatively narrow range of fairly limited means of verifying AI-generated multimedia content, inferior technical execution skills; They generally rely more heavily on authority-based signals when accepting data;

When facing a large amount of information, heuristic strategies are often used rather than a thorough examination whenever emotionally stimulated.

Generally speaking, this kind of situation is rarely due to insufficient skills; as these systems are complex and display methods vary along with unique circumstances during delivery, the hazards have been magnified [6].

In terms of the mental stage, older people have increased execution costs for data validation. Origin traceability, multi-platform inquiry and backward checking generally involve several process stages and require specific digital abilities; it usually takes some time to complete. In addition, because the fake information appears in a "specialised" form, that is, through established interpretive frameworks, professional jargon and graphic presentations, older people tend to associate the official credibility of these expert formats with their authenticity; Therefore, they often do not question the authenticity of the data source or the sufficiency of evidence [15]. At this time, "superficial verisimilitude" has a perturbing effect on people's perception.

The wish for companionship, fear of physical decline or frailty in old age, and doubt about social change are necessary conditions that influence older people's acceptance of new technologies. In today's digital environment, short-cut discussions that reach quick conclusions, offer simple solutions, or promptly console people are more likely to capture public attention and gain trust. Whether the style of AI-generated content is defined as "compassionate warning", "expert analysis" or "in your best interest", it may reduce protection mental barriers and enable viewers to accept data more easily, rather than taking a cautious attitude [8]. The effect of the interaction among affect, perception and logical reasoning on the persuasiveness for older people is substantial.

One less important reason is the lack of support from family and friends. Upon identifying suspicious information and unable to contact relatives or verify with trustworthy companions, older adults are more inclined to make judgments based on instinct or trust in the disseminator. At this time, due to the delay or forgetfulness of the validation action, misleading information is more likely to pass undetected [19].

Table 1. Difficulties in identifying types of synthetic media misuse

	Deepfake and face-swapping videos	Voice cloning and identity impersonation	Explanatory content featuring synthesized "authoritative" identities	Mass-generated marketing inducement and emotional mobilization content
Ease of Information Verification	Hard	Moderately Hard	Hard	Moderately Hard
Credibility	Relatively High	Relatively High	High	High
Emotional Dependence	High	High	High	Relatively High

4.4. Negative impacts

Monetary Security: Most current research has recorded the threat of deception by family members or officials [4] and the financial loss caused by uncontrollable shopping or tipping desires. Especially when both "crisis situations" and "identity fraud" happen at once, the victims usually cannot act immediately without personal information; After losing money, they are unable to recall suspicious persons or gather materials that help the police crack the case. After the incident, it takes some time to restore because of the complex financial path and multi-platform operations [20].

Mental Health: After a long period of exposure to fear-tending information, older people may experience feelings of anxiety; Such situations often cause emotional distress and affect their sleep at night; At the same time, they may also give rise to other mental problems. After senior citizens realise they have been deceived, they may feel ashamed, criticise themselves for being too cautious, or question their own ability to think clearly, thus reducing their desire to seek help or formally handle the problem [21].

Health practices: misuse of drugs, failure to regular attend the clinic because of wrong health prevention data; increasing medical expenses due to misunderstanding one's own symptoms. When articulated through "professional analysis" or "reliable advice" in integrated media, the obfuscation of its fraudulence is usually greater, and audiences may not be able to discern problems merely based on their first impressions [22].

Cognitive framework and social confidence: Continuous contact with unverified or frequently repeated false information may cause older people to doubt reliable data channels and authoritative institutions, and they may seek close-knit or single information sources. Once dependence is generated, there will be a strengthened tendency of selective perception, the reinforcement of the echo-chamber effect, and an increase in the incidence of false information; consequently, it has a negative impact on the ability to evaluate information and on the construction of social trust structures over time [23].

4.5. Evidence and limitations of response strategies

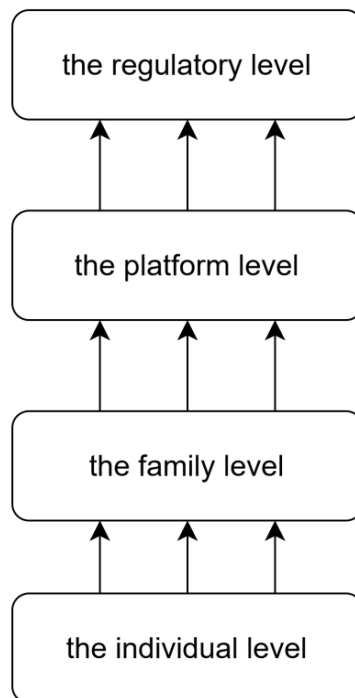


Figure 2. The general paths currently used to mitigate the risks of AI-generated content

Based on the above research, the general paths currently used to mitigate the risks of AI-generated content are mainly divided into four levels, as shown in Figure 2: personal digital competence education [24], family multi-generational assistance, platform supervision and

legislative cooperation. At the same time, due to differences in operating costs, the range of coverage, validity period, etc., at these levels, they target different risk factors and stages.

In terms of individuals, digital literacy education is the most widespread, and there is the most evidence that any intervention works. Generally speaking, research shows that if the coaching of identifying disinformation, verifying origin and enhancing hazard awareness is strengthened, older adults will be more vigilant about suspicious information in the short term and eager to verify it [24]. Most of the data comes from laboratory tests, questionnaires or simulated situations; There is an uncertain factor that these results cannot be replicated in actual information environments. Considering that new types of risks such as deepfake films and voice copying are emerging, it is still unknown after an in-depth exploration over a long period and during practical application whether the existing training resources will continue to be effective against this kind of cross-media deception caused by the novel synthetic media [24].

In terms of the domestic situation, intergenerational support provides a critical buffer for individuals to avoid making high-risk financial choices. Kinship units can develop particular earning plans, standardised safety codes, and so on; when conducting large-scale transactions, there is generally a process of "double verification before each discharge", which prolongs the individual's cognitive process for monetary flow changes and helps users regain logicity after sudden feelings or temporary psychological pressure. At present, there are already some methods to solve problems such as auditory replication and mimicry deception in semi-protected dialogue, crisis scenarios, etc., [19] Although the construction process has not been clearly stated. After all, family-level involvement still belongs to the situation of domestic configuration and regular interaction among its members; therefore, not every older person can obtain help from their relatives or enjoy continuous assistance.

In terms of the platform framework, many studies have proposed various governance strategies to reduce the dissemination power and spread capability of fake media, such as deepfake warnings, media tags, hazard blocking, user identity management, and prioritising "suspicious materials", etc., which may be used before the content is viewed by individuals or in the early stage to prevent extensive harm. Platform regulation entails several trade-offs: First, it is necessary to regulate the false positive rate and avoid over-censorship of legal production and expression; At the same time, improving the transparency and clarity of the regulatory process can also enhance users' trust in the service. Based on existing literature, there is still a problem of data notification and hazard warning for the elderly, which has not been thoroughly addressed; how to provide relevant alerts without overloading their mental capacity continues to be a significant issue in this area at the platform level [25].

In terms of policy-level governance, some research has proposed that it is necessary to establish a liability chain across multiple links such as content creation and distribution through platforms, financial circulation monitoring, etc., to define the responsibility framework for different entities in the generation of AI-generated media fraud. In addition, by enabling the exchange of illegal patterns and risk information across various platforms, it will help improve the recognition effect on a large volume of deceptive and solicitation activities. Based on this premise, specifically design the older generation as the main protected group to help develop more targeted defensive strategies and initiatives, including instructions, technology application, supervision and management, etc., so as to establish a synchronous regulatory mechanism of "instruction - technology - monitoring" [26]. The operational effect of supervisory actions depends on statutes, rules and enforcement power; currently there are problems with harmonising diverse areas, adapting to changes in technology, etc., so it is unclear if they can be sustained for a long time.

5. Conclusion

This article gathers various manifestations of adverse effects caused by the misuse of AI-generated content among older adults and constructs a comprehensive "technology-cognition-behaviour" framework to systematically map these risks. The research indicates that the heightened authenticity and official-like appearance of synthetic media significantly increase the probability of misinterpretation among seniors, who often lack the specialized digital skills required for verification. This vulnerability is further exacerbated by the synergy between platform recommendation engines and peer-to-peer recognition within social networks, which together expand the reach and perceived credibility of misinformation. Furthermore, high verification costs—including the time and technical ability needed for origin traceability—coupled with socio-psychological factors like loneliness and the desire for companionship, jointly amplify the potential harm caused by digital deceptions.

References

- [1] Mehta, P., Jagatap, G., Gallagher, K., Timmerman, B., Deb, P., Garg, S., ... & Dolan-Gavitt, B. (2023). Can deepfakes be created by novice users?. *arXiv preprint arXiv: 2304.14576*. [21]
- [2] Patel, N., Jain, R., & Mathew, G. J. (2025). Unveiling the Influence of Misinformation and Deceptive AI-Generated Content on Gen Z: A Comprehensive Study. *Advances in Consumer Research*, 2(3).
- [3] Marchal, N., Xu, R., Elasmir, R., Gabriel, I., Goldberg, B., & Isaac, W. (2024). Generative AI misuse: A taxonomy of tactics and insights from real-world data. *arXiv preprint arXiv: 2406.13843*.
- [4] Kang, F., Cao, Y., & Chen, H. (2025). Face2VoiceSync: Lightweight Face-Voice Consistency for Text-Driven Talking Face Generation. *arXiv preprint arXiv: 2507.19225*.
- [5] Yu, L., Mottola, G., Kieffer, C. N., Mascio, R., Valdes, O., Bennett, D. A., & Boyle, P. A. (2023). Vulnerability of older adults to government impersonation scams. *JAMA Network Open*, 6(9), e2335319-e2335319.
- [6] Corsi, G. (2024). Evaluating Twitter's algorithmic amplification of low-credibility content: an observational study. *EPJ Data Science*, 13(1), 1-15.
- [7] Yunita, N., Juhana, A., & Sari, I. P. (2025). Enhancing digital literacy in the elderly through media and technology to prevent hoaxes in Indonesia: A systematic literature review. *Literasi: Jurnal Ilmu Pendidikan*, 16(1), 87-101.
- [8] Majerczak, P., & Strzelecki, A. (2022). Trust, media credibility, social ties, and the intention to share towards information verification in an age of fake news. *Behavioral Sciences*, 12(2), 51.
- [9] Gu, C., Li, X., & Xiang, Q. (2025). The infinite monkey theorem in AIGC advertising: Matching effects between AI disclosure and advertising appeals on consumer advertising avoidance intention. *Journal of Research in Interactive Marketing*, 1-20.
- [10] Al-Khazraji, S. H., Saleh, H. H., Khalid, A. I., & Mishkhal, I. A. (2023). Impact of deepfake technology on social media: Detection, misinformation and societal implications. *The Eurasia Proceedings of Science Technology Engineering and Mathematics*, 23(429-441), 2.
- [11] Ye, Z., Zou, X., Post, T., Mo, W., & Yang, Q. (2022). Too old to plan? Age identity and financial planning among the older population of China. *China Economic Review*, 73, 101770.
- [12] Mak, S., & Thomas, A. (2022). Steps for conducting a scoping review. *Journal of graduate medical education*, 14(5), 565-567.
- [13] Kaur, A., Noori Hoshayr, A., Saikrishna, V., Firmin, S., & Xia, F. (2024). Deepfake video detection: challenges and opportunities. *Artificial Intelligence Review*, 57(6), 159.
- [14] Wang, K., Chen, M., Lu, L., Feng, J., Chen, Q., Ba, Z., ... and Chen, C. (2025, May) One borrowed phrase: Evaluate the risks of voice synthesis in the AIGC period. In *2025 IEEE Symposium on Security and Privacy (SP)*: pp. (4663-4681). IEEE is an academic community.
- [15] Zhang, Y., Song, W., Koura, Y. H., & Su, Y. (2023). Social bots and information propagation in social networks: Simulating cooperative and competitive interaction dynamics. *Systems*, 11(4), 210.
- [16] Spasova, L. (2025). Susceptibility to authority figures in advertising with persuasive strategies—age differences in perception. *Business & Management Compass*, 69(3), 33-47.
- [17] Narayanan, A. (2023). Understanding social media recommendation algorithms.

- [17] Pathak, R., Spezzano, F., & Pera, M. S. (2023). Understanding the contribution of recommendation algorithms on misinformation recommendation and misinformation dissemination on social networks. *ACM Transactions on the Web*, 17(4), 1-26.
- [18] José-Cabezudo, R. S., & Camarero-Izquierdo, C. (2012). Determinants of opening-forwarding e-mail messages. *Journal of Advertising*, 41(2), 97-112.
- [19] Cui, S., Zhao, Y., & Qie, R. (2025). Who needs What Support? Exploring the relationship between intergenerational support and digital media use among Chinese older adults: A latent profile analysis. *Computers in Human Behavior*, 164, 108506.
- [20] Ueno, D., Arakawa, M., Fujii, Y., Amano, S., Kato, Y., Matsuoka, T., & Narumoto, J. (2022). Psychosocial characteristics of victims of special fraud among Japanese older adults: A cross-sectional study using scam vulnerability scale. *Frontiers in Psychology*, 13, 960442.
- [21] Kemp, S., & Erades Pérez, N. (2023). Consumer fraud against older adults in digital society: Examining victimization and its impact. *International Journal of Environmental Research and Public Health*, 20(7), 5404.
- [22] Zhou, J., Xiang, H., & Xie, B. (2023). Better safe than sorry: a study on older adults' credibility judgments and spreading of health misinformation. *Universal Access in the Information Society*, 22(3), 957-966.
- [23] Ecker, P. A. (2025). *The Digital Reinforcement of Bias and Belief: Understanding the Cognitive and Social Impact of Web-Based Information Processing*. Springer Nature.
- [24] Choudhary, H., & Bansal, N. (2022). Addressing digital divide through digital literacy training programs: A systematic literature review. *Digital Education Review*, (41), 224-248.
- [25] Siapera, E. (2022). Platform governance and the "infodemic". *Javnost-The Public*, 29(2), 197-214.
- [26] Rebovich, D., & Corbo, L. (2022). The distillation of national crime data into a plan for elderly fraud prevention: A quantitative and qualitative analysis of US postal inspection service cases of fraud against the elderly. In *The new technology of financial crime* (pp. 126-149). Routledge.