

Calculated Trust: The Dilemma of Older Adults in Judging the Authenticity of AI-Generated Videos in the Age of Algorithms

Zhengxi Cao

*College of Arts and Media, Tongji University, Shanghai, China
2644430866@qq.com*

Abstract. With the explosion of large AIGC video models, AI videos are increasingly appearing on short video platforms, and the older adult population is gradually showing signs of discomfort and difficulties in the digital environment. Taking WeChat Channels as the research context and drawing on system trust theory, this paper systematically explores the formation mechanism behind older adults' dilemmas in verifying the authenticity of AI-generated videos. Through a questionnaire survey (N=210), the study found that the dilemma faced by older adults is actually a transfer of system trust — from "trust in the platform and algorithm system" to "trust in specific AI video content." The vulnerable position of older adults determines the intensity of this transfer, while the type of AI video content triggers this transfer. By introducing algorithm trust into the study of older adults, this research provides theoretical reference for digital literacy interventions for this group.

Keywords: System trust, authenticity judgment, older adults, AI video, algorithm awareness

1. Introduction

Short video platforms like WeChat Channels have become a primary medium for daily entertainment and social interaction among middle-aged and elderly groups. However, fueled by the advancement of large-scale AI video models, AIGC-generated content has proliferated extensively on these platforms impacting older internet users' ability to discern information and their trust perceptions. With its low entry barrier embedded within WeChat, the natural trust of social circles, and content aligned with user preferences, WeChat Channels has become a core platform for older adults to engage in digital life. Against the backdrop of the Internet evolving from a pure information intermediary to an algorithm-driven intelligent platform, the validity of traditional trust mechanisms is facing unprecedented academic scrutiny. On platforms like WeChat Channels, system trust is reflected in users' trust in the platform's recommendation mechanism and recommended content. When the platform fails to filter out AI-generated false content, users' system trust can easily turn into a source of bias in judging the authenticity of video content—a phenomenon known as improper transfer of system trust. In recent years, research on older adults' AI literacy and identification challenges has begun to emerge. Studies have explored educational pathways for AI literacy among older adults, age-friendly learning environment design, and the role of social support

and intergenerational collaboration in enhancing digital resilience [1-4]. However, there is still a lack of systematic theoretical explanation regarding the mechanisms behind older adults' biases in judging AI video authenticity. Drawing on the System Trust Theory, this study conducts an empirical analysis of 210 valid questionnaires, so as to systematically explore the formation mechanism of older adults' cognitive bias in identifying AI-generated short videos. By applying system trust theory to the context of AI dissemination, this research provides a reference for older adults' adaptation to digital media and contributes to building a more inclusive and safe digitally aging society.

2. Literature review and research questions

2.1. Algorithmic non-neutrality and characteristics of AIGC misinformation

The non-neutrality of algorithms leads to algorithmic biases and information cocoons, which in turn distort platform recommendations and create an environment for the spread of false AIGC videos. AIGC-produced content is mainly characterized by multimodal deception and content homogenization, causing users' information environments to become increasingly narrow and misleading [5]. Various types of AI videos have also spread widely on short video platforms.

When middle-aged and elderly users are long exposed to misleading, straightforward AI content, their critical thinking and deep thinking abilities tend to gradually erode, making them more inclined to rely on "authority certification" [6]. When emotional information is consistent across audio and video channels, observers can integrate natural audio cues with synthetic video cues to make judgments [7]. While the internet breaks through spatial boundaries, it also undermines traditional trust boundaries. The perceived legitimacy of a system is closely tied to the traits and professional competence demonstrated by its agents [8]. The pre-existing trust that middle-aged and elderly users place in familiar celebrities can directly transfer to AI face-swapped content, thereby reducing their willingness to verify its authenticity. Therefore, this paper proposes the following hypothesis:

H1: Different types of AI video presentations have significant differences in the bias of authenticity judgments among middle-aged and elderly users.

2.2. Algorithmic disadvantage and information judgment dilemma among middle-aged and elderly groups

The middle-aged and elderly are attracted to AI Face-Swap content, which essentially stems from a significant emotional compensation phenomenon. Their content preferences often correspond to emotional deficits and social difficulties in real life, aiming to gain 'substitute satisfaction' and express symbolic rights [9]. Meanwhile, the digital divide and generational digital gap explain why they are easily deceived. Differences in digital cognition across generations, combined with the inherent physical and mental disadvantages of middle-aged and elderly individuals, inadequate platform adaptation for elderly users, and a lack of family digital support, collectively weaken their risk assessment capabilities [10]. The older the respondents, the lower they self-rate their AI literacy; their technical understanding and application skills lag significantly behind younger groups, showing insufficient confidence in 'technical controllability,' along with a strong sense of helplessness and tendency to avoid technology [11]. Therefore, this study proposes the following hypotheses:

H2: The vulnerable characteristics of middle-aged and elderly people (age, education, and algorithm awareness) are significantly positively correlated with their bias in judging the authenticity of AI videos.

H2-1: Older middle-aged and elderly people tend to have higher trust in AI videos, with more obvious judgment bias.

H2-2: Middle-aged and elderly people with weaker algorithm awareness are more likely to believe in the authenticity of AI videos.

The selective sharing behavior of elderly groups is associated with their self-efficacy. Elderly individuals with high self-efficacy tend to actively verify information to reduce uncertainty and determine whether to share content based on their own judgments [12]. Therefore, this study proposes the following hypothesis:

H3: The bias in middle-aged and elderly people's judgment of AI video authenticity has a significant positive effect on their behavior (liking, sharing, reposting, and purchasing).

H3-1: The greater the judgment bias, the more likely middle-aged and elderly people are to like and share AI videos.

H3-2: The greater the judgment bias, the more likely middle-aged and elderly people are to take actual actions due to AI video content (such as buying products or believing health information).

The research on the AI literacy and recognition dilemma of elderly people involves topics such as literacy education pathways, age-friendly learning environment design, social support, and intergenerational collaboration, but few studies systematically explain the mechanisms behind elderly people's bias in judging the authenticity of AI videos from a theoretical perspective [1-4]. This study seeks to link algorithms, AIGC content, and the middle-aged and elderly population together, focusing on how the vulnerable status of this group affects their judgment of AI video authenticity, the role of different types of AI videos in this process, and behaviors and attitudes after trusting the algorithms (Figure 1). It systematically analyzes the formation mechanisms and related influences of the AI video authenticity dilemma among middle-aged and elderly people in an algorithm-driven environment.

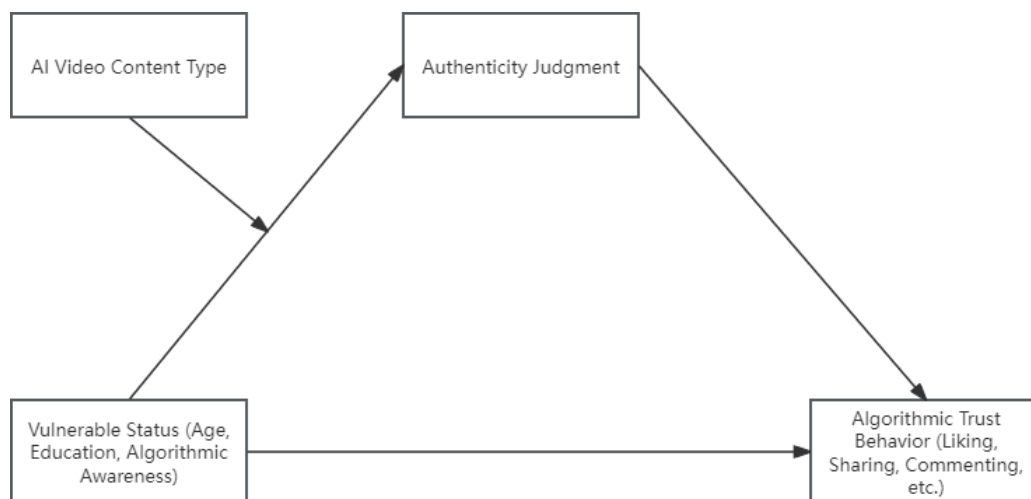


Figure 1. Model of AI video trust and judgment in older adults

3. Research methods

A questionnaire survey was employed in this study to address the proposed research questions. The questionnaire design was based on mature scales and contextually adapted for the characteristics of middle-aged and elderly groups. The questionnaire consisted of six sections: respondents' basic demographic information, algorithm awareness, viewing frequency and authenticity trust judgment for four categories of AI videos, algorithmic inducement and usage behaviors, post-viewing

behaviors related to AI videos, and overall perceptions and attitudes after exposure to AI-generated content.

In this study, indicators of vulnerable status were quantified as age, education level, and algorithm awareness. The algorithm awareness scale was adapted from the Algorithm Media Content Awareness (AMCA) scale developed by Zarouali et al., which has a Cronbach's α of 0.82 and a KMO value of 0.79. AI video content type was used as a moderating variable, and authenticity judgment bias was quantified as the user's trust score in AI videos, to study the impact of different video presentation styles on middle-aged and elderly people's bias in judging authenticity [13]. This referred to the Face-Swap trust scale developed by Plohl et al., with reliability coefficients for each dimension ranging from 0.85 to 0.92 [14]. This study also referred to the algorithm inducement scale developed by Li, Yu, and Tian to measure users' perception of platform algorithm inducement behavior, with a Cronbach's α of 0.88 and a KMO value of 0.84 [15].

4. Research findings

This study used SPSS for data analysis. Through descriptive statistics, group difference tests, regression analysis, and mediation effect tests, all the research hypotheses mentioned earlier were verified.

4.1. Testing differences in trust across different types of AI videos

The results (Table 1) reveal highly significant differences in trust across the four AI video types. Face-Swap videos had the highest trust level ($M = 4.38$), followed by AI videos with pictures and voiceovers ($M = 4.22$) and AI digital human videos ($M = 4.10$), while fully AI-generated scenario videos had the lowest trust level ($M = 3.14$). The data showed $F(3,627) = 68.86$, $p < 0.001$, $\omega^2 = 0.171$. The mean trust ratings for the first three types of videos were all above 4.10, while the purely AI-animated videos only scored 3.14, a gap of 1.24 points. When middle-aged and older adults judge the authenticity of AI videos, they rely on cues of anthropomorphism or realism in the videos. Celebrities in deepfake videos, virtual anchors in digital human videos, and real images/voices in AI picture-voice videos offer trust-transferring cues, while purely AI-animated videos lack these anthropomorphic features, failing to build such trust.

Table 1. Descriptive statistics and repeated measures ANOVA of trust levels for four types of AI videos

Video Type	Mean	Standard Deviation	95% Confidence Interval	F	p	ω^2
AI Digital Human	4.1	1.12	[3.95, 4.25]			
AI Illustrated + AI Voiceover	4.22	1.13	[4.07, 4.37]			
AI Face-Swap	4.38	1.02	[4.24, 4.52]			
AI Fully Generated Scenario	3.14	0.94	[3.01, 3.27]			
Overall Difference				68.86		

Table 1. (continued)

	<0.001	0.171
--	--------	-------

4.2. Testing the relationship between vulnerable status factors and trust

Research shows that age is positively correlated with trust, and algorithm awareness is negatively correlated with trust, while education level does not show a significant correlation (Table 2). Age exhibits a moderately strong significant positive correlation with trust ($r = 0.564$, $p < 0.001$): older respondents tend to trust AI videos more, and readily extend their platform trust to platform-recommended content. Algorithm awareness has a moderately strong significant negative correlation with trust ($r = -0.477$, $p < 0.001$). Middle-aged and elderly users with higher algorithm awareness tend to trust AI videos less and are more inclined to question and intervene in recommended content. However, the correlation between education level and trust is not significant ($r = -0.025$, $p = 0.717$), possibly because the education distribution in this study's sample is relatively concentrated, or perhaps education is not a key indicator of vulnerable status. Middle-aged and elderly people's ability to judge AI videos may rely more on experience, like feedback from family or help from the community, rather than formal education. Age and algorithm awareness show a strong negative correlation ($r = -0.596$, $p < 0.001$), indicating that factors within vulnerable status are not isolated but interrelated.

Table 2. Pearson correlation matrix between variables and trust level

Variables	1	2	3	4	5
1. Age	—				
2. Education	-0.012	—			
3. Usage Duration	—	—	—		
4. Algorithm Awareness	-0.596*	0.032	—	—	
5. Average Trust Score	0.564*	-0.025	—	-0.477*	—

The results (Table 3) show that the model is overall significant ($F = 36.77$, $p < 0.001$). Age, education, and algorithm awareness together explain 34.9% of the variance in trust ($R^2 = 0.349$). Age is the biggest factor affecting trust ($\beta = 0.434$, $p < 0.001$); for each increase in age group, the average trust score goes up by 0.327 points, suggesting that as people get older, they rely more on long-term experience and habits to make judgments. Algorithm awareness is a significant protective factor ($\beta = -0.218$, $p = 0.002$); for each 1-point increase in algorithm awareness, the average trust score drops by 0.240 points. Enhancing algorithm literacy among middle-aged and older adults serves as a viable strategy to mitigate their authenticity judgment bias toward AI videos. Users with strong algorithm awareness know that algorithms personalize content for them, so they don't accept platform recommendations uncritically. Consequently, H2, H2-1 and H2-2 are fully supported, whereas H2-3 fails to obtain statistical support.

Table 3. Multiple linear regression analysis of structural vulnerability factors on trust level

Variables	Unstandardized Coefficient B	Standard Error	Standardized Coefficient β	t Value	P Value	95% CI
(Constant)	3.228	0.38	—	8.498	<0.001	[2.480, 3.976]
Age	0.327	0.053	0.434	6.194	<0.001	[0.223, 0.431]
Algorithm Awareness	-0.24	0.077	-0.218	-3.108	0.002	[-0.392, -0.088]
Education	-0.008	0.033	-0.013	-0.232	0.817	[-0.073, 0.058]

4.3. Testing the relationship between trust and behavior

The results show (Table 4) that trust has a significant positive effect on behavior ($\beta = 0.461$, $p < 0.001$), with the model explaining 21.2% of the variance ($R^2 = 0.212$). For every 1-point increase in trust, the behavior score goes up by 1.065 points. Middle-aged and older adults' trust in AI videos does translate into actual actions. This also shows that the risk of AI videos misleading older adults is real. But it's worth noting that trust only accounts for 21.2% of the behavior variance ($R^2 = 0.212$), meaning nearly 80% of the variance is explained by other factors, leaving room for future research. So, H3 is supported.

Table 4. Linear regression analysis of trust on behavior

Variables	Unstandardized Coefficient B	Standard Error	Standardized Coefficient β	t Value	P Value	95% CI
(Constant)	0.757	0.493	—	1.536	0.126	[-0.214, 1.728]
Average Trust Score	1.065	0.142	0.461	7.482	<0.001	[0.785, 1.345]

4.4. Mediation effect test

This section further examines whether vulnerability factors influence behavior via trust. In this study, age, education level, and algorithm awareness are the independent variables, trust level is treated as the mediator, and behavior scores are the dependent variable. This mediation analysis adopts a three-model regression framework (Table 5). Model 1 (total effect) shows that age ($\beta=0.285$, $p<0.001$) and algorithm awareness ($\beta=-0.598$, $p<0.001$) have significant direct effects on behavior, and the model explains 64.1% of the variance in behavior ($R^2=0.641$). Model 2 (path a) shows that age ($\beta=0.434$, $p<0.001$) and algorithm awareness ($\beta=-0.218$, $p=0.002$) have significant effects on trust, and the model explains 34.9% of the variance in trust ($R^2=0.349$). Model 3 (direct effect) shows that after including trust, the direct effects of age ($\beta=0.274$, $p<0.001$) and algorithm awareness ($\beta=-0.593$, $p<0.001$) on behavior are still significant, but trust has no significant effect on behavior ($\beta=0.023$, $p=0.652$). Although path a ($X \rightarrow M$) is significant, path b ($M \rightarrow Y$) is not

significant ($p=0.652$), and the 95% confidence interval of the indirect effect includes 0, so no mediation effect is present.

The effects of age and algorithm awareness on behavior are significant, but these effects are not transmitted through trust. When age and algorithm awareness are included in the model, the independent effect of trust disappears. This suggests that the influence of trust on behavior is largely due to differences in age and algorithm awareness. Those with higher algorithm awareness tend to have lower trust and lower behavior frequency. Vulnerability factors affect both trust and behavior directly, and trust does not play a decisive role.

Table 5. Results of mediation effect analysis

Variables	Model 1 (Total Effect)	Model 2 (Path a)	Model 3 (Direct Effect)
Behavior	Average Trust Score	Behavior	
Age	0.285***	0.434***	0.274***
Education	0.012	-0.013	0.013
Algorithm Awareness	-0.598*	-0.218	-0.593*
Average Trust Score			0.023
R ²	0.641	0.349	0.642
F Value	122.714***	36.773***	91.730***

5. Conclusion

The study found significant differences in trust levels among the four types of AI videos. Videos with humanoid features are trusted much more than those without such features, suggesting that the celebrity image's trust transfer effect is not the strongest trust cue; convincing images and human voices can also make viewers believe the video content. Such AI videos activate a judgment pattern in middle-aged and older adults that relies on source credibility heuristics from the traditional media era, essentially transferring system trust into the algorithmic environment. Age and algorithm awareness are key factors influencing judgments of AI video authenticity, but they work in opposite directions. Older users, with relatively limited cognitive resources, rely more on trust-based heuristics and skip individual verification of each video. The study also found that once age and algorithm awareness are included in the model, trust itself loses its influence. Middle-aged and older adults with high algorithm awareness behave less because they maintain a cautious attitude toward AI content, meaning the suppressive effect of algorithm awareness on behavior doesn't depend on reduced trust. Older users may have developed habits of automatically liking, sharing, and commenting through long-term use of WeChat Video Channels and Douyin. These habits aren't formed through rational evaluation. In fact, the behavioral variations explained by trust are actually absorbed by age and algorithm awareness. There are also some limitations in this study. The online-sourced sample consisted of participants with pre-existing digital skills, potentially underestimating the overall bias. Behavior data was self-reported, so social desirability bias and recall bias are present.

References

- [1] Ko, E. G., Nanayakkara, S., & Huff Jr, E. W. (2025). Litti: Empowering older adults' AI literacy through generative AI experience. In *CHI EA '25: Extended abstracts of the CHI conference on human factors in computing systems*. New York: ACM.
- [2] Liu, D., & Li, J. (2026). How teacher–learner–AI collaboration empowers AI literacy of older adults: A case study of an older adult school in China. *Humanities and Social Sciences Communications*.
<https://doi.org/10.1057/s41599-026-08471-7>
- [3] Kırmacı, Ö., & Florez-Revuelta, F. (2026). Features for designing AI-supported learning environments for older people. *Studies in the Education of Adults*, 1–22. <https://doi.org/10.1080/02660830.2026.2669407>
- [4] Stanford Impact Labs. (2025). Empowering diverse digital citizens: Intergenerational digital storytelling project. <https://impact.stanford.edu/investment/empowering-diverse-digital-citizens>
- [5] Duan, Y. H., Zhang, H. J., et al. (2025). Principles and industry applications of AI video generation technology. *Peking University*.
- [6] Zhang, J. (2026). AI virtual digital human influencers vs human influencers: The impact of health short videos on older adults' cognition and attitude. *Frontiers in Public Health*.
- [7] Gao, Y., et al. (2025). Trust in virtual agents: Exploring the role of stylization and voice. *IEEE Transactions on Visualization and Computer Graphics*, 31(5), 3623–3633.
- [8] Hu, B. J., & Li, Y. J. (2015). Internet and the reconstruction of trust. *Contemporary Communication*, (4), 19–25.
- [9] Pan, K. W., Huang, C. C., & Li, X. L. (2025). Study on micro-drama viewing behavior and its influencing factors among middle-aged and elderly groups. *Modern Publishing*, (1), 50–64.
- [10] Tao, X. D., & Sun, Y. Z. (2025). The digital generation gap and bridging strategies in family communication: A study based on the use of emojis. *Modern Communication (Journal of Communication University of China)*, 47(2), 159-168. <https://doi.org/10.19997/j.cnki.xdcb.2025.02.005>
- [11] School of Media and Communication, Shanghai Jiao Tong University. (2025). Research report on the development of generative artificial intelligence and digital communication. *Shanghai Jiao Tong University*.
- [12] Cairang, Z. M., & Song, Q. X. (2025). Between real and fake: How the elderly judge and respond to the authenticity of AIGC short videos. *Philosophy and Social Science Preprint Platform*.
<https://zsybyb.cn/abs/202511.03775> [PSSXiv:202511.03775V1]
- [13] Zarouali, B., Boerman, S. C., & de Vreese, C. H. (2021). Is this recommended by an algorithm? The development and validation of the algorithmic media content awareness scale (AMCA-scale). *Telematics and Informatics*, 62, Article 101607.
- [14] Plohl, N., Mlakar, I., Aquilino, L., et al. (2025). Development and validation of the perceived deepfake trustworthiness questionnaire (PDTQ) in three languages. *International Journal of Human-Computer Interaction*, 41(11), 6786–6803.
- [15] Li, Z., Yu, X., & Tian, K. (2026). Algorithm-associated digital addiction among older adults: Mechanisms and public health implications for healthy aging. *Frontiers in Public Health*, 13, Article 1746304.