
Original Paper

AI Awareness and Tobacco Policy Messaging Among US Adults: Electronic Experimental Study

Julia Mary Alber^{1*}, MPH, PhD; David Askay^{2*}, PhD; Anuraj Dhillon^{2*}, PhD; Lauren Sandoval^{1*}, BS; Sofia Ramos^{3*}; Katharine Santilena^{1*}, MS, MPH, PhD

¹Department of Kinesiology and Public Health, California State Polytechnic University, San Luis Obispo, CA, United States

²Department of Communication Studies, California Polytechnic State University, San Luis Obispo, CA, United States

³Department of Biological Sciences, California Polytechnic State University, San Luis Obispo, CA, United States

*all authors contributed equally

Corresponding Author:

Julia Mary Alber, MPH, PhD
Department of Kinesiology and Public Health
California State Polytechnic University
1 Grand Ave
San Luis Obispo, CA 93407
United States
Phone: 1 8057561779
Email: jmalber@calpoly.edu

Abstract

Background: Despite public health efforts, tobacco use remains the leading cause of preventable death in the United States and continues to disproportionately affect underrepresented populations. Public policies are needed to improve health equity in tobacco-related health outcomes. One strategy for promoting public support for these policies is through health messaging. Improvements in artificial intelligence (AI) technology offer new opportunities to create tailored policy messages quickly; however, there is limited research on how the public might perceive the use of AI for public health messages.

Objective: This study aimed to examine how knowledge of AI use impacts perceptions of a tobacco control policy video.

Methods: A national sample of US adults (N=500) was shown the same AI-generated video that focused on a tobacco control policy. Participants were then randomly assigned to 1 of 4 conditions where they were (1) told the narrator of the video was AI, (2) told the narrator of the video was human, (3) told it was unknown whether the narrator was AI or human, or (4) not provided any information about the narrator.

Results: Perceived video rating, effectiveness, and credibility did not significantly differ among the conditions. However, the mean speaker rating was significantly higher ($P=.001$) when participants were told the narrator of the health message was human (mean 3.65, SD 0.91) compared to the other conditions. Notably, positive attitudes toward AI were highest among those not provided information about the narrator; however, this difference was not statistically significant (mean 3.04, SD 0.90).

Conclusions: Results suggest that AI may impact perceptions of the speaker of a video; however, more research is needed to understand if these impacts would occur over time and after multiple exposures to content. Further qualitative research may help explain why potential differences may have occurred in speaker ratings. Public health professionals and researchers should further consider the ethics and cost-effectiveness of using AI for health messaging.

JMIR AI 2025;4:e72987; doi: [10.2196/72987](https://doi.org/10.2196/72987)

Keywords: artificial intelligence; tobacco control; health communication; generative artificial intelligence; health policies

Introduction

Despite some progress in tobacco control and prevention, tobacco remains the leading cause of preventable death in the United States [1], with persistent inequities in health outcomes among youth, rural residents, people with mental

health illnesses, communities of color, and the Lesbian, Gay, Bisexual, Transgender, Queer or Questioning, Intersex, Asexual + community [2,3]. Secondhand smoke exposure disproportionately impacts Black people compared to any other racial group, putting them at higher risk for asthma and long-term lung disease [2,4]. Public policies are essential for

addressing health inequities associated with tobacco use and exposure [2].

While there is evidence that policies can decrease smoking behaviors [5], there is a need to implement more comprehensive policies to promote and achieve health equity in the United States. Given this need, more countries and US jurisdictions are moving toward “Endgame” policies [6]. Endgame policies focus on strategies that address existing health inequities and target commercial sales of tobacco products rather than individual use. Existing research in New Zealand indicates some support for Endgame policies [7-9]. For example, flavored tobacco products have been specifically targeted to communities of color and youth, which in turn has led to higher tobacco-related disease and death among these communities [10-13]. One potential Endgame strategy to reduce these inequities is a policy banning the sale of tobacco products.

While research has examined the overall effectiveness of some tobacco control policies, there is limited research on how to best promote support for these types of policies. Developing and testing messaging for promoting tobacco Endgame policies is essential for promoting local, state, and national efforts in the United States to address tobacco-related inequities. However, developing and testing high-quality content, particularly video content, can be expensive and time-consuming, given that it can require a significant investment of resources [14].

With limited resources, public health organizations such as local tobacco control coalitions often lack the time, funding, or expertise to design high-quality content. Generative artificial intelligence (AI) offers tools to address these challenges by enabling the rapid creation of targeted messages, supporting simple analyses of communication features (eg, pause counts), and producing high-quality videos [15-18]. Platforms such as ChatGPT (OpenAI) can generate scripts or social media text within seconds, while other programs can create videos featuring digital avatars that deliver messages with customizable voices and appearances. These capabilities can reduce costs, shorten production time, and allow messages to be tailored to specific audiences. For example, one study found that AI-generated health messages were rated higher in quality and clarity than human-generated messages from Twitter (xAI; [19]), although the findings were limited to a sample of college and young adult women.

With this in mind, AI messages can help tailor messages to a specific population. In addition, the use of AI, such as ChatGPT, can be used to decrease language barriers in health care and simplify medical jargon to make it more accessible [20]. Incorporating AI in health messaging has been associated with several positive outcomes, including increased clarity, improved message tailoring, and reduced language barriers in health communication [19,20]. Within tobacco control, AI is already being used to personalize quitting interventions (eg, chatbots for cessation support) and to monitor tobacco-related trends (eg, social media surveillance) [21].

Overall, AI has the potential to revolutionize health communication and allow public health organizations to create quality content more quickly and cost-effectively. However, understanding how the public perceives the use of AI for health messaging is important. Research highlights several barriers to accepting messages involving AI, including concerns that such messages may be perceived as less trustworthy or relatable, and discomfort associated with AI entities mimicking human characteristics (commonly referred to as the “uncanny valley”) [22]. For example, though traditionally applied to visual AI, the “uncanny valley” may also occur with synthetic voices that mimic human voices but may lack natural emotion or fluctuation, potentially eliciting audience discomfort [22]. These factors can diminish the perceived effectiveness and credibility of public health messages.

Existing theoretical frameworks can help explain how audiences perceive AI in health messaging. The elaboration likelihood model (ELM) describes 2 routes of information processing: the central and peripheral [23]. When individuals have little interest in or connection to a topic, they are less likely to critically think about the information presented [23] and instead rely on heuristics such as source credibility to make decisions [23]. For example, the source or speaker in a message and their perceived credibility could influence individuals’ decisions to follow or not follow the message. In the case of AI-delivered messages, positive perceptions of AI may increase acceptance of the message, whereas negative perceptions of AI may lead audiences to discount or ignore it.

The computers as social actors (CASA) paradigm further shows that people apply social expectations such as trust and empathy to digital agents [24]. Recent work has extended CASA to health communication contexts, demonstrating that perceived characteristics of AI sources can shape message credibility and audience engagement [25]. In health messages, when trust and credibility are vital to the processing of a message, audiences may evaluate AI speakers as more or less credible than human ones, even when delivering identical content. These frameworks highlight key factors shaping how audiences interpret and respond to AI-delivered health information.

To the authors’ knowledge, no existing research has specifically investigated how perceptions of an AI-generated speaker may impact message outcomes in an experimental design using both qualitative and quantitative data. Therefore, the purpose of this study was to explore whether awareness that a speaker is AI-generated influences the message outcomes of a tobacco policy video. We hypothesized that participants who were informed that the narrator was AI-generated would rate the video lower in overall quality, perceived effectiveness, credibility, and speaker evaluations compared with those who were told that the narrator was human, unknown, or not informed.

Methods

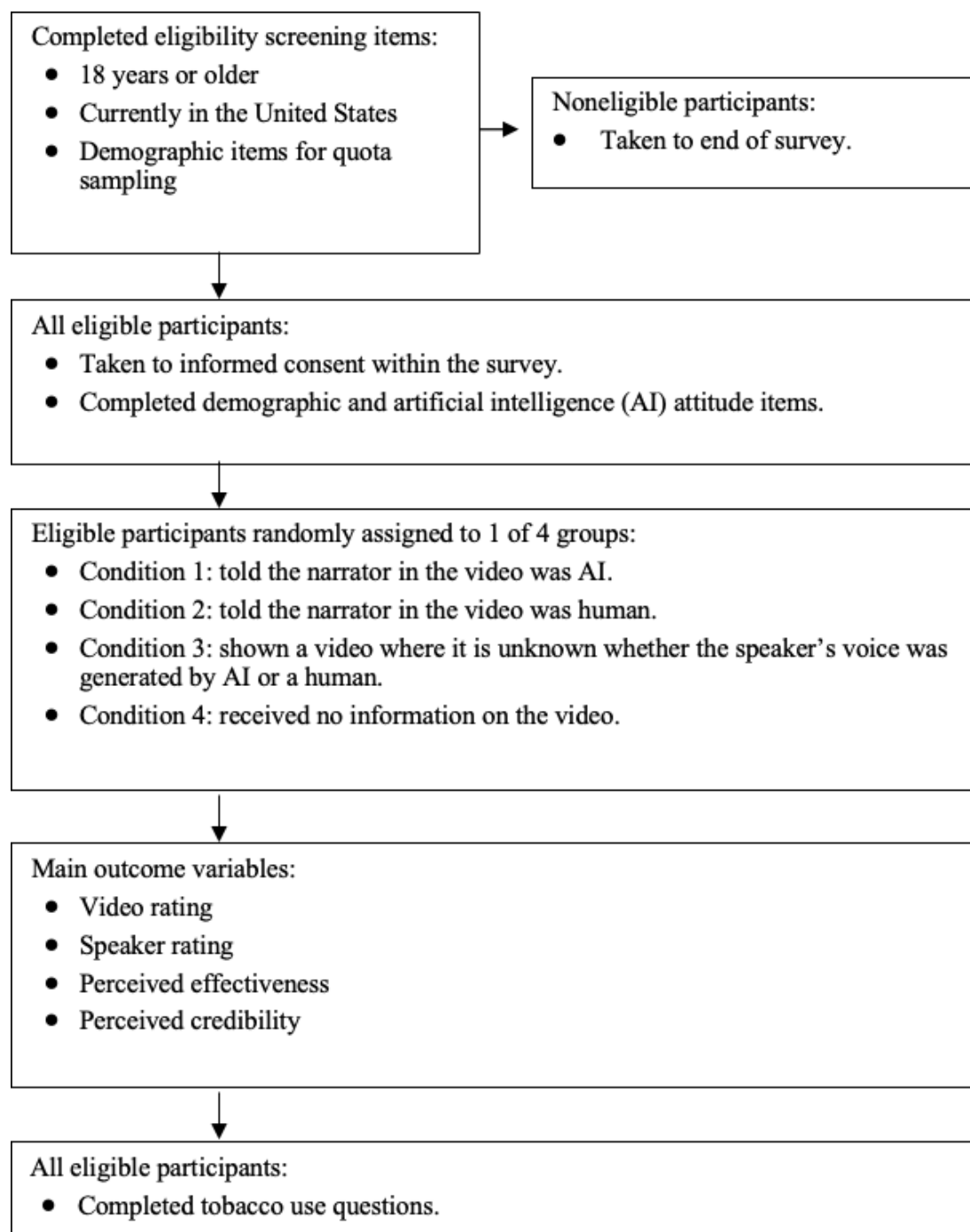
Recruitment and Data Collection

Data were collected between January 25 and February 1, 2024. US adults (N=500) were recruited through the Qualtrics Survey Panel for an electronic survey. Quota sampling was used to approximate the US adult population distribution for age, sex, race, and ethnicity based on the American Community Survey from the US Census Bureau [26]. A priori power analysis using G*Power (Erdfelder, Faul, and Buchner) indicated that a sample of this size was sufficient to detect a small effect ($f=0.15$) with 80% power at $\alpha=.05$ in a 4-condition ANOVA.

Participants were recruited through an electronic survey panel (Qualtrics Panel) that maintains a large pool of individuals who voluntarily register to participate in research. Qualtrics Panels use a standard recruitment email message that includes information on compensation, time to complete, and an anonymous survey link, which is sent to the registered pool of participants. Participants who meet eligibility criteria, with quotas applied to match US Census benchmarks for age, sex, and race and ethnicity, are profiled and invited to studies. Participants are compensated for participation according to the panel's standard compensation options, which can vary based on participation. The panel company managed recruitment and provided researchers with the final sample. The Qualtrics Panel has been used to collect data for other studies examining health beliefs and message testing [27-29].

The procedures for the survey are summarized in [Figure 1](#). Within the survey, each participant was randomly assigned to 1 of 4 conditions. First, participants were exposed to a video narrated by an AI developed using PlayHT (voice-generating software; Mahmoud Felfel and Hammad Syed). The video depicted a narrator telling a story about how their addiction to flavored vapes started and how this addiction controlled their life. At the end of the video, a message promoting the end of flavored tobacco sales was displayed. The only thing participants saw in the video was the words being spoken and a consistent background. Participants were then presented with different text information from 1 of 4 conditions: (1) told the narrator in the video was AI, (2) told the narrator in the video was human, (3) told it was unknown whether the narrator's voice was generated by AI or a human, and (4) given no mention of the narrator's source, allowing participants to experience the video naturally without bias or preconceived notions about its creation. Despite participants being informed that the narrator's voice was human, AI-generated, unknown, or not mentioned, deception was used. The same video featuring an identical AI-narrator voice was used across all conditions. This strategy ensured consistency, minimized variability, and allowed the researchers to focus exclusively on the effects of the informational context provided to participants. Several strategies were used to identify inattentive or fraudulent participants, including attention-check measures, checking the accuracy of open-ended responses, and speed checks ([Multimedia Appendix 1](#)).

Figure 1. Overview of study procedures for a randomized electronic survey assessing perceptions of an artificial intelligence (AI)-narrated tobacco control video among US adults, 2024.



Measures

Demographic variables were used to collect information on participants' sex, race, ethnicity, marital status, income, and education, as well as their tobacco use, based on established surveys [1,28]. In addition, AI attitudes were assessed by having participants rate their agreement from 1 (strongly disagree) to 5 (strongly agree) for 8 statements (eg, "I think artificial intelligence is dangerous") [30]. An average was then calculated for positive and negative AI attitudes (Cronbach $\alpha=0.79$ and 0.71 , respectively). Four main message outcome measures were used: video rating,

perceived effectiveness, perceived credibility, and speaker rating.

Video rating was measured by asking participants how much they liked or disliked the video on a scale from 1 (like a great deal) and 7 (dislike a great deal) [31]. To measure perceived effectiveness, participants rated their agreement from 1 (strongly disagree) to 5 (strongly agree) for 6 statements about the video being convincing, compelling, persuasive, effective, saying something important to them, and putting thoughts in their mind about supporting a flavored tobacco sales ban. The items were developed based on previous research (Cronbach $\alpha=0.91$) [31,32].

Perceived credibility (Cronbach $\alpha=0.89$) was assessed by asking participants to rate their level of agreement from 1 (strongly disagree) to 5 (strongly agree) that the information was accurate, believable, and factual [33]. To measure speaker rating (Cronbach $\alpha=0.97$), participants rated the person speaking in the video from 1 (not at all) to 5 (extremely) on 20 characteristics (bright, cares about me, competent, concerned with me, credible, ethical, expert, genuine, has my interest at heart, honorable, informed, intelligent, knowledgeable, likable, moral, not self-centered, sensitive, trained, trustworthy, and understanding). These items were adapted from previous research [28,34,35]. Scores for speaker rating, perceived effectiveness, and perceived credibility were created by averaging across the items in each scale.

Statistical Analysis

SPSS (version 29; IBM Corp) was used to examine frequency and central tendency measures of demographic characteristics and tobacco use variables. To check randomization efficacy, chi-square and 1-way ANOVA tests were used to examine if any condition differences existed among participant characteristics. One-way ANOVAs were used to examine if condition differences existed among video rating, perceived effectiveness, perceived credibility, and speaker rating. In addition to ANOVA statistical tests, linear regression analyses were conducted to determine whether conditions predicted message outcomes (ie, speaker rating, perceived effectiveness, and credibility) while controlling for variables (ie, age, gender, Black or African American, Latinx, education, marital status, flavored product use, lifetime cigarette use, and AI attitudes). Dummy coding was used for the regression analysis, with the AI-narrated condition (Condition 1) set as the reference group, allowing comparisons between this condition and each of the other conditions.

Qualitative Analysis

A thematic analysis was performed to categorize participants' written comments for each video [36]. The process involved

2 trained research assistants creating initial codes using a sample of comments, defining codes, pilot testing the codes until acceptable intercoder reliability was established ($\geq 90\%$ agreement), coding the remaining data, and grouping the codes into themes. Only codes representing at least 5% of participants were included in the themes. Theme frequency was then explored by participant conditions using NVivo Qualitative Data Analysis Software (version 16; Lumivero).

Ethical Considerations

Before data collection, California Polytechnic State University's Institutional Review Board approval was obtained (approval #2023-223-OL). All participants provided informed consent electronically before participating in the study. Data were collected anonymously and participants were compensated according to their survey panel's normal compensation options.

Results

Demographic Characteristics

Demographic information is presented in Table 1. Participants were, on average, the age of 48.42 (SD 17.66) years, with an equal split identifying as male ($n=245$, 49%) and female ($n=250$, 50%). The majority of participants identified as White ($n=385$, 77.0%), while approximately 48.7% ($n=215$) of participants reported making less than US \$50,000 annually. The majority of participants had more than a high school education ($n=398$, 79.8%).

In terms of tobacco use, 175 (35%) participants reported having smoked at least 100 cigarettes in their lifetime, and 146 (29.2%) participants reported ever using flavored tobacco products.

Table 1. Demographic characteristics, tobacco use, and message outcome measures among US adults ($N=500$) participating in an electronic survey on artificial intelligence (AI)-narrated tobacco control messaging, United States, 2024.

Characteristic or outcome variables	Condition 1: told video is AI ($n=116$) ^a	Condition 2: told video is human ($n=138$)	Condition 3: told video is unknown ($n=121$)	Condition 4: no additional information ($n=125$)	Difference by condition, <i>P</i> value
Age (years), mean (SD)	47.93 (17.24)	49.92 (18.91)	49.02 (17.7)	46.65 (17.67)	.48
Sex, <i>n</i> (%)					.50
Male	63 (54.3)	63 (45.7)	61 (50.4)	58 (46.4)	
Female	52 (44.8)	74 (53.6)	58 (47.9)	66 (52.8)	
Income, <i>n</i> (%)					.44
<US \$25,000/year	19 (16.7)	30 (21.9)	29 (24.2)	30 (24.6)	
$\geq$ US \$25,000/year	95 (83.3)	107 (78.1)	91 (75.8)	92 (75.4)	
Race, <i>n</i> (%)					
White	91 (78.4)	106 (76.8)	97 (80.2)	91 (72.8)	.56
Black	15 (12.9)	17 (12.3)	16 (13.2)	16 (12.8)	.10

Characteristic or outcome variables	Condition 1: told video is AI (n=116) ^a	Condition 2: told video is human (n=138)	Condition 3: told video is unknown (n=121)	Condition 4: no additional information (n=125)	Difference by condition, <i>P</i> value
American Indian or Alaska Native	2 (1.7)	2 (1.4)	1 (0.8)	3 (2.4)	.80
Asian	6 (5.2)	10 (7.2)	4 (3.3)	9 (7.2)	.49
Pacific Islander	1 (0.9)	0 (0)	1 (0.8)	0 (0)	.53
Other	4 (3.4)	6 (4.3)	3 (2.5)	9 (7.2)	.30
Latinx	12 (10.3)	28 (20.3)	22 (18.2)	28 (22.4)	.08
Education, n (%)					.57
High school or below	20 (17.2)	29 (21.2)	29 (24.0)	23 (18.4)	
Above high school	96 (82.8)	108 (78.8)	92 (76.0)	102 (81.6)	
Marital status, n (%)					.79
Married	48 (41.4)	56 (40.6)	56 (46.3)	55 (44.0)	
Other	68 (58.6)	82 (59.4)	65 (53.7)	70 (56.0)	
Flavored product use among ever-users, n (%)					.99
Everyday or someday	16 (45.7)	20 (43.5)	15 (46.9)	16 (45.7)	
Not at all or do not know	19 (54.3)	26 (56.5)	17 (53.1)	19 (54.3)	
Lifetime cigarette use, n (%)					.14
Yes	33 (28.9)	56 (42.4)	46 (38.0)	40 (33.1)	
No	81 (71.1)	76 (57.6)	75 (62.0)	81 (66.9)	
Video rating, mean (SD)	3.22 (1.62)	2.73 (1.53)	3.08 (1.64)	2.82 (1.68)	.07
Perceived effectiveness, mean (SD)	3.51 (1.02)	3.75 (1.01)	3.71 (0.98)	3.79 (0.10)	.14
Message credibility, mean (SD)	4.03 (0.93)	4.20 (0.81)	4.07 (0.87)	4.19 (0.86)	.29
Speaker rating, mean (SD)	3.25 (0.93)	3.65 (0.91)	3.35 (0.99)	3.60 (0.96)	.001 ^b
Positive AI attitudes, mean (SD)	2.84 (0.88)	2.93 (0.88)	2.82 (0.79)	3.04 (0.90)	.17
Negative AI attitudes, mean (SD)	2.44 (0.87)	2.49 (0.88)	2.44 (0.86)	2.67 (0.90)	.15

^aAI: artificial intelligence.

^bStatistically significant.

Message Outcomes

Speaker ratings showed a significant difference across conditions ($P=.001$), with the highest ratings observed in Condition 2 (told the narrator was human; mean 3.65, SD 0.91). Condition 2 (told the narrator human; mean 3.65, SD 0.91) had significantly higher speaker rating compared to Condition 1 (told the narrator was AI; mean 3.25, SD 0.93). Condition 4 (no additional information provided; mean 3.60, SD 0.96) also had a significantly higher speaker rating compared to Condition 1 (told the narrator was AI; mean 3.25, SD 0.93). In contrast, there were no significant differences across conditions in video rating, perceived effectiveness, or perceived credibility. While not statistically significant, the largest difference in video rating was between Condition 1 (told the narrator was AI; mean 3.22, SD 1.62) and Condition 2 (told the narrator was human; mean 2.73, SD 1.53). Condition 1 also showed the lowest perceived effectiveness (mean 3.51, SD 1.02). Interestingly, positive AI attitudes were highest in Condition 4 (no information provided about the narrator; mean 3.04, SD 0.90), though this difference was also not statistically significant.

A linear regression analysis revealed that being told the narrator was human was significantly associated with higher speaker ratings compared to being told the narrator was AI ($B=0.69$; $P=.003$; $F_{13,125}=2.78$; $P=.002$). Similarly, positive AI attitudes were also associated with higher speaker ratings ($B=0.40$; $P<.001$). For perceived effectiveness, there were no significant differences between conditions. However, participants with more positive attitudes toward AI were more likely to rate the message as effective ($B=0.43$; $P<.001$). Regression models for video rating and message credibility outcomes were not significant.

Qualitative Comments

The results from the qualitative comment analysis produced 4 themes, including Vaping Outcomes, Video Negative Content, Voice, and Video Positive Content. The themes, theme definitions, corresponding codes, code frequencies within the dialog, and examples of corresponding quotes are outlined in Table 2. There were some notable differences and similarities in the qualitative comments across conditions.

Table 2. Thematic analysis of participant responses to a tobacco control video narrated by artificial intelligence (AI): summary of themes, definitions, and example codes from a randomized electronic survey of US adults, 2024.

Themes and their codes	Representative quote
Vaping Outcome	
Participants described their beliefs or attitudes toward tobacco products and vaping in general. This includes their views on the addictive qualities, financial costs, or health implications of these substances, as well as general criticisms and negative perceptions.	
Health	"Vaping ruins your health"
Cost	"...vaping is expensive."
Addiction	"...the thought was the importance of someone choosing a drug over enjoying a good time is a sad case of addiction"
Vaping negative	"People should listen, vaping is dangerous and addictive, don't believe commercials about vaping not being addictive."
Video Negative Content	
Participants provided negative feedback on the video related to its credibility, engagement, persuasiveness, or personal relevance. This feedback also involved comments about possible improvements to the video and what the video lacked.	
No factual	"The ad did not use fact or science to back up its point."
Boring	"It didn't apply to me or anything I would find useful and was actually pretty boring."
Disagreement	"It could help someone who is thinking of quitting... but it also just makes me feel more determined that having a ban on vape flavors is not the answer..."
Not convincing	"I thought it was sounding more like a novel, a fictitious one at that. I also thought that if the advertisement was trying to convince someone not to vape, then they failed because there was no sign of it having anybody's best interests in mind. It sounds like they were reading a story."
Not applicable	"It was targeted at adults. This is more of a youth or teenager issue in my opinion."
Voice	
Participants mentioned the narrator's voice being AI ^a or sounding monotone.	
AI detected	"... you could tell it was AI spoken"
Monotone	"The voice is monotone and dead. Creepy."
Video Content Positive	
Participants described the content of the video in a positive way. This included how well the ad conveys its message, connects with viewers on a personal level, and motivates action or awareness.	
Factual	"...gives facts"
Take action	"People need to be more informed about the dangers of vaping."
Compassion	"...this person lost everything, spent thousands, and damaged her lungs."
Personal	"It was very powerful. My husband died as a result of his nicotine addiction. This message hits home."
Convincing	"The ad was convincing to me, explaining the results of vaping addiction."
Genuine	"...is informative, genuine, vulnerable, and authentic"
Informative	"This video is very Informative, true, and inspiring."
Positive	"It was empowering."

^aAI: artificial intelligence.

For Condition 1 (participants told the video was AI), the most commonly referenced code was "informative" (n=21). This was closely followed by the perception of the theme "voice" being "monotoned," referenced 20 times, and the theme "vaping outcomes" coded as "addictive," referenced 19 times within Condition 1's dialog.

Participants who were told the narrator in the video was human (Condition 2) most commonly referenced with the code "vaping negative" within the theme "vaping outcome" (n=36). The second most referenced code within the condition was "positive" within the theme "video content," referenced 29 times.

In Condition 3 (participants told it was unknown whether the narrator was human or AI), the most recurring codes were

within the theme "video content," coded as "informative" and "positive," and the theme "vaping outcomes," which produced the code "negative outcomes," both referenced 25 times in the comments. This was followed closely by the theme "video content" being perceived as "convincing" and "positive," occurring 22 times.

Similar to Condition 2, Condition 4 (participants were given no information about the narrator) most commonly included the code "vaping negative" within the theme "vaping outcome," referenced 46 times within that condition. The second most recurring code in Condition 4 was "positive" within the theme "video content," referenced 36 times. This result was similar to Condition 2, where the highest recurring codes were "Theme: vaping outcome, code: vaping negative" and "Theme: video content, code: positive."

Across the conditions, the most frequent was “vaping negative” within the theme “vaping outcome,” referenced 12 times in Condition 1, 36 times in Condition 2, 25 times in Condition 3, and 46 times in Condition 4. Another common code across conditions was “informative” within the theme “video content positive” (Table 3).

Table 3. Frequency of qualitative themes coded from open-ended responses by experimental condition (narrator artificial intelligence [AI] identity disclosure).

Themes and their codes	Condition 1: told video is AI ^a (n=116)	Condition 2: told video is human (n=138)	Condition 3: told video is unknown (n=121)	Condition 4: no additional information (n=125)
Video Content Positive, n (%)				
Informative	21 (18)	21 (15)	25 (20)	21 (17)
Compassion	15 (13)	12 (8)	9 (7)	8 (6)
Positive	15 (13)	29 (21)	22 (18)	33 (26)
Personal	14 (12)	17 (12)	15 (12)	9 (7)
Genuine	10 (8)	14 (10)	22 (18)	18 (14)
Convincing	9 (7)	14 (10)	22 (18)	11 (9)
Take action	8 (6)	7 (5)	9 (7)	8 (6)
Factual	6 (5)	10 (7)	8 (6)	5 (4)
Voice, n (%)				
Monotone	20 (17)	11 (8)	17 (14)	14 (11)
AI detected	14 (12)	4 (3)	7 (5)	3 (2)
Real voice	4 (3)	1 (0.7)	3 (2)	4 (3)
Vaping outcome, n (%)				
Addictive	19 (16)	20 (14)	16 (13)	29 (23)
Negative	12 (10)	36 (26)	25 (20)	46 (37)
Cost	10 (8)	13 (9)	13 (10)	17 (13)
Health	7 (6)	8 (5)	4 (3)	8 (6)
Video content negative, n (%)				
Not applicable	9 (7)	12 (8)	11 (9)	6 (4)
Not convincing	9 (7)	5 (3)	11 (9)	7 (5)
Not factual	9 (7)	6 (4)	5 (4)	7 (5)
Boring	5 (4)	6 (4)	7 (5)	9 (7)
Disagreement	5 (4)	6 (4)	8 (6)	7 (5)

^aAI: artificial intelligence.

Discussion

Principal Findings

This study explored how perceptions of narrator identity (AI vs human) influence audience evaluations of a health messaging video focused on a tobacco control policy. It was hypothesized that participants informed that the narrator was AI-generated would rate the video lower on video rating, perceived effectiveness, credibility, and speaker ratings compared to those told the narrator was human, unknown, or not mentioned. Importantly, across all conditions, the narration was AI-generated, even when participants were told that it was human. In addition, the study found similarities and differences in qualitative feedback that differed by condition.

The findings partially supported the hypothesis. Quantitative results revealed that speaker ratings were significantly higher when participants were told the narrator was human,

suggesting that perceived narrator identity may change perceptions toward the narrator, even when the narration itself is identical. The significant differences in speaker ratings, despite the identical message, suggest that source cues like whether a narrator is perceived as AI may shape how individuals evaluate a speaker. These findings indicate that perceptions of the narrator can influence speaker evaluations, though the quality of the content itself may help mitigate some of these biases. Trends showing slightly lower ratings for videos where the narrator was described as AI further align with existing research that highlights biases against AI-driven communication in trust-sensitive contexts, such as health messaging [37].

Qualitative findings added depth to the results of this study by illustrating how perceptions of narrator identity shaped audience engagement. Participants in the AI-narrator condition (Condition 1) frequently described the voice as “monotone” or “unnatural,” potentially detracting from the video’s emotional resonance. In contrast, participants in

the human-narrator condition (Condition 2) described the same content as “genuine,” “convincing,” and “relatable.” This suggests that perceived authenticity and relatability are critical to audience engagement. Interestingly, participants in the no-information condition (Condition 4) focused primarily on the video’s content, referencing its informative nature and the health impacts of vaping. Across all conditions, participants consistently highlighted the video’s “informative” content, underscoring the importance of factual accuracy and relevance in health messaging.

In addition, participants’ overall attitudes toward AI did not show any significant difference among the conditions. This may suggest that overall perceptions of AI may be less influential than specific perceptions toward the narrator of a message. This study offers a novel contribution by examining how knowledge of an AI narrator may influence perceptions of a tobacco control policy message, using a controlled design where the message content remains constant across conditions. Unlike previous research that analyzed existing human-generated content (ie, Twitter posts), this design isolates the effect of the perception of the speaker by removing confounding factors (eg, message length and quality of the message) [19]. In addition, the inclusion of qualitative data, which is currently scarce in AI message research, provides deeper insight into audience reactions to AI-generated speakers. Taken together, these findings suggest that public health professionals should prioritize shaping perceptions of the speaker as a central factor in developing effective AI-based health communication strategies, rather than focusing solely on broad attitudes toward AI.

Limitations

There are several limitations to note that should be addressed in future research. Participants were only exposed to one video, which may not have provided sufficient opportunity to fully process and evaluate the content. In addition, while participants were told the narrator was AI or human, all conditions featured the same AI-generated voice. This approach controlled for variability but did not allow for direct comparison with an actual human-narrated video. Future studies should include both human and AI narrators to better understand audience preferences and perceptions.

Comparison With Prior Work

Other research has shown that negative attitudes toward AI may influence message perceptions. In one study comparing the evaluation of AI versus human-generated messages, researchers found moderating effects of negative attitudes toward AI on the evaluation of the messages, but not on which message was ultimately selected [20]. Given the limitations of this study and existing research in terms of study design, exposure, and lack of measurement over time, more research is needed to better understand the role of general AI attitudes in message outcomes.

While not significant, there was a trend in the overall video rating that was similar. With the current findings, there is some indication of preferences toward using human-generated videos compared to AI. Other studies have found that

there may be mistrust and suspicion regarding AI involvement in health content, as it is a newer generative digital media tool [37]. As AI continues to evolve and update, AI-related concerns persist and could contribute to the study’s result of lower ratings among the videos where participants were told there was an AI narrator. In a recent study comparing human-made and AI-generated teaching videos, researchers found similar results of participant preferences for human-made videos. However, results between conditions demonstrated participants acquired the same level of knowledge whether the video was human-made or not [38]. Another study examining the effectiveness of health messages, specifically related to COVID-19 vaccination, between human and nonhuman sources found participants perceived the message as more credible and more influential if the message came from a human in comparison to an AI communicator [39]. These findings are similar to the results of this study, finding a preference for human interaction over AI, but demonstrate the potential of AI-generated materials as an adequate means of effectively portraying information for an intended audience.

In the context of tobacco control policies, this research provides some evidence that generative AI could be used to develop messages to promote policy support. Although differences were observed in speaker ratings, other outcomes, such as video likability and perceived effectiveness, were not significant. Previous work on tobacco prevention campaigns has emphasized the importance of tailoring narratives to resonate with specific communities [7,8]. Generative AI offers potential for such tailoring, particularly for communities disproportionately affected by flavored tobacco [16]. In addition, it has the potential to be used to counteract misinformation and pro-tobacco content on social media [40]. In a recent study, researchers compared generative AI messages to existing organizational messages for vaping prevention [41]. Contrary to our current findings, AI-generated messages outperformed the existing messages. It is important to note, however, that this previous work compared different messages, while this study used the same video across conditions.

While this study did not find differences in perceived credibility among conditions, researchers and practitioners should remain attentive to the accuracy of generative AI output [42]. For example, one content analysis found that ChatGPT-generated content for smoking cessation was accurate only 57% of the time [43]. Given the lack of published research studying the effectiveness and cost-effectiveness of using generative AI for promoting tobacco control policies, more research is needed to understand how this technology can be leveraged most effectively by public health professionals.

Implications for Practice

The findings highlight important considerations for public health practitioners seeking to integrate AI into health messaging campaigns. One key takeaway is the significant influence of perceived narrator identity on speaker ratings. Participants rated the narrator higher when they believed it

to be human, even though the narration was AI-generated in all conditions. This underscores the need to address biases against AI in contexts where trust and emotional engagement are critical. However, it was noted that general AI attitudes were not different across the conditions. Strategies such as framing AI as a collaborative tool or emphasizing human oversight in content creation may help mitigate these biases. For example, campaigns could position AI as enhancing efficiency and personalization while maintaining human involvement to reassure audiences.

The qualitative results emphasize the importance of emotionally engaging and factually accurate content. Participants across conditions frequently described the video as “informative,” indicating that high-quality messaging can resonate with audiences, regardless of the narrator’s perceived identity. This suggests that the strength of the content itself can help overcome biases against AI-generated communication. Public health practitioners should prioritize creating clear, relevant, and accurate messages that address audience needs and expectations.

Finally, findings from the no-information condition (Condition 4) suggest that omitting explicit mention of AI involvement can help audiences engage more neutrally with the content. When ethical considerations allow, avoiding disclosure of the narrator’s source may reduce bias and encourage audiences to focus on the message itself. However, transparency remains essential in some contexts, and practitioners must carefully balance trust-building with strategies to mitigate audience biases.

Overall, these insights highlight the potential for AI to serve as a valuable tool in health messaging. Its cost-effectiveness, scalability, and ability to tailor content for diverse populations may make it a particularly useful resource for organizations with limited budgets or technical expertise [38]. However, more research is needed to fully understand the feasibility and cost-effectiveness of using AI within public health organizations. By addressing biases and emphasizing content quality, practitioners can leverage AI to deliver impactful, audience-centered health messages.

Implications for Theory

This study contributes to the theoretical understanding of audience responses to AI in health communication. The significant preference for human narrators in speaker ratings reinforces existing theories that emphasize the importance of trust and relatability in communication [39]. These findings align with the Social Presence Theory [44], which suggests that perceptions of the communicator’s presence—whether real or artificial—can significantly impact how messages are received. In this study, even though the content and delivery were identical, participants’ evaluations were shaped by their perceptions of narrator identity, demonstrating that trust and relatability remain pivotal in audience engagement. The results also somewhat align with ELM. It appeared that the “cue” of knowing whether the speaker was AI or not influenced how participants perceived the speaker and processed the information in terms of intention to follow the

recommendations of the video. Interestingly, using the CASA paradigm, knowledge of AI did not appear to impact trust or perceived credibility in the video.

The qualitative findings reveal how perceptions, rather than actual differences in narration, drive audience engagement. Participants who believed the narrator was human described the content as “genuine” and emotionally resonant, while those who told the narrator was AI frequently noted a lack of authenticity. This underscores the critical role of perceived authenticity in shaping trust and engagement, which is supported by the Heuristic-Systematic Model [45]. According to this model, individuals often rely on heuristics, such as perceived human authenticity, to evaluate messages when systematic processing is not fully engaged. The findings reinforce the importance of narrative framing in health communication, as even subtle cues about the narrator’s identity can influence audience perceptions and ultimately the effectiveness of the message.

Interestingly, the lack of significant differences in video rating, perceived effectiveness, and credibility across conditions suggests that high-quality content can mitigate some biases against AI. This finding underscores the importance of focusing on the intrinsic value of the message itself—clarity, relevance, and informativeness—over the medium or delivery method. Theoretical models, such as the ELM [23], emphasize that when the content of a message is highly relevant and cognitively engaging, audiences are more likely to process the message through the central route, leading to deeper understanding and less susceptibility to peripheral biases, such as those linked to the narrator’s identity. Similarly, research on health communication has shown that message quality and perceived relevance are among the strongest predictors of message effectiveness, even in digital or mediated contexts [5]. These findings suggest that public health campaigns leveraging AI should prioritize delivering clear and relevant information to ensure audience engagement and mitigate potential biases.

Future Research

This study fills a gap in research examining reactions to AI-involved health messaging videos. Future research could build on these findings to examine whether message outcomes differ over time with multiple exposures to different types of videos. In addition, studies could compare different AI visuals, such as avatars, graphics, or cartoons, to determine likability across AI technologies. Examining emotional reactions or features of AI-generated emotions is also needed. Finally, examining the usability and cost-effectiveness of AI within public health settings would better inform the overall utility of using it for health messaging. This study did not examine the cost-effectiveness; however, some research shows that AI can be a cost-effective and rapid means of producing content, which is a common reason more industries use it in outreach materials [37]. It would be important to determine the cost-effectiveness and usability of AI within public health settings where resources and training may be limited. Overall, AI is a tool with great potential in future health messaging and outreach, but obstacles remain

related to the intended population's attitude, trust, and future acceptance.

Conclusions

As AI continues to evolve and programs become more human-like, more organizations are adapting to use AI. This study found that speaker ratings were significantly influenced by the perceived identity of the narrator, with participants

rating the speaker significantly higher when they believed the voice was human. However, no significant differences were observed for perceived effectiveness, credibility, or overall video ratings across conditions. These findings suggest that message sources, particularly regarding AI, can influence an audience's perceptions of the speaker even when content remains constant.

Acknowledgments

We are grateful to Brielle Vail for her assistance with analyzing the qualitative data. We used the generative AI tool, PlayHT, to record the script for the video used in the study. AI was not used in the writing of this manuscript. This research was generously supported by the Bill and Linda Frost Foundation at the California Polytechnic State University, San Luis Obispo, Bailey College of Science and Mathematics.

Data Availability

The datasets analyzed during this study are available from the corresponding author on reasonable request.

Authors' Contributions

JMA, DA, and AD conceptualized the study, and JMA acquired funding for the study. LS created the video and survey used in the study. SR conducted the formal quantitative analysis and prepared the tables for the quantitative data with support from JMA. KJS conducted the qualitative analysis and prepared the themes table with support from JMA. LS and SA drafted the initial manuscript, and all authors contributed to editing and reviewing the manuscript before submission.

Conflicts of Interest

None declared.

Multimedia Appendix 1

Link to the video shown to participants.

[\[DOCX File \(Microsoft Word File\), 29 KB-Multimedia Appendix 1\]](#)

References

1. BRFSS questionnaires. Centers for Disease Control and Prevention. Sep 4, 2024. URL: <https://www.cdc.gov/brfss/questionnaires/index.htm> [Accessed 2024-10-04]
2. Eliminating tobacco-related disease and death: addressing disparities. US Department of Health and Human Services. Nov 19, 2024. URL: <https://www.hhs.gov/sites/default/files/2024-sgr-tobacco-related-health-disparities-exec-summary.pdf> [Accessed 2024-11-19]
3. LGBTQ+ people and commercial tobacco: health disparities and ways to advance health equity. US Centers for Disease Control and prevention. 2024. URL: <https://www.cdc.gov/tobacco-health-equity/collection/lgbtq.html> [Accessed 2025-10-09]
4. Health problems caused by secondhand smoke. US Centers for Disease Control and Prevention. Jun 17, 2024. URL: <https://www.cdc.gov/tobacco/secondhand-smoke/health.html> [Accessed 2024-12-04]
5. Hoffman SJ, Tan C. Overview of systematic reviews on the health-related effects of government tobacco control policies. BMC Public Health. Aug 5, 2015;15:744. [doi: [10.1186/s12889-015-2041-6](https://doi.org/10.1186/s12889-015-2041-6)] [Medline: [26242915](https://pubmed.ncbi.nlm.nih.gov/26242915/)]
6. Public Health Law Center at Mitchell Hamline School of Law. Endgame policy platform. 2023. URL: <https://www.publichealthlawcenter.org/sites/default/files/inline-files/Endgame%20Policy%20Platform%20-%20November%202023.pdf> [Accessed 2024-12-04]
7. Payán DD, Burke NJ, Persinger J, Martinez J, Jones Barker L, Song AV. Public support for policies to regulate flavoured tobacco and e-cigarette products in rural California. Tob Control. Apr 2023;32(e1):e125-e129. [doi: [10.1136/tobaccocontrol-2021-057031](https://doi.org/10.1136/tobaccocontrol-2021-057031)] [Medline: [35064014](https://pubmed.ncbi.nlm.nih.gov/35064014/)]
8. Edwards R, Wilson N, Peace J, Weerasekera D, Thomson GW, Gifford H. Support for a tobacco endgame and increased regulation of the tobacco industry among New Zealand smokers: results from a National Survey. Tob Control. May 2013;22(e1):e86-93. [doi: [10.1136/tobaccocontrol-2011-050324](https://doi.org/10.1136/tobaccocontrol-2011-050324)] [Medline: [22535362](https://pubmed.ncbi.nlm.nih.gov/22535362/)]
9. González-Marrón A, Koprivnikar H, Tisza J, et al. Tobacco endgame in the WHO European Region: Feasibility in light of current tobacco control status. Tob Induc Dis. 2023;21:151. [doi: [10.18332/tid/174360](https://doi.org/10.18332/tid/174360)] [Medline: [38026503](https://pubmed.ncbi.nlm.nih.gov/38026503/)]
10. Kreslake JM, Wayne GF, Alpert HR, Koh HK, Connolly GN. Tobacco industry control of menthol in cigarettes and targeting of adolescents and young adults. Am J Public Health. Sep 2008;98(9):1685-1692. [doi: [10.2105/AJPH.2007.125542](https://doi.org/10.2105/AJPH.2007.125542)] [Medline: [18633084](https://pubmed.ncbi.nlm.nih.gov/18633084/)]

11. Gardiner PS. The African Americanization of menthol cigarette use in the United States. *Nicotine & Tobacco Res.* Feb 2004;6(1):55-65. [doi: [10.1080/14622200310001649478](https://doi.org/10.1080/14622200310001649478)]
12. Gately I. Tobacco: A Cultural History of How an Exotic Plant Seduced Civilization. Grove Press; 2001. URL: <https://books.google.com/books?id=x41jVocj05EC&lpg=PP1&pg=PP8#v=onepage&q&f=false> [Accessed 2025-10-19]
13. Mendez D, Le TTT. Consequences of a match made in hell: the harm caused by menthol smoking to the African American population over 1980-2018. *Tob Control.* Sep 16, 2021;16:tobaccocontrol-2021-056748. [doi: [10.1136/tobaccocontrol-2021-056748](https://doi.org/10.1136/tobaccocontrol-2021-056748)] [Medline: [34535507](https://pubmed.ncbi.nlm.nih.gov/34535507/)]
14. Hicks NJ, Nicols CM. Learning B, editor. Health Industry Communication. 2nd ed. Jones & Bartlet Learning; 2017. URL: <https://www.jblearning.com/catalog> [Accessed 2024-12-04]
15. Menichetti J, Hillen MA, Papageorgiou A, Pieterse AH. How can ChatGPT be used to support healthcare communication research? *Patient Educ Couns.* Oct 2023;115:107947. [doi: [10.1016/j.pec.2023.107947](https://doi.org/10.1016/j.pec.2023.107947)]
16. Miller MR, Sehat CM, Jennings R. Leveraging AI for public health communication: opportunities and risks. *J Public Health Manag Pract.* 2024;30(4):616-618. [doi: [10.1097/PHH.0000000000001986](https://doi.org/10.1097/PHH.0000000000001986)]
17. Alabduljabbar R. User-centric AI: evaluating the usability of generative AI applications through user reviews on app stores. *PeerJ Comput Sci.* 2024;10:e2421. [doi: [10.7717/peerj-cs.2421](https://doi.org/10.7717/peerj-cs.2421)] [Medline: [39650468](https://pubmed.ncbi.nlm.nih.gov/39650468/)]
18. Zhang W, Lu G, Chen Z, Li GY. Generative video communications: concepts, key technologies, and future research trends. *Engineering (Beijing).* Jun 2025. [doi: [10.1016/j.eng.2025.06.018](https://doi.org/10.1016/j.eng.2025.06.018)]
19. Lim S, Schmäzle R. Artificial intelligence for health message generation: an empirical study using a large language model (LLM) and prompt engineering. *Front Commun.* 2023;8. [doi: [10.3389/fcomm.2023.1129082](https://doi.org/10.3389/fcomm.2023.1129082)]
20. Teixeira da Silva JA. Can ChatGPT rescue or assist with language barriers in healthcare communication? *Patient Educ Couns.* Oct 2023;115:107940. [doi: [10.1016/j.pec.2023.107940](https://doi.org/10.1016/j.pec.2023.107940)]
21. Olawade DB, Aienobe-Asekharen CA. Artificial intelligence in tobacco control: a systematic scoping review of applications, challenges, and ethical implications. *Int J Med Inform.* Oct 2025;202:105987. [doi: [10.1016/j.ijmedinf.2025.105987](https://doi.org/10.1016/j.ijmedinf.2025.105987)] [Medline: [40409168](https://pubmed.ncbi.nlm.nih.gov/40409168/)]
22. Shin M, Kim SJ, Biocca F. The uncanny valley: no need for any further judgments when an avatar looks eerie. *Comput Human Behav.* May 2019;94:100-109. [doi: [10.1016/j.chb.2019.01.016](https://doi.org/10.1016/j.chb.2019.01.016)]
23. PettyRE, CacioppoJT. The elaboration likelihood model of persuasion. In: *Adv Exp Soc Psychol.* Vol 19. Academic Press; 1986:123-205. [doi: [10.1016/S0065-2601\(08\)60214-2](https://doi.org/10.1016/S0065-2601(08)60214-2)]
24. Reeves B, Nass C. The Media Equation: How People Treat Computers, Television, and New Media like Real People and Places. Cambridge University Press; 1996. URL: <https://psycnet.apa.org/record/1996-98923-000> [Accessed 2025-10-18]
25. Zhao X, Sun Y, Liu W, Wong CW. Tailoring generative AI chatbots for multiethnic communities in disaster preparedness communication: extending the CASA paradigm. *J Comput-Mediat Commun.* Jan 24, 2025;30(1). [doi: [10.1093/jcmc/zmae022](https://doi.org/10.1093/jcmc/zmae022)]
26. U.S. census bureau releases 2018-2022 ACS 5-year estimates. United States Census Bureau. 2023. URL: <https://www.census.gov/programs-surveys/acs/news/updates/2023.html> [Accessed 2025-08-28]
27. Pogue K, Jensen JL, Stancil CK, et al. Influences on attitudes regarding potential COVID-19 vaccination in the United States. *Vaccines (Basel).* Oct 3, 2020;8(4):582. [doi: [10.3390/vaccines8040582](https://doi.org/10.3390/vaccines8040582)] [Medline: [33022917](https://pubmed.ncbi.nlm.nih.gov/33022917/)]
28. Alber JM, Glanz K. Does the screening status of message characters affect message effects? *Health Educ Behav.* Feb 2018;45(1):14-19. [doi: [10.1177/1090198117708232](https://doi.org/10.1177/1090198117708232)] [Medline: [28548601](https://pubmed.ncbi.nlm.nih.gov/28548601/)]
29. Alber JM, Cohen C, Bleakley A, et al. Comparing the effects of different story types and speakers in hepatitis B storytelling videos. *Health Promot Pract.* Sep 2020;21(5):811-821. [doi: [10.1177/1524839919894248](https://doi.org/10.1177/1524839919894248)] [Medline: [31955614](https://pubmed.ncbi.nlm.nih.gov/31955614/)]
30. Schepman A, Rodway P. The General Attitudes towards Artificial Intelligence Scale (GAAIS): confirmatory validation and associations with personality, corporate distrust, and general trust. *International Journal of Human-Computer Interaction.* Aug 9, 2023;39(13):2724-2741. [doi: [10.1080/10447318.2022.2085400](https://doi.org/10.1080/10447318.2022.2085400)]
31. Kang Y, Cappella JN, Fishbein M. The effect of marijuana scenes in anti-marijuana public service announcements on adolescents' evaluation of ad effectiveness. *Health Commun.* Sep 2009;24(6):483-493. [doi: [10.1080/10410230903104269](https://doi.org/10.1080/10410230903104269)] [Medline: [19735026](https://pubmed.ncbi.nlm.nih.gov/19735026/)]
32. Green MC, Donahue JK. Persistence of belief change in the face of deception: the effect of factual stories revealed to be false. *Media Psychol.* Jan 2011;14(3):312-331. [doi: [10.1080/15213269.2011.598050](https://doi.org/10.1080/15213269.2011.598050)]
33. Eastin MS. Credibility assessments of online health information: the effects of source expertise and knowledge of content. *J Comput-Mediat Commun.* Jul 1, 2001;6(4):JCMC643. [doi: [10.1111/j.1083-6101.2001.tb00126.x](https://doi.org/10.1111/j.1083-6101.2001.tb00126.x)]
34. McCroskey JC, Teven JJ. Goodwill: a reexamination of the construct and its measurement. *Commun Monogr.* Mar 1999;66(1):90-103. [doi: [10.1080/03637759909376464](https://doi.org/10.1080/03637759909376464)]

35. Jones LW, Sinclair RC, Courneya KS. The effects of source credibility and message framing on exercise intentions, behaviors, and attitudes: an integration of the elaboration likelihood model and prospect theory1. *J Applied Social Psychol.* Jan 2003;33(1):179-196. URL: <https://onlinelibrary.wiley.com/toc/15591816/33/1> [Accessed 2025-10-22] [doi: [10.1111/j.1559-1816.2003.tb02078.x](https://doi.org/10.1111/j.1559-1816.2003.tb02078.x)]
36. Braun V, Clarke V. Using thematic analysis in psychology. *Qual Res Psychol.* Jan 2006;3(2):77-101. [doi: [10.1191/1478088706qp063oa](https://doi.org/10.1191/1478088706qp063oa)]
37. Williamson SM, Prybutok V. The era of artificial intelligence deception: unraveling the complexities of false realities and emerging threats of misinformation. *Information.* 2024;15(6):299. [doi: [10.3390/info15060299](https://doi.org/10.3390/info15060299)]
38. Wang Y, Pan Y, Yan M, Su Z, Luan TH. A survey on ChatGPT: AI-generated contents, challenges, and solutions. *IEEE Open J Comput Soc.* 2023;4:280-302. [doi: [10.1109/OJCS.2023.3300321](https://doi.org/10.1109/OJCS.2023.3300321)]
39. Dai Y, Lee J, Kim JW. AI vs. human voices: how delivery source and narrative format influence the effectiveness of Persuasion messages. *Int J Hum Comput Int.* Dec 16, 2024;40(24):8735-8749. [doi: [10.1080/10447318.2023.2288734](https://doi.org/10.1080/10447318.2023.2288734)]
40. Kong G, Ouellette RR, Murthy D. Generative artificial intelligence and social media: insights for tobacco control. *Tob Control.* Dec 6, 2024;6:2024. [doi: [10.1136/tc-2024-058813](https://doi.org/10.1136/tc-2024-058813)] [Medline: [39643443](https://pubmed.ncbi.nlm.nih.gov/39643443/)]
41. Leung J, Sun T, Stjepanovic D, et al. Generative artificial intelligence with youth codesign to create vaping awareness advertisements. *JAMA Netw Open.* Jul 1, 2025;8(7):e2514040. [doi: [10.1001/jamanetworkopen.2025.14040](https://doi.org/10.1001/jamanetworkopen.2025.14040)] [Medline: [40742592](https://pubmed.ncbi.nlm.nih.gov/40742592/)]
42. Andrew A. Overview of the emerging role of chatbots, including large language models, in supporting tobacco smoking and vaping cessation: a narrative review. *Global Health Journal.* Mar 2025;9(1):6-11. [doi: [10.1016/j.glohj.2025.02.003](https://doi.org/10.1016/j.glohj.2025.02.003)]
43. Abroms LC, Yousefi A, Wysota CN, Wu TC, Broniatowski DA. Assessing the adherence of ChatGPT chatbots to public health guidelines for smoking cessation: content analysis. *J Med Internet Res.* Jan 30, 2025;27:e66896. [doi: [10.2196/66896](https://doi.org/10.2196/66896)] [Medline: [39883917](https://pubmed.ncbi.nlm.nih.gov/39883917/)]
44. Short J, Williams E, Christie B. *The Social Psychology of Telecommunications.* John Wiley & Sons; 1976. URL: <https://search.worldcat.org/title/2585964> [Accessed 2024-12-04]
45. Chaiken S, Ledgerwood A. Heuristic and systematic information processing. In: *Handbook of Theories of Social Psychology.* Vol . 2012;1. 246-266. [doi: [10.4135/9781446249215.n13](https://doi.org/10.4135/9781446249215.n13)]

Abbreviations

AI: artificial intelligence

CASA: computers as social actors

ELM: elaboration likelihood model

Edited by Andrew Coristine; peer-reviewed by Kwanho Kim, Yuxian Cui; submitted 22.Feb.2025; final revised version received 25.Sep.2025; accepted 25.Sep.2025; published 27.Oct.2025

Please cite as:

Alber JM, Askay D, Dhillon A, Sandoval L, Ramos S, Santilena K

AI Awareness and Tobacco Policy Messaging Among US Adults: Electronic Experimental Study

*JMIR AI*2025;4:e72987

URL: <https://ai.jmir.org/2025/1/e72987>

doi: [10.2196/72987](https://doi.org/10.2196/72987)

© Julia Mary Alber, David Askay, Anuraj Dhillon, Lauren Sandoval, Sofia Ramos, Katharine Santilena. Originally published in JMIR AI (<https://ai.jmir.org>), 27.Oct.2025. This is an open-access article distributed under the terms of the Creative Commons Attribution License (<https://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution, and reproduction in any medium, provided the original work, first published in JMIR AI, is properly cited. The complete bibliographic information, a link to the original publication on <https://www.ai.jmir.org/>, as well as this copyright and license information must be included.