

Impact of AI-Generated Images on Social Media Trust Across Generations

Tarik Doğan / Ph.D. 

Selçuk University, Faculty of Communication, Department of Advertising
tdogan@selcuk.edu.tr

Abstract

This study investigates the impact of AI-generated images circulating on social media on users' perceived trust in visual content, with a focus on generational differences. The rapid development of generative artificial intelligence technologies and the proliferation of hyper-realistic synthetic visuals have intensified debates surrounding reality, credibility, and verification in digital media environments. The study is theoretically grounded in Baudrillard's theory of hyperreality, Manovich's approach to digital media, and debates on algorithmic authority. A quantitative method and quasi-experimental design were employed. Real photographs and AI-generated visuals were presented through an asymmetric block design (2-2-1-2-2-1) developed to minimize learning effects. The final analytic sample consisted of 399 voluntary online participants classified into Generations X, Y, and Z, with responses screened using an attention-check criterion. Data were analyzed using paired-samples t-tests, one-way ANOVA, and post-hoc analyses. Findings showed

that participants rated AI-generated visuals as significantly more trustworthy than real photographs. The highest individual trust score was recorded for an AI-generated landscape image, while the strongest pairwise effect was observed in the portrait category. Significant generational differences were also identified: Generation X showed lower trust in AI-generated visuals, whereas Generations Y and Z expressed greater confidence in synthetic content. No statistically significant differences were found by gender or education level. Overall, the findings suggest a reconfiguration of perceived visual trust in this experimental context. However, the results are specific to the present sample and ten-image stimulus set and should not be interpreted as evidence of a society-wide shift in trust.

Keywords: AI-Generated Images, Hyperreality, Digital Trust, Social Media, Generational Differences

JEL Codes: L82, D83, D91, C83, L86

Citation: Doğan, T. (2026). Impact of AI-generated images on social media trust across generations. *Researches on Multidisciplinary Approaches (ROMAYA Journal)*, 2026(2).

1. Introduction

With the development of digital media technologies, visual production processes have undergone a radical transformation; in particular, Generative Artificial Intelligence (GenAI) systems have reopened debates on the concept of visual reality. The widespread use of artificial intelligence models capable of generating images and videos from text has sparked significant discussion about the accuracy, reliability, and potential for content manipulation circulating on social media platforms (Aziz et al., 2024). As AI-supported images become increasingly realistic, it becomes more difficult for users to distinguish between real photographs and synthetic content; consequently, the concept of “visual evidence” is being transformed within the digital media environment (Karakaya & Dayı, 2024). In this context, generative artificial intelligence has become not merely a technical innovation but also a sociocultural phenomenon that affects social trust, media literacy, and information verification processes. Generative adversarial networks (GANs), developed by Goodfellow et al. (2014), constitute a significant turning point in the development of generative artificial intelligence technologies. Through GAN- and diffusion-based models, people, objects, and events that have never existed in the physical world can be produced in hyper-realistic forms. Today, systems such as Midjourney, DALL-E, and Stable Diffusion transform digital visual culture by generating synthetic images with high aesthetic quality (Arielli, 2024). This transformation affects not only modes of visual production but also representational practices in fields such as art, media, advertising, and political communication (Manovich, 2001). Baudrillard’s theory of simulation offers an important theoretical framework for explaining the new regime of reality created by AI-generated images. According to Baudrillard, in contemporary society, simulation replaces reality, producing a new order called “hyperreality” (Baudrillard, 2018). Images generated by artificial intelligence also reinforce the order of simulacra by creating a sense of reality without necessarily referring to the physical world. Synthetic content circulating especially on social media platforms weakens users’ visual verification reflexes and reshapes trust in visual media (Gül Ünlü & Küçükşabanoğlu, 2023).

Experimental studies have shown that users often fail to distinguish AI-generated faces with a high level of photorealism from photographs of real people. This technical perfection turns synthetic images into not only an aesthetic form of production but also an instrument for manipulating trust (Nightingale & Farid, 2022). Similarly, Osińska et al. (2025) state that users have difficulty distinguishing certain AI-generated images from real photographs and those high-resolution human faces, in particular, generate a sense of trust. Göring et al. (2023) further demonstrate that the aesthetic appeal of AI-generated images direct-

ly affects users’ perception of reality. This indicates that the technical perfection of AI-generated content facilitates perceptual manipulation in social media environments. The interaction-oriented structure of social media algorithms increases the visibility of AI-generated content and deepens the risk of disinformation (Doğan, 2024). Castells (2008) argues that, in the network society, information now circulates through digital networks and that new power relations are established through algorithms. Pasquale (2015) states that algorithmic systems have become black boxes, and that users consume digital platforms without knowing how or why certain content is prioritized. This situation particularly facilitates the use of hyper-realistic synthetic images for manipulative purposes on social media. Fake content produced through deepfake technologies is increasingly an effective tool in propaganda, disinformation, and perception management (Zan & Zan, 2025).

Although studies focusing on technical realism, aesthetic qualities, and perceptual effects of AI-generated images have increased in the literature, empirical research examining how such content transforms user trust in social media across generations remains limited. In particular, the generational differences in trust perceptions toward synthetic content are noteworthy. While Prensky (2001) defines individuals born in the digital age as “digital natives,” Buckingham (2003) emphasizes that familiarity with digital technologies alone does not necessarily indicate critical media literacy. Generation Z, which has high levels of social media use, encounters AI-generated content more frequently and consumes it as a natural part of their everyday digital experience. Gürel’s (2023) study shows that young users are not fully certain about the accuracy of online content and experience concerns about information security in the digital media environment. Similarly, Hobbs (2017) notes that digital literacy requires not only technical skills of use, but also the capacity for content verification and critical evaluation. However, the effect of AI-generated content on user perception is not limited to trust alone. Some studies show that highly realistic images that do not appear entirely natural create an “uncanny valley” effect, generating feelings of discomfort and distrust in users (Kishnani, 2025). This demonstrates that AI-generated images are associated not only with technical quality but also with cognitive, psychological, and cultural processes of perception. Therefore, synthetic media content constitutes a multilayered field of digital experience that transforms users’ perception of reality.

This study aims to examine the effect of AI-generated images circulating in social media environments on user trust. Within the scope of the research, real photographs and AI-generated images are compared through a controlled experimental design, and users’ levels of trust are analyzed. In particular, through an asymmetric block structure in which ima-

ges are presented in specific sequences, the study seeks to measure participants' tendencies to distinguish synthetic content and their perceptions of trust. In this respect, the study aims to make an empirical contribution to the literature on visual reality, digital trust, and media literacy in the age of generative artificial intelligence.

2. Theoretical Framework

Images generated by artificial intelligence represent a logical evolution of digital media algorithms. According to Manovich (2024), the technical infrastructure of this process is based on the algorithmic decomposition of shape, light, and motion in the digital construction of objects, as well as on the computer-based simulation of traditional artistic effects through graphic software such as Photoshop in the 1990s. However, this process is not merely a technical advancement; it also constitutes a sociotechnical rupture in which decision-making mechanisms turn into a non-transparent "black box" (Pasquale, 2015). This structure of digital media, built through the interaction between databases and algorithms (Manovich, 2001), has today become a fundamental force shaping the everyday practices of "born-digital" generations who use technology as if it were their native language (Palfrey & Gasser, 2008). Today, this process of simulation and data processing has reached the stage of "cinematic algorithms," in which traditional concepts of creativity are surpassed. According to Hutson and Smith (2025), in this new stage, production has evolved not into a fully automated process but into forms of "hybrid creation," in which the artist becomes a curator who directs and refines machine outputs. Ultimately, this technological ecosystem points to a new phase of "artificial aesthetics," in which machines measure and rank aesthetic value, while human creativity is pushed toward adapting to machine judgments through "aesthetic alignment" (Arielli, 2024).

2.1. Generative Artificial Intelligence Technologies and the Proliferation of Visual Content Production

Generative artificial intelligence (GenAI) technologies are radically transforming visual production practices, particularly through diffusion models that enable text-to-image generation (Aziz et al., 2024). Generative adversarial networks (GANs), developed by Goodfellow et al. (2014), and the zero-shot learning capability introduced by Ramesh et al. (2021) constitute the technical foundation of this transformation. Advances in infrastructure are no longer limited to still images but extend to sophisticated architectures capable of generating realistic video scenes with high visual coherence from textual

prompts (Sanchez-Acedo et al., 2026). The capacity of these systems has reached such a massive scale that more than 15 billion synthetic images generated by artificial intelligence in 2023 alone have been found to surpass the volume of production achieved by traditional photography over 150 years (Aziz et al., 2024). These technological possibilities have created a field of "democratization" in production processes by enabling visual content production to spread to broader user groups without requiring technical knowledge or specialized expertise (Rapp et al., 2025; Yıldız, 2025). However, these automated processes sever the physical referential relationship between the digital image and the objective world, thereby producing an ontological ambiguity (Karakaya & Dayı, 2024).

The dimension of this transformation, described as "cultural automation," presents a new structure that standardizes aesthetic judgment and transfers human creativity onto an algorithmic plane (Manovich & Arielli, 2024). This situation undermines the traditional status of photography as "evidence of reality" and reinforces the era of simulation and "hyperreality," in which computer-based representations without a referent are perceived as more real than the original (Baudrillard, 2018; Sontag, 2008; İpek & İpek, 2025). In this universe of synthetic reality, where the boundaries between the virtual and physical worlds become blurred, the fact that highly realistic deepfake content and hyper-realistic synthetic images have become almost impossible to distinguish with the human eye lays the groundwork for an epistemological crisis at the societal level (Osińska et al., 2025; Sanchez-Acedo et al., 2026). Indeed, the approximately 900% annual increase in deepfake videos worldwide clearly reveals the manipulative power and mass impact of this technology (Zan & Zan, 2025). When this environment of perceptual risk is combined with the target-audience segmentation and microtargeting algorithms of digital platforms, it enables disinformation to spread much faster, more personalized, and more persuasively (Gül Ünlü & Küçükşabanoğlu, 2023). It has been observed that AI-generated depictions of events that never occurred, such as fabricated images of political leaders or synthetic explosion scenes, can receive millions of interactions in a short period and cause fluctuations in global markets or public confusion (Karakaya & Dayı, 2024). The ability of algorithms to control information flows and create echo chambers provides an environment in which false information can be reproduced without sufficient oversight (Arslan, 2026). As a result, visual disinformation and the accompanying crisis of trust are legitimized by the flawless technical quality of synthetic images and transformed into a systemic problem that threatens information security, individual privacy, and democratic foundations (Gül Ünlü & Küçükşabanoğlu, 2023; Zan & Zan, 2025).

2.2. Theories of Visual Perception and the Uncanny Valley: Ontological Ambiguity

As the level of human likeness in visual perception increases, the “uncanny valley” effect emerges when the resemblance of an artificial figure to a human being reaches an “almost perfect” threshold (Mori, 2012). This effect is triggered not only in physical robotic systems but also by the viewer’s subconscious recognition of subtle artifacts in AI-generated images with mid-range realism, such as asymmetrical facial features, inconsistent lighting, and unrealistic skin textures (Kishnani, 2025). Indeed, uncanniness is also defined as a profound sense of “estrangement” arising from the presence of an extremely foreign and irrational element within a familiar form (Rapp et al., 2025). Contemporary systems attempt to overcome this valley through opaque data architectures and algorithms that prioritize fictional aesthetics (Hutson & Smith, 2025; Pasquale, 2015).

The increase in photorealism erases the ontological distinction between real photographs and synthetic images, thereby neutralizing cognitive skepticism (Karakaya & Dayı, 2024; Osińska et al., 2025). Indeed, empirical reception studies have shown that, beyond being indistinguishable from real faces, perfected synthetic faces are perceived as more authentic and more trustworthy than real human faces (Nightingale & Farid, 2022; Miller et al., 2023). The automation of “average aesthetics” by artificial intelligence dulls the human tendency to search for imperfections and manipulates perceptual defense mechanisms (Arielli, 2024). This is because the smooth images synthesized by generative models trained on massive datasets exploit existing expectations and biases in human cognition, thereby completely dissolving ontological boundaries (van Hees et al., 2025). Structural elements in the grammar of visual design, such as lighting and texture, construct a reality effect that carries the synthetic into the realm of truth (Kress & Van Leeuwen, 2020). This resulting hyper-realistic structure does not merely create an aesthetic illusion; it also gives rise to a complex crisis environment in which visual manipulation becomes legitimized on digital platforms, individuals are drawn into epistemological blindness, and social trust in the accuracy of information is undermined (Gül Ünlü & Küçükşabanoğlu, 2023; İpek & İpek, 2025).

2.3. Generational Theories and Digital Skepticism: The X, Y, and Z Paradox

Generation Z, referred to as digital natives, possesses a high level of operational self-confidence due to its “native language”-like relationship with technology (Prensky, 2001). However, technical proficiency does not always correspond to the level of critical literacy required to decode media messages (Buckingham, 2003). Digital literacy is not merely a technical competence but a critical evaluation process

that questions the ontological status of the image (Kellner & Share, 2007; Sezgin, 2025).

Generation Y, described in the literature as a bridge between Generation X and newer digital generations, stands out as a group born into an analog world, witnessing the digital revolution during adulthood, and positioned between these two worlds as “digital hybrids” (Yıldız, 2012, as cited in Canbolat, 2022). Generation Y displays a hybrid attitude by combining its ability to adapt to the technological speed of the succeeding Generation Z (Boztunç, 2014) with the traditional security and skepticism reflexes of the preceding Generation X (Canbolat, 2022). Although this generation integrates advanced digital tools into professional and everyday life at a high level, it tends to use the knowledge and experience it acquires in the most efficient way by subjecting them to a rational filtering process (Seymen, 2017).

Therefore, Generation Y’s approach to content in the digital ecosystem is characterized by cautious acceptance, balancing technological inclination and traditional skepticism regarding data privacy (Canbolat, 2022). By contrast, Generation X preserves its traditional information-seeking reflexes and “analog skepticism” (Prensky, 2001). This generation develops a structural skepticism toward digital content and subjects messages to a more rational filtering process (Hobbs, 2017). In particular, older age groups, especially those aged 60 and above, exhibit resistance to digital disinformation by verifying the information they encounter through strategies such as cross-checking multiple sources and consulting expert opinions (Arslan, 2026).

2.4. Algorithmic Authority and Trust: Collective Trust in Social Media Content

Algorithmic systems have become fundamental mechanisms of authority that structure individuals’ perception of reality and aesthetic judgments in the network society. This new social morphology redefines power relations and social hierarchies through algorithms that control the flow of information (Castells, 2008). In this context, algorithms function directly as objects of surveillance and domination from the perspective of panoptic power (Foucault, 1992), since what is considered “real,” “true,” or “visible” is no longer determined by central authorities, but by non-transparent data architectures described as “black boxes” (Pasquale, 2015). Faced with the complex and uncontrollable nature of these systems, users tend to exhibit a collective form of “algorithmic surrender” rather than developing rational skepticism, falling under the illusion that algorithmic outputs are objective, neutral, and error-free (Hashmi et al., 2024).

This process of surrender takes place within the ecosystem defined as “surveillance capitalism,” in which individual experiences are transformed into

behavioral data and commodified (Zuboff, 2019). Social media platforms reinforce a constructed and idealized perception of reality, referred to as “cinematic algorithms,” by prioritizing highly interactive content, including synthetic images (Hutson & Smith, 2025). In AI-generated content, the balance between “aesthetic quality” and “text-image alignment” disables the user’s cognitive filter (Aziz et al., 2024), while the visual fidelity created by synthetic images manipulates social trust by making the ontological distinction between them and real photographs nearly impossible (Osińska et al., 2025). This situation is characterized as a form of domination over the perception of reality, in which content-based and visual reading take precedence over technical reading (Karakaya & Dayı, 2024). Moreover, algorithms reproduce structural biases, thereby undermining the myth of systemic neutrality (Noble, 2018). Nevertheless, metrics such as likes, shares, and engagement on social networks construct an unshakable form of “social proof,” regardless of the accuracy of the content (Metzger & Flanagin, 2013). As a result, users tend to accept algorithmic filters as absolute truth rather than subjecting images to analytical scrutiny. Overcoming this epistemological crisis depends on establishing algorithmic accountability and independent auditing mechanisms (Diakopoulos, 2016).

3. Methodology

3.1. Research Design: Quantitative Analysis and Cross-Sectional Survey

This research employed a quasi-experimental design to examine causal relationships between variables using controlled stimuli and standardized measurement instruments. The study is based entirely on quantitative data collection and analysis and is structured within the cross-sectional survey framework, which aims to describe the attitudes, perceptions, and behavioral tendencies of large populations through numerical data (Creswell, 2014; Saunders et al., 2019). In accordance with the basic axioms of the quantitative research tradition, analyses of variance (ANOVA) and post-hoc tests were integrated into the research model in order to preserve the structural integrity of the data and identify meaningful intergenerational differences among Generations X, Y, and Z (Büyüköztürk et al., 2017; Pandey & Pandey, 2015).

The main element ensuring the methodological originality of the study and the internal validity of the data is the architecture of the Asymmetric Block Design, specifically constructed by the researcher for the presentation of synthetic (AI-generated) and real images (Kluger et al., 2025). To prevent participants’ tendency to form a mental template and make predictions through pattern recognition in visual perception tests and repeated-measures designs, the

stimuli were asymmetrically blocked in a 2-2-1-2-2-1 sequence. Creswell (2014) emphasizes that, in such experimental interventions, systematic errors such as order bias and the learning effect, in which participants alter their responses as they become familiar with the stimuli, may compromise the purity of the data. This asymmetric architecture requires participants to evaluate each image as an independent “ontological reality test,” thereby minimizing cognitive transfer and attention competition among survey items (Kluger et al., 2025; Saunders et al., 2019).

In addition, the analysis of these complex and unbalanced datasets was based on the logic of the Social Relations Model (SRM), which enables the decomposition of variance components arising from the interaction between the participant as actor and the image as partner. This advanced modeling approach makes it possible to isolate individual perceptual differences from “measurement error” and to reveal the pure statistical projection of trust in AI-generated images (Kluger et al., 2025; Singh, 2007).

3.2. Data Collection Instrument

In this study, a five-point Likert-type scale was used as the data-collection instrument to measure participants’ perceptions of synthetic and real images in digital environments in a psychometrically objective manner (Gürel, 2023). Likert scales, which constitute the foundation of attitude and perception measurement in the social sciences, convert participants’ tendencies toward variables into standardized numerical values, thereby optimizing construct validity and internal consistency coefficients (Creswell, 2014; Singh, 2007). The scale design was structured to sensitively differentiate the intensity and direction of individuals’ trust in AI-generated images and to transfer participants’ attitudes, ranging from “I completely trust” to “I do not trust at all,” into numerical data on a linear continuum (Singh, 2007). In this context, the scale’s sensitivity was based on a psychometric balance that translates participants’ cognitive thresholds into statistically meaningful variance (Pandey & Pandey, 2015). In this study, “trust” is operationalized narrowly as the perceived credibility and reliability assigned to the presented visual stimulus as a source of visual information; it does not refer to generalized interpersonal, institutional, or societal trust.

To filter out participant inattention, which is frequently encountered in online quantitative measurement designs and threatens the structural integrity of the data, a directed query/attention-check item was included in the measurement instrument as a methodological necessity. Participants who marked responses at random without focusing on all survey items, or who displayed careless-responding tendencies that disrupted the dataset’s statistical normality, were identified through this structured filtering mechanism and excluded from the analysis

(Creswell, 2014). Designed within the framework of the reliability criteria for data collection instruments emphasized by Pandey and Pandey (2015), this procedure filters out measurement errors in the dataset by requiring participants to focus solely on the concept intended to be measured; it thereby directly increases the statistical power of the research and the external reliability of the findings (Singh, 2007).

The ten images constituting the core component of the measurement instrument, five real photographs and five AI-generated images, were subjected to strict standardization as visual stimuli within the research design. To evaluate the ontological distinction between synthetic images and real photographs independently of external variables, the resolution, lighting balance, and color saturation of the images

were brought into perceptual equivalence (Aziz et al., 2024). Presenting the stimuli using a standardized protocol that controls their technical characteristics, such as size and background, eliminates attentional competition that may arise during visual perception (Qin et al., 2020). This perceptual standardization architecture, which prevents external variables from turning into confounding effects, preserves the purity of the experimental manipulation and maximizes the methodological internal validity of the research (Creswell, 2014; Pandey & Pandey, 2015). In this way, the measurement instrument goes beyond merely being a technical data-collection form and becomes a scientific instrument that tests the ontological boundaries of artificial aesthetics (Singh, 2007).



Figure 1. Real Photographs Used as Stimuli.



Figure 2. AI-Generated Images Used as Stimuli

3.3. Participants and Sampling Procedure

A non-probability online volunteer sampling approach was used. The final analytic sample comprised 399 participants who could be classified into one of the three predefined generational groups used in the study (Generation X, Generation Y/Millennials, or Generation Z) according to the birth-year boundaries reported in Table 1. Participation was voluntary. Responses were retained for analysis when participants completed the visual evaluation task and met the attention-check criterion described above; responses indicating random or careless responding were excluded. Because recruitment was online and voluntary rather than probability-based, the sample should be interpreted as an experimental sample rather than as a population-representative sample.

3.4. Principles of Visual Pairing: Paired Design

In this study, the selection and presentation of visual stimuli were based on a “paired design” protocol to control for external factors and ensure methodological rigor. In quantitative designs, maximizing internal validity, which refers to the accuracy of causal relationships, requires the structural control of confounding variables that may manipulate the dependent variable (Creswell, 2014). Including synthetic and authentic images in the experiment as non-random paired sets minimizes confounding effects from participants’ individual differences in judgment, thereby providing a statistical layer of protection for the reliability of the research data (Singh, 2007).

The core of the paired design is the principle of “stimulus equivalence.” In accordance with the principles of measurement sensitivity emphasized by Pandey and Pandey (2015), each AI-generated image used in the study was content-wise paired with a real photograph from the same category, such as a nature landscape paired with a nature landscape. Equalizing the structural characteristics of the images, including technical elements such as perspective depth, object positions, and lighting conditions, at a perceptual standard prevents the “noise” created by external variables and enables only the ontological difference, whether synthetic or real, to be isolated and measured (Aziz et al., 2024). This method eliminates threats to internal validity by ensuring that the intended effect is measured solely by the stimulus itself (Creswell, 2014). To disrupt the guessing strategies and cognitive templates that participants may develop during repeated-measurement processes, the images were presented in an asymmetric sequence specific to the study, following a 2-2-1-2-2-1 format. This presentation strategy was constructed in accordance with the logic of the Asymmetric Block Design, which enables the modeling of asymmetric data structures obtained from the dyadic interactions of two different categories, namely artificial and real images (Kluger et al., 2025). Combining the paired design with the asymmetric presentation architecture prevents participants from forming a pattern in their minds and requires them to evaluate each stimulus as an independent “partner” (Kluger et al., 2025). This integrated methodological structure filters individual perceptual differences from measurement error and allows the trust

threshold created by synthetic media to be identified in its purest form (Singh, 2007).

4. Findings

In this section, the findings obtained from the data collected within the framework of the quasi-experimental design are presented in line with the pre-determined theoretical framework, namely Baudrillard’s (1994) theory of simulation, Manovich’s (2024) model of algorithmic decomposition, and Kluger et al.’s (2025) Asymmetric Block Design protocol. The order of reporting follows the study’s hypothesis hierarchy, beginning with descriptive statistics and proceeding to increasingly complex inferential analyses. The significance threshold was set at $\alpha = .05$, and effect sizes were reported through Cohen’s d and eta-squared (η^2), respectively.

4.1. Sample Profile

The study sample consisted of 399 voluntary participants recruited online. The sample exhibited marked heterogeneity in generational distribution: the majority of participants were Generation Y ($n = 231$, 57.9%), followed by Generation Z ($n = 98$, 24.6%) and Generation X ($n = 70$, 17.5%). Although the sample was unbalanced by gender, this variable was included in the analysis: 68.7% of participants were female ($n = 274$), while 31.3% were male ($n = 125$). In terms of educational level, the majority of participants had an undergraduate education (65.9%, $n = 263$), while a substantial proportion had a postgraduate education (28.3%, $n = 113$). The sample profile is presented in detail in Table 1.

Table 1. Sample Profile: Distribution by Generation, Gender and Education Variables and Confidence Score Averages

Variable / Category		n	%	AI Trust M	Real Trust M	SD (AI)
Generation	Generation X (1965–1980)	70	17.5	15.50	12.41	2.93
	Millennials (1981–1996)	231	57.9	16.79	13.35	3.11
	Generation Z (1997–2012)	98	24.6	17.47	13.71	3.54
Gender	Women	274	68.7	16.64	13.36	3.01
	Male	125	31.3	16.94	13.08	3.72
Education Status	High school and below	23	5.8	16.09	—	3.91
	Undergraduate	263	65.9	16.87	—	3.34
	Graduate	113	28.3	16.53	—	2.87
Total		399	100.0	16.73	13.27	3.25

Note. AI Trust M = Average total confidence score of artificial intelligence images (5 items, 5–25 score range). Actual Trust M = Average total confidence score of real images. SS = Standard deviation. The small sample size ($n = 23$) in the high school and below group should be considered in the comments.

The first notable pattern in Table 1 is that the mean level of trust in AI-generated images is systematically higher than that in real images across all su-

bgroups. This structural asymmetry will be examined in depth in the following sections.

4.2. General Trust Comparison: Real And AI-Generated Images

4.2.1. Ontological paradox: the hyperreality effect

The study's primary hypothesis predicted that participants' trust scores for AI-generated images would

differ significantly from those for real images. The results of the paired-samples t-test confirmed this hypothesis with an extremely strong effect size (see Table 2): $t(398) = -21.12$, $p < .001$, $d = -1.058$. The absolute value of Cohen's d clearly exceeds the threshold of 0.80, which indicates a large effect size; this ratio suggests that approximately 47% of the variance is explained by group membership.

Table 2. The Difference Between Trust in Real and AI Images: Matched Sample t-Test Results

Image Type	n	M	SD	t	p	d
Real Images (5 items total)	399	13.27	3.25	-21.12	<.001	-1.058
Artificial Intelligence Images (5 items total)	399	16.73	3.25			
Difference (AI – Reality)	—	+3.46	—			

Note. Likert scale: 1 = I don't trust at all, 5 = I trust completely. Both total scores represent a 5-item composite measurement (range 5–25). d = Cohen's d effect size. The analysis was conducted on 399 participants with complete data for both variables.

The total mean trust score for real images was $M = 13.27$ ($SD = 3.25$), whereas the mean score for AI-generated images was $M = 16.73$ ($SD = 3.25$). The 3.46-point difference between the means corresponds to approximately 17.3% of the theoretical scale range (5–25), indicating not only a statistical but also a substantive divergence.

This finding directly resonates with Baudrillard's (2018) concept of simulacra. Baudrillard argues that, in the postmodern visual regime, simulation no longer merely represents reality but surpasses it, and that hyperreality—that is, the model appearing more real than the reality it represents—has become a normative perceptual condition. The present findings bring this proposition onto an empirical plane; participants' attribution of higher trust to AI-generated images than to real photographs can be interpreted as a measurable manifestation of the Baudrillardian hyperreality effect in the context of social media. In Manovich's (2024) terms, "algorithmic image generation automatically produces a stronger signal of credibility when human percep-

tion is left without a stable point of reference"; the present results are consistent with this mechanism among participants in the current sample, without establishing population-level generalizability.

4.3. Intergenerational Variance: ANOVA and Gender / Education Analyses

4.3.1. One-way ANOVA: generation $\times$ trust in ai-generated images

A one-way analysis of variance (ANOVA) was conducted to test whether the total trust score for AI-generated images (AI_Total, ranging from 5 to 25) differed by generation. The results indicated that the main effect of generation was statistically significant: $F(2, 396) = 7.869$, $p < .001$, $\eta^2 = .038$. The eta-squared value indicates that 3.8% of the variance is attributable to the generation variable; according to Cohen's (1988) criteria, this corresponds to a small-to-medium effect size.

Table 3. Generational Effect on Total Trust in AI Visuals: One-Way ANOVA and Tukey HSD Post-Hoc Test Results

Source of Variance	SS	df	MS	F	p	η^2
Intergroup (Generation)	146.96	2	73.48	7.869	<.001	.038
Within Groups (Error)	3704.37	396	9.35			
Total	3851.33	398				

Note. SS = Sum of squares. df = Degrees of freedom. MS = Mean of squares. η^2 = Eta-square effect size. Since no violation of the assumption of the Levene's homogeneity test was detected, standard ANOVA was applied.

Impact of AI-Generated Images on Social Media Trust Across Generations

Table 3b. Tukey HSD Post-Hoc Comparisons

Comparison	M _i	M _j	MD (i - j)	q	p	η ²
Generation X vs Millennials	15.50	16.79	-1.288	4.183	.009	.038
Generation X vs Generation Z	15.50	17.47	-1.969	5.577	< .001	
Millennials vs Gen Z	16.79	17.47	-0.682	2.505	.181 ns	

Note. M = Average of the total confidence score of artificial intelligence. MD = Mean difference. q = Studentized interval statistic. p < .01. p < .001. ns = Not statistically significant.

The Tukey HSD post-hoc test provides a more detailed explanation of the significant ANOVA finding: Generation X exhibited significantly lower AI trust scores than both Generation Y (MD = -1.288, q = 4.183, p = .009) and Generation Z (MD = -1.969, q = 5.577, p < .001). The difference between Generations Y and Z, however, did not reach the statistical significance threshold (MD = -0.682, p = .181).

This pattern reflects a mechanism conceptualized as analog skepticism. Generation X grew up in a cognitive environment shaped before the widespread adoption of digital photography and internalized the visual verification practices of the analog media era. During this period, manipulating a photograph involved both high technical and economic thresholds. By contrast, Generation Y, which emerged in the middle of the digital age, and Generation Z,

which encountered generative artificial intelligence technologies at earlier ages, approach visual media through a different epistemic framework. Paradoxically, although Generation Z is the generation most closely associated with technology, it displayed the highest mean level of trust in AI-generated images (M = 17.47). This finding indicates that digital competence does not automatically translate into visual critical capacity.

4.3.2. The effect of gender and educational levels on trust perception

In addition to the generation effect, the possible influence of gender and educational level variables on trust perception was examined. The findings are summarized in Table 3c.

Table 3c. Perception of Trust by Gender and Education Level Variables: Independent Sample t-Test and One-Way ANOVA Results

Independent Variable / Group	n	M (AI)	M (Real)	t	p
Gender: Female	274	16.64	13.36	-0.859	.391 ns
Gender: Male	125	16.94	13.08		
Education ANOVA (3 groups)	399	—	—	F = 0.911	.403 ns
High school and below	23	16.09	—		
Undergraduate	263	16.87	—		
Graduate	113	16.53	—		

Note. AI M = Average total trust in AI images. Actual M = Average total confidence in real images. For the training variable, only the AI confidence score was subjected to ANOVA. η² (training) = .005. d = Not statistically significant. In terms of gender, the independent-samples t-test revealed no significant difference in AI confidence score between female (M = 16.64, SD = 3.01) and male (M = 16.94, SD = 3.72) participants: t(397) = -0.859, p = .391. Along the same lines, actual visual confidence scores did not differ significantly according to gender: t(397) = 0.803, p = .423. The one-way ANOVA conducted for the education level also did not reveal a significant effect: F(2, 396) = 0.911, p = .403, η² = .005.

These findings indicate that individual differences in visual reality perception are not shaped by gender or education level, but primarily by generational identity and accompanying media socialization experiences. The pattern in question supports Kluger et al.'s (2025) discussion of the "observer component" within the framework of the Social Relations Model (SRM); they argue that the generation is the strongest observer variable representing the individual's systematic tendency towards visual stimuli.

4.4. Paired Image Analysis: Set-Level Comparison

4.4.1. Set-based findings and the uncanny valley effect

The five image sets within the scope of the Paired Design protocol are arranged to enable a real-artificial intelligence comparison. Paired sample t-test was applied for each set; mean difference (MD), t statistic, and Cohen's d effect size are reported in Table 4.

Table 4. Confidence Score Comparison Based on Paired Image Sets: Set-Level t-Test Results

Set/Category	n	Real M	G_SS	AI M	AI_SD	MD	t	d
Set 1 — Portrait (G1 / Y1)	397	2.18	0.91	3.53	0.93	−1.35	−21.53	−1.080
Set 2 — Landscape (G2 / Y2)	395	2.53	1.07	3.39	0.92	−0.86	−13.37	−0.673
Set 3 — Architecture (G3 / Y3)	394	2.47	1.06	3.32	1.03	−0.85	−13.17	−0.664
Set 4 — Portrait (G4 / Y4)	396	3.05	1.05	2.94	0.99	+0.11	1.625 ns	0.082
Set 5 — Landscape (G5 / Y5)	392	3.13	0.95	3.66	0.84	−0.53	−10.29	−0.520

Note. G = Real photo, Y = AI-generated image. MD = Mean difference (True − AI). d = Cohen's d. $p < .001$. ns = Not significant. Category assignments are based on researcher classification.

The most striking pattern in Table 4 is reflected in the heterogeneity among the five sets. In four of the sets—Set 1 ($d = -1.080$), Set 2 ($d = -0.673$), Set 3 ($d = -0.664$), and Set 5 ($d = -0.520$) AI-generated images received significantly higher trust scores than real photographs. By contrast, Set 4 (Portrait: G4 / Y4) was the only exception to this general tendency: $t(395) = 1.625$, $p = .105$, $d = 0.082$; this value did not reach statistical significance.

The fact that Set 1 (Portrait) showed the strongest effect, with $d = -1.080$, is a theoretically instructive finding, although it does not align with Mori's (2012) Uncanny Valley theory. Mori argues that humanoid faces evoke an intense feeling of discomfort and uncanniness when they approach a certain threshold of resemblance. According to this model, AI-generated portraits would be expected to receive lower trust scores than real faces. However, Set 1 displays precisely the opposite pattern: the AI-generated portrait (Y1, $M = 3.53$) received a trust score that was 1.35 points higher than that of the real portrait (G1, $M = 2.18$). This finding may be interpreted either as an Inverted Uncanny Valley or as an indication that generative artificial intelligence has now completely surpassed the threshold of visual perfection. Contemporary diffusion models are known to over-refine facial asymmetries, pore texture, and lighting gradients in human faces; this process effectively inc-

reases perceived realism while paradoxically reversing the uncanniness threshold described by Mori. By contrast, Set 4 (G4 / Y4) reveals the boundary of this Uncanny Valley inversion. The non-significant difference between G4 ($M = 3.05$) and Y4 ($M = 2.94$) ($p = .105$) likely reflects detectable AI artifacts in Y4, such as anatomical inconsistencies, lighting errors, or background collapse, and suggests that participants may still retain the capacity for detection when a visual falls below a certain level of quality.

4.5. Asymmetric Block Design Analysis and the SRM Perspective

4.5.1. Control of the learning effect

Kluger et al.'s (2025) Asymmetric Block Design proposes that visual presentations should be ordered according to a specific algorithm to control the practice and carryover effects, both of which are frequently encountered in repeated-measures studies. The 2-2-1-2-2-1 sequence applied in this research limits the number of consecutive images of the same type within each block, thereby preventing participants from developing a systematic expectation regarding the real/AI distinction. Table 5 presents the mean trust scores according to block and pivot positions.

Table 5. Confidence Score Distribution by Block and Pivot Positions within the Scope of Asymmetric Block Design (2-2-1-2-2-1)

Blocks / Positions	n	Image(s)	Type	M (trust)	SD
Block 1 (Positions 1–2)	795	G1, G2	Fact	2.35	1.01
Block 2 (Positions 3–4)	793	Y1, Y2	AI	3.46	0.93
Pivot 1 (Position 5)	396	G3	Fact	2.47	1.06
Block 3 (Positions 6–7)	793	Y3, G4	Mixed	3.19	1.05
Block 4 (Positions 8–9)	791	G5, Y4	Mixed	3.04	0.97
Pivot 2 (Position 10)	396	Y5	AI	3.66	0.84

Note. Intra-block averages represent the combined mean across all items in the relevant block. The pivot positions (Positions 5 and 10) contain only one image; the other blocks contain two images each. Mixed blocks (3 and 4) contain both real and artificial intelligence images.

The block design analysis reveals several critical patterns. First, the 1.11-point difference between Block 1, which consisted only of real images ($M = 2.35$), and Block 2, which consisted only of AI-generated images ($M = 3.46$), replicates the overall real–AI divergence even at the early stages of the experiment. This weakens the hypothesis that the difference was artificially produced by a learning effect. Second, G3 in the Pivot 1 position ($M = 2.47$) and Y5 in the Pivot 2 position ($M = 3.66$) indicate that placing a single image in the middle of the blocks was effective. The isolated AI-generated image (Y5) received the highest mean trust score among all images, suggesting that AI-generated images may be perceived as even more convincing in the absence of a block context. Evaluated through the terminology of Kluger et al.'s (2025) Social Relations Model (SRM), two main components can be distinguished. The partner effect refers to the systematic response produced by the visual stimulus across all participants, regardless of its

type; the fact that Y5 received the highest score across all generations provides strong evidence for this component. The observer effect, on the other hand, reflects the individual participant's general tendency toward all images. The generation-based ANOVA findings indicate that Generation X's systematically lower trust attribution represents an observer effect, showing that this tendency persists consistently and independently of stimulus type. These two effects overlap to produce the structural trust asymmetry that constitutes the study's central finding.

4.6. Image-Based Descriptive Statistics

Table 6 presents the descriptive statistics of the trust scores for each of the 10 visual stimuli used in the study, in accordance with the order of set pairing. This table describes the image-level variance and provides a micro-level reference framework for the inferential analyses presented above.

Table 6. Confidence Score Descriptive Statistics for All Visual Stimuli (Question B, $N = 394\text{--}398$)

Image	Category (Couple)	Type	N	M	SD	Min	Max
G1	Set 1 — Portrait	Fact	397	2.18	0.91	1	5
G2	Set 2 — Landscape	Fact	398	2.53	1.07	1	5
Y1	Set 1 — Portrait	AI	398	3.53	0.93	1	5
Y2	Set 2 — Landscape	AI	395	3.39	0.92	1	5
G3	Set 3 — Architecture	Fact	396	2.47	1.06	1	5
Y3	Set 3 — Architecture	AI	396	3.32	1.03	1	5
G4	Set 4 — Portrait	Fact	397	3.05	1.05	1	5
G5	Set 5 — Landscape	Fact	394	3.13	0.95	1	5
Y4	Set 4 — Portrait	AI	397	2.94	0.99	1	5
Y5	Set 5 — Landscape	AI	396	3.66	0.84	1	5

Note. G = Real photo; Y = Artificial intelligence-generated image. Scale: 1 = I don't trust at all, 5 = I trust completely. Set matching was carried out according to the Asymmetric Block Design protocol. The rate of missing data remains at most 1.3% per image.

When the descriptive statistics at the image level are examined, a notable pattern emerges among the real images: G1 ($M = 2.18$) and G2 ($M = 2.53$) on the one hand, and G4 ($M = 3.05$) and G5 ($M = 3.13$) on the other. This difference indicates that the real images positioned in the early block received lower trust scores and supports the hypothesis that block position plays a contextual role in the perception of trust. Among the AI-generated images, Y5 ($M = 3.66$) stands out with the highest mean, whereas Y4 ($M = 2.94$) has the lowest mean. This divergence suggests that participants were partially able to perceive quality variance among generative AI outputs; however, this perception did not translate into a systematic ability to differentiate between them.

4.7. General Evaluation: Theoretical Integration of the Findings

The findings presented in this section can be synthesized within the framework of three complementary theoretical layers. At the first layer, Baudrillard's (2018) thesis of hyperreality finds a limited empirical correspondence: within the tested stimuli, AI-generated images received higher average trust ratings than real photographs. Rather than documenting a society-wide displacement of reality, this pattern can be interpreted as an experimental instance consistent with the hyperreality thesis. At the second layer, Manovich's (2024) model of algorithmic decomposition becomes relevant. Manovich argues

that machine vision and generative artificial intelligence optimize visual statistics according to human intuitive standards; therefore, images produced by algorithms generate a visual economy that meets or exceeds the expectations of the average photographic viewer. In the present study, the relatively low standard deviations for AI-generated images, particularly those in the portrait pair in Set 1 ($SD = 0.84\text{--}0.93$), indicate a comparatively homogeneous distribution of responses and suggest that algorithmic image production creates a form of consensus among participants' reactions. At the third layer, Kluger et al.'s (2025) Asymmetric Block Design and SRM framework explain the structural source of individual differences in the findings. The fact that the generation effect ($X < Y \approx Z$) was both significant and independent of gender and education supports Kluger et al.'s proposition regarding the observer component and suggests that, among the socio-demographic variables tested here, generation was the clearest predictor of visual trust perception. The convergence of these three layers demonstrates that the perception of visual reality on social media is a multilayered process that requires analysis at the macro level through Baudrillard, the meso level through Manovich, and the micro level through Kluger et al.'s SRM framework.

5. Discussion and Conclusion

This study examined the transformative effects of generative artificial intelligence (GenAI) technologies on perceptions of visual reality and how these effects differ across generations, using a quasi-experimental design. The growing accessibility of deepfake technologies has enabled the rapid production and dissemination of hyper-realistic synthetic images, increasing their potential to manipulate perceptions of reality (Söğütülüler, 2025). Within the present experimental context, the findings suggest that photographic origin did not function as a guarantee of perceived credibility: for several tested stimuli, participants rated AI-generated images as more trustworthy than real photographs. Participants' total trust score for AI-generated images ($M = 16.73$) was significantly higher than their trust score for real images ($M = 13.27$), revealing a clear "trust reversal" in the perception of visual reality. This result is consistent with Baudrillard's (2018) approach to simulation and hyperreality, since AI-generated images constitute a new visual regime with no direct counterpart in the physical world, yet they produce a sense of reality. Manovich's (2001) account of digital media, in terms of the interaction between databases and algorithms, provides a theoretical framework for interpreting how algorithmically optimized image surfaces can shape perceived visual credibility.

One of the strongest findings of the study is the large effect size favoring AI-generated images, as indicated by the paired-samples t-test: $t(398) = -21.12$, $p < .001$, $d = -1.058$. This finding points to a significant rupture not only statistically but also theoretically. In parallel with Nightingale and Farid's (2022) study, which showed that AI-generated faces can be indistinguishable from real faces and may even be perceived as more trustworthy, the present study suggests that, within this sample, some synthetic images produced a stronger impression of authenticity and credibility than matched real photographs. Similarly, Miller et al. (2023) show that some AI-generated faces are perceived as more real than human faces, while Osińska et al. (2025) note that the photorealistic quality of AI-generated images challenges users' visual discrimination abilities. Therefore, the findings suggest that, for some stimuli in this experiment, synthetic images can achieve a perceptual advantage over matched real photographs.

The analyses conducted through paired visual sets showed that this trust reversal does not operate with the same intensity across all visual categories. In four of the five sets, AI-generated images received significantly higher trust scores than real images, whereas this difference was not significant in only Set 4. In particular, the fact that the AI-generated image in the portrait comparison in Set 1 received much higher trust than the real portrait is a striking result in relation to Mori's (2012) "uncanny valley" approach. While the classical theory of the uncanny valley argues that humanoid artificial representations produce discomfort at a certain threshold of realism, the present study shows that some AI-generated portraits surpass this threshold and generate trust. This situation is consistent with Kishnani's (2025) findings that uncanniness increases at medium levels of realism but may decrease at high levels. In other words, contemporary generative AI systems not only narrow the uncanny valley; in some cases, they reverse it by making the artificial appear more trustworthy.

This result can also be explained through Arielli's (2024) concepts of "artificial aesthetics" and "aesthetic alignment." AI models produce perfected surfaces by optimizing the average aesthetic patterns they learn from massive visual datasets. This perfection weakens the viewer's reflex to search for visual flaws and transforms the synthetic image into a trustworthy visual object. The criteria of photorealism, image quality, and text-image alignment emphasized by Aziz et al. (2024) also show that user trust is related not only to content accuracy but also to the technical and aesthetic coherence of the visual surface. Therefore, the findings of this study demonstrate that visual trust is no longer tied solely to the "real source" but also to the aesthetically persuasive power generated algorithmically. Intergenerational analyses showed that trust in AI-generated

images differs most clearly by socio-demographic generation. The one-way ANOVA revealed a significant generation effect, showing that Generation X's trust in AI-generated images was significantly lower than that of both Generations Y and Z. By contrast, the difference between Generations Y and Z was not significant. This finding is important because it shows that proximity to digital technology does not necessarily overlap with critical visual literacy. While Prensky's (2001) concept of "digital natives" suggests that younger generations establish a more natural relationship with technological tools, Buckingham (2003) and Hobbs (2017) emphasize that technical proficiency does not necessarily mean the capacity to critically analyze media messages. The findings of this study likewise reveal that Generation Z, despite its technical familiarity, may be more likely to trust synthetic images.

At this point, Generation X's lower trust scores can be interpreted within the framework of "analog skepticism." Individuals socialized in an analog media environment may have experienced the documentary status of photography within a different historical context and developed a more cautious perception toward digital manipulation. By contrast, Generations Y and Z, who are continuously exposed to visual flows on social media platforms, may evaluate AI-generated images as an ordinary part of the digital experience. Yıldız's (2025) study on Generation Z and synthetic media shows that young people adopt AI tools for ease of use and practical benefits yet also develop epistemic concerns about manipulative content. The quantitative findings of the present study complement this phenomenological observation by demonstrating that young users may assign higher visual trust scores even when they experience such concerns. It is also noteworthy that gender and educational level did not yield significant differences in trust toward AI-generated images. This result suggests that the perception of visual reality is related less to classical demographic variables than to media socialization and generational experience. The fact that educational level did not serve as a significant protective variable indicates that digital literacy cannot be explained solely by formal education level. In this context, Kellner and Share's (2007) treatment of critical media literacy as an essential skill for democratic societies is particularly important. The present findings show that a high level of education alone is insufficient to distinguish synthetic images or to attribute lower trust to them; rather, specific forms of AI literacy, visual verification, and algorithmic awareness training are required.

The findings on the Asymmetric Block Design strengthen the study's methodological originality. Presenting the images in a 2-2-1-2-2-1 sequence prevented participants from detecting a simple pattern and adjusting their responses accordingly; never-

theless, the trust difference in favor of AI-generated images was maintained. The fact that the mean score for real images in Block 1 was $M = 2.35$, while the mean score for AI-generated images in Block 2 was $M = 3.46$, shows that the trust reversal emerged even at the early stage of the experiment. In addition, the Y5 image in the Pivot 2 position received the highest mean trust score ($M = 3.66$), indicating that certain AI-generated images can be highly persuasive even when presented without contextual support. Evaluated within the framework of Kluger et al.'s (2025) Social Relations Model, this finding points both to the partner effect arising from the visual stimulus itself and to the observer effect associated with generation. The findings of this study are also significant for debates on algorithmic authority and disinformation in social media environments. Castells's (2008) argument that information flows restructure power relations in the network society becomes even more relevant with the circulation of AI-generated images through platform algorithms. Pasquale's (2015) "black box" metaphor describes an environment in which users consume visual information without knowing by which criteria specific content is made visible. In this context, AI-generated images create not merely an individual perceptual error, but a systematic trust problem reinforced by platform architectures. Gül Ünlü and Küçükşabanoğlu's (2023) observation that artificial intelligence plays a dual role in both the production of disinformation and the struggle against it aligns with this study's findings. From the perspective of the literature on deepfakes and synthetic media, the study's results show that the visual trust crisis should be addressed not only at the level of individual misperception but also at the levels of social trust and the democratic public sphere. While Altuncu et al. (2024) state that deepfake technologies are rapidly evolving in their definitions, measurement, datasets, and standards, Zan and Zan (2025) emphasize that these technologies pose serious pedagogical, ethical, and social risks. İpek and İpek (2025) likewise note that deepfake images transform the perception of reality within synthetic reality. The fact that trust in AI-generated images exceeded trust in real images in this study demonstrates that these risks are not merely theoretical but also empirically observable.

In conclusion, this study shows that the perception of visual reality is undergoing a serious ontological and epistemological transformation in the age of generative artificial intelligence. Participants' perception of AI-generated images as more trustworthy than real photographs reveals that photography's status as "evidence of truth" has weakened and that algorithmic image production has created a new regime of trust. The significance of intergenerational differences shows that this transformation is not experienced in the same way across all user groups; Generation X appears more cautious, whereas Ge-

nerations Y and Z are more likely to trust synthetic images. By contrast, the fact that gender and educational level did not yield significant differences suggests that the solution lies not merely in raising the overall level of education, but in developing generation-specific strategies for critical digital and visual literacy. Within this framework, overcoming the trust crisis created by AI-generated images requires intervention at three levels. First, critical media literacy programs should be expanded to improve users' ability to recognize, verify, and question synthetic content. Second, mechanisms for labeling AI-generated content, ensuring source transparency, and strengthening algorithmic accountability on social media platforms should be reinforced (Diakopoulos, 2016). Third, by taking into account generation-specific media use practices, safe verification practices should be prioritized for Generation X, professional-digital balance skills for Generation Y, and synthetic media awareness and AI literacy for Generation Z. In this way, while preserving the creative possibilities offered by generative artificial intelligence, the potential of synthetic images to produce disinformation, perception management, and the erosion of social trust can be limited.

5.1. Limitations and Future Research

Several limitations should be considered when interpreting the findings. First, the study relied on a 399-participant online volunteer sample that was not probability-based and was unevenly distributed across generations, gender, and education; therefore, the results should not be generalized to all social media users or to society at large. Second, the experiment used only ten standardized visual stimuli (five real photographs and five AI-generated images). Image-specific characteristics may have influenced the observed effect sizes, and the stimulus set cannot represent the full diversity of synthetic media. Third, the design did not treat the image-generation platform, model version, prompting strategy, post-processing, or platform-specific presentation as experimental factors. Because different systems may produce different aesthetic tendencies and levels of photorealism, AI-generated imagery should not be interpreted as a homogeneous category on the basis of this study.

Fourth, generative image technologies and users' familiarity with them are changing rapidly, which means that the trust pattern observed here may be time-sensitive. Finally, the Likert-type measure captures perceived credibility and trustworthiness under controlled exposure rather than actual verification, sharing, or belief behavior in a naturalistic social media feed. Future research should therefore use larger and more diverse samples, broader and periodically updated stimulus sets, multi-platform and multi-model comparisons, longitudinal replications, and

field-based designs that examine how source cues, labels, and platform context affect visual credibility over time.

5.2. Conclusion and Recommendations

In conclusion, within the boundaries of this 399-participant, ten-image quasi-experimental study, the findings suggest that highly realistic AI-generated imagery can alter perceived visual credibility. In the experimental context, participants rated AI-generated images as more trustworthy than real photographs overall; this pattern suggests that photography's conventional role as "evidence of truth" can be challenged in some digital-media contexts and that algorithmically generated imagery may produce a reconfigured pattern of perceived trust. These results should not be interpreted as evidence of a society-wide transformation. The observed generational differences further suggest that this pattern is not experienced uniformly: Generation X participants were more cautious, whereas Generations Y and Z assigned higher trust scores to synthetic images. The absence of significant gender and education effects indicates that general educational attainment alone may be insufficient and that more targeted forms of critical digital, visual, and AI literacy are needed. Accordingly, a multi-level response is recommended. At the individual level, media-literacy programs should strengthen visual verification skills, source checking, reverse-search habits, and awareness of common synthetic-media cues. At the platform level, visible labeling of AI-generated content, clearer source and provenance information, accessible verification tools, and stronger algorithmic accountability should be reinforced (Diakopoulos, 2016). At the educational and institutional levels, training should be adapted to generation-specific media practices while avoiding assumptions that technological familiarity automatically produces critical literacy. For researchers, repeated studies with updated stimuli, multiple AI-generation systems, and longitudinal or field designs are necessary to determine whether the present pattern persists as generative technologies and user expectations evolve. Such measures can help preserve the creative benefits of generative AI while reducing the risks of disinformation, manipulation, and misplaced trust in synthetic visual content.

References

- Altuncu, E., Franqueira, V. N. L., & Li, S. (2024). Deepfake: Definitions, performance metrics and standards, datasets, and a meta-review. *Frontiers in Big Data*, 7, 1400024. <https://doi.org/10.3389/fdata.2024.1400024>
- Arielli, E. (2024). Made by and for humans? The issue of aesthetic alignment. In L. Manovich & E. Arielli (Eds.), *Artificial aesthetics: Generative AI, art and visual media* (pp. 170–191). <https://manovich.net/index.php/projects/artificial-aesthetics>
- Arslan, E. (2026). Dezenformasyonla karşılaşma deneyimi, ileti-

şimsel tepkiler ve psikolojik yansımalar: 60+ Tazelenme Üniversitesi öğrencileri örneği. *The Turkish Online Journal of Design, Art and Communication*, 16(1), 219–236. <https://doi.org/10.7456/tojdac.1797644>

Aziz, M., Rehman, U., Safi, S. A., & Abbasi, A. Z. (2024). Visual verity in AI-generated imagery: Computational metrics and human-centric analysis [Preprint]. *arXiv*. <https://doi.org/10.48550/arXiv.2408.12762>

Baudrillard, J. (2018). *Simülakrlar ve simülasyon* (O. Adanır, Trans.). Doğu Batı Yayınları.

Bińkowski, M., Sutherland, D. J., Arbel, M., & Gretton, A. (2018). Demystifying MMD GANs [Preprint]. *arXiv*. <https://doi.org/10.48550/arXiv.1801.01401>

Boztunç, N. (2014). Y kuşağını elde tutabilme üzerine bir çalışma [Unpublished graduation project]. Ege Üniversitesi Sosyal Bilimler Enstitüsü.

Buckingham, D. (1998). Media education in the UK: Moving beyond protectionism. *Journal of Communication*, 48(1), 33–43. <https://doi.org/10.1111/j.1460-2466.1998.tb02735.x>

Buckingham, D. (2003). *Media education: Literacy, learning and contemporary culture*. Polity Press.

Büyüköztürk, Ş., Kılıç Çakmak, E., Akgün, Ö. E., Karadeniz, Ş., & Demirel, F. (2017). *Bilimsel araştırma yöntemleri* (23. baskı). Pegem Akademi. <https://doi.org/10.14527/9789944919289>

Canbolat, K. (2022). Kuşaklar arası dijital mahremiyet algısının dönüşümü [Unpublished master's thesis]. Çanakkale Onsekiz Mart Üniversitesi.

Castells, M. (2008). Enformasyon çağı: Ekonomi, toplum ve kültür, ağ toplumunun yükselişi (E. Kılıç, Trans.). İstanbul Bilgi Üniversitesi Yayınları.

Castells, M. (2009). *Communication power*. Oxford University Press.

Cohen, J. (1988). *Statistical power analysis for the behavioral sciences* (2nd ed.). Routledge. <https://doi.org/10.4324/9780203771587>

Creswell, J. W. (2014). *Research design: Qualitative, quantitative, and mixed methods approaches* (4th ed.). SAGE Publications.

Diakopoulos, N. (2015). Algorithmic accountability: Journalistic investigation of computational power structures. *Digital Journalism*, 3(3), 398–415. <https://doi.org/10.1080/21670811.2014.976411>

Diakopoulos, N. (2016). Accountability in algorithmic decision making. *Communications of the ACM*, 59(2), 56–62. <https://doi.org/10.1145/2844110>

Doğan, S. (2024). Sosyal medyada görsel üretim için yapay zekâ kullanımı. In H. N. Tarakçı & N. Tufan Yeniçikti (Eds.), *Sosyal medyada iletişim-II* (pp. 229–243). Palet Yayınları.

Foucault, M. (1992). *Hapishanenin doğuşu*. İmge Kitabevi.

Goodfellow, I. J., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., & Bengio, Y. (2014). Generative adversarial nets [Preprint]. *arXiv*. <https://doi.org/10.48550/arXiv.1406.2661>

Göring, S., Rao, R. R. R., Merten, R., & Raake, A. (2023). Analysis of appeal for realistic AI-generated photos. *IEEE Access*, 11, 38999–39012. <https://doi.org/10.1109/ACCESS.2023.3267968>

Gül Ünlü, D., & Küçükşabanoglu, Z. (2023). Dezenformasyon ve yapay zekâ: Dezenformasyonla mücadele yollarına yapay zekâ uzmanlarının gözünden bakmak. *İletişim ve Diplomasi*, 11, 83–106. <https://doi.org/10.54722/iletisimvediplomasi.1375478>

Gürel, E. (2023). *Dijital medyada siyasal iletişim ve Z kuşağı üniversite öğrencileri: İstanbul özelinde çevrimiçi siyasal katılım üzerine bir araştırma* [Master's thesis]. Maltepe Üniversitesi.

Hashmi, A., Shahzad, S. A., Lin, C.-W., Tsao, Y., & Wang, H.-M. (2024). Understanding audiovisual deepfake detection: Techniques, challenges, human factors and perceptual insights [Preprint]. *arXiv*. <https://doi.org/10.48550/arXiv.2411.07650>

Hobbs, R. (1998). Teaching with and about film and television: Integrating media literacy concepts into management education. *Journal of Management Development*, 17(4), 259–272. <https://doi.org/10.1108/02621719810210136>

doi.org/10.1108/02621719810210136

Hobbs, R. (2017). *Create to learn: Introduction to digital literacy*. John Wiley & Sons. <https://doi.org/10.1002/9781394260201>

Hutson, J., & Smith, A. A. (2025). *Cinematic algorithms: The rise of generative AI in video art and visual culture*. CRC Press. <https://doi.org/10.1201/9781003601128>

İpek, A. R., & Ermiş İpek, S. (2025). Sentetik gerçeklik bağlamında yapay zekâ destekli derin sahte görsellerin gerçeklik algısını dönüştürmesi. *Sanat Yazıları*, (53), 725–744. <https://doi.org/10.61742/sanatyzalari.1761815>

Karakaya, Ü., & Dayı, H. (2024). Yapay zeka destekli görüntü üretimi ve gerçeklik ilişkisi. *Akdeniz Sanat*, 18(33), 9–31. <https://doi.org/10.48069/akdenizsanat.1389695>

Karaoğlu, G. (2022). *Dijital okuryazarlık*. In G. Karaoğlu (Ed.), *Dijital okuryazarlık* (pp. 43–58). Efe Akademi.

Kellner, D., & Share, J. (2007). Critical media literacy is not an option, it's a necessity. *Learning Inquiry*, 1(1), 59–69. <https://doi.org/10.1007/s11519-007-0004-2>

Kishnani, D. K. (2025). The uncanny valley: An empirical study on human perceptions of AI-generated text and images [Master's thesis, Massachusetts Institute of Technology]. DSpace@MIT. <https://hdl.handle.net/1721.1/159096>

Kluger, A. N., Ackerman, R. A., Kenny, D. A., Malloy, T. E., & Eastwick, P. W. (2025). The social-relations model for asymmetric-block-design data: A tutorial with R. *Advances in Methods and Practices in Psychological Science*, 8(1), 1–40. <https://doi.org/10.1177/25152459241279522>

Kress, G., & van Leeuwen, T. (2021). *Reading images: The grammar of visual design* (3rd ed.). Routledge. <https://doi.org/10.4324/9781003099857>

Manovich, L. (2001). *The language of new media*. MIT Press.

Manovich, L. (2024). AI aesthetics and media evolution. In L. Manovich & E. Arielli (Eds.), *Artificial aesthetics: Generative AI, art and visual media* (pp. 119–144). <https://manovich.net/index.php/projects/artificial-aesthetics>

Manovich, L., & Arielli, E. (2024). *Artificial aesthetics: Generative AI, art and visual media*. <https://manovich.net/index.php/projects/artificial-aesthetics>

Metzger, M. J., & Flanagin, A. J. (2013). Credibility and trust of information in online environments: The use of cognitive heuristics. *Journal of Pragmatics*, 59, 210–220. <https://doi.org/10.1016/j.pragma.2013.07.012>

Miller, E. J., Steward, B. A., Witkower, Z., Sutherland, C. A. M., Krumhuber, E. G., & Dawel, A. (2023). AI hyperrealism: Why AI faces are perceived as more real than human ones. *Psychological Science*, 34(12), 1390–1403. <https://doi.org/10.1177/09567976231207095>

Mori, M. (2012). The uncanny valley (K. F. MacDorman & N. Kageki, Trans.). *IEEE Robotics & Automation Magazine*, 19(2), 98–100. <https://doi.org/10.1109/MRA.2012.2192811>

Nightingale, S. J., & Farid, H. (2022). AI-synthesized faces are indistinguishable from real faces and more trustworthy. *Proceedings of the National Academy of Sciences*, 119(8), e2120481119. <https://doi.org/10.1073/pnas.2120481119>

Noble, S. U. (2018). *Algorithms of oppression: How search engines reinforce racism*. New York University Press. <https://doi.org/10.18574/nyu/9781479833641.001.0001>

Osińska, V., Kortas, W., Szalach, A., & Welter, M. (2025). AI images vs. real photographs: Investigating visual recognition and perception. *Journal of Eye Movement Research*, 18(6), 61. <https://doi.org/10.3390/jemr18060061>

Palfrey, J., & Gasser, U. (2008). *Born digital: Understanding the first generation of digital natives*. Basic Books.

Pandey, P., & Pandey, M. M. (2015). *Research methodology: Tools and techniques*. Bridge Center.

Pasquale, F. (2015). *The black box society: The secret algorithms that control money and information*. Harvard University Press. <https://doi.org/10.4159/harvard.9780674736061>

Prensky, M. (2001). Digital natives, digital immigrants part 1. On the

- Horizon, 9(5), 1–6. <https://doi.org/10.1108/10748120110424816>
- Qin, N., Xue, J., Chen, C., & Zhang, M. (2020). The bright and dark sides of performance-dependent monetary rewards: Evidence from visual perception tasks. *Cognitive Science*, 44(3), e12825. <https://doi.org/10.1111/cogs.12825>
- Ramesh, A., Pavlov, M., Goh, G., Gray, S., Voss, C., Radford, A., Chen, M., & Sutskever, I. (2021). Zero-shot text-to-image generation. *Proceedings of the 38th International Conference on Machine Learning*, 139, 8821–8831.
- Rapp, A., Di Lodovico, C., Torrielli, F., & Di Caro, L. (2025). How do people experience the images created by generative artificial intelligence? An exploration of people's perceptions, appraisals, and emotions related to a Gen-AI text-to-image model and its creations. *International Journal of Human-Computer Studies*, 193, 103375. <https://doi.org/10.1016/j.ijhcs.2024.103375>
- Sanchez-Acedo, A., Carbonell-Alcocer, A., Cascarano, P., Hajahmadi, S., Gertrudix, M., & Marfia, G. (2026). Exploring the impact of visual components on the perceived realism of generative AI videos. *Online Journal of Communication and Media Technologies*, 16(1), e202603. <https://doi.org/10.30935/ojcm/17737>
- Saunders, M., Lewis, P., & Thornhill, A. (2019). *Research methods for business students* (8th ed.). Pearson.
- Seymen, A. F. (2017). Y ve Z kuşak insanı özelliklerinin Millî Eğitim Bakanlığı 2014–2019 stratejik programı ve TÜBİTAK Vizyon 2023 öngörülleri ile ilişkilendirilmesi. *Kent Akademisi*, 10(4), 467–489.
- Sezgin, S. (2025). Yapay zekâ çağında dijital uçurumu yeniden düşünmek: Eğitime yönelik altı boyutlu bir kavramsal çerçeve. *Mehmet Akif Ersoy Üniversitesi Eğitim Bilimleri Enstitüsü Dergisi*, 13(17), 1–28. <https://dergipark.org.tr/pub/ebed/article/1832400>
- Singh, K. (2007). *Quantitative social research methods*. SAGE Publications. <https://doi.org/10.4135/9789351507741>
- Sontag, S. (2008). *Fotoğraf üzerine* (O. Akınhay, Trans.). Agora Kitaplığı.
- Söğütlüler, T. (2025). Deepfake manipulation on social media: A case study of fraudulent activities in Türkiye. *Intermedia International E-Journal*, 12(22), 237–265. <https://doi.org/10.56133/intermedia.1614843>
- van Hees, J., Grootswagers, T., Quek, G. L., & Varlet, M. (2025). Human perception of art in the age of artificial intelligence. *Frontiers in Psychology*, 15, 1497469. <https://doi.org/10.3389/fpsyg.2024.1497469>
- Yıldız, G. (2025). Yapay zekâ çağında gençlik: Sentetik medya içerikleri ve eleştirel okuryazarlık üzerine fenomenolojik bir inceleme. *İmgelem, Yeni Medya Çalışmaları Özel Sayısı*, 315–350. <https://doi.org/10.53791/imgelem.1732361>
- Zan, N., & Zan, B. U. (2025). Yapay zekâ destekli derin sahtecilik teknolojisinin pedagojik, etik ve toplumsal boyutları. *Turkish Studies*, 20(4), 2961–2983. <https://doi.org/10.7827/TurkishStudies.83049>
- Zuboff, S. (2019). The age of surveillance capitalism: The fight for a human future at the new frontier of power. *PublicAffairs*.