



INSTITUTO
UNIVERSITÁRIO
DE LISBOA

Exploring the Impact of Deepfake Advertisements on Ad Avoidance and Consumer Behavior in the Fashion Industry

Inês Costa Rodrigues

MSc in Marketing

Supervisor:

PhD João Ricardo Paulo Marques Guerreiro, Associate Professor,
Department of Marketing, Operations & General Management at ISCTE
Business School

September 2024

Department of Marketing, Operations and General Management

Exploring the Impact of Deepfake Advertisements on Ad Avoidance and Consumer Behavior in the Fashion Industry

Inês Costa Rodrigues

MSc in Marketing

Supervisor:

PhD João Ricardo Paulo Marques Guerreiro, Associate Professor,
Department of Marketing, Operations & General Management at ISCTE
Business School

Resumo

A publicidade deepfake é uma realidade iminente para a futura remodelação da indústria da publicidade. Por conseguinte, o principal objetivo desta investigação é investigar o potencial impacto dos anúncios deepfake no comportamento e na percepção do consumidor, com um enfoque específico na no ato de evitar os anúncios na indústria da moda. À medida que a IA continua a evoluir como uma tecnologia disruptiva, surgem novas oportunidades para os profissionais de marketing e anunciantes de moda explorarem. O estudo visa fornecer informações sobre este fenómeno emergente e as suas implicações para o futuro da publicidade, centrando-se na forma como a tecnologia deepfake pode influenciar a confiança do consumidor e as percepções de autenticidade. Foi realizada uma experiência em linha que envolveu 268 participantes para explorar a forma como a percepção de verosimilhança, criatividade e relevância influenciam as reacções dos consumidores a anúncios deepfake em comparação com anúncios não deepfake. Os resultados do estudo sugerem que, embora a publicidade deepfake ofereça oportunidades criativas, também coloca desafios relacionados com o ceticismo e a confiança dos consumidores. Esta investigação constitui um pequeno, mas promissor, passo para explorar o potencial dos conteúdos gerados por deepfake. Em última análise, este estudo fornece informações valiosas para aproveitar a tecnologia deepfake de forma responsável e estabelece as bases para futuras investigações sobre as suas implicações no comportamento do consumidor e na publicidade na indústria da moda.

Palavras-chave: Deepfake; Publicidade gerada por Deepfake; Inteligência artificial; Anúncios; Publicidade; Indústria da moda

Classification JEL: M31 Marketing & M37 Advertising

Abstract

Deepfake advertisement is an imminent reality for the future reshaping of the advertising industry. Therefore, this research's primary objective is to investigate the potential impact of deepfake advertisements on consumer behavior and perception, with a specific focus on ad avoidance within the fashion industry. As AI continues to evolve as a disruptive technology, new opportunities emerge for fashion marketers and advertisers to explore. The study aims to provide insights into this emerging phenomenon and its implications for the future of advertising by focusing on how deepfake technology can influence consumer trust and perceptions of authenticity. An online experiment involving 268 participants was conducted to explore how perceived verisimilitude, creativity, and relevance influence consumers' responses to deepfake advertisements compared to non-deepfake ones. Findings from the study suggest that while deepfake advertisement offer creative opportunities, they also pose challenges related to consumer skepticism and trust. This research provides a small, yet promising, step towards exploring the potential of deepfake-generated content. Ultimately, this study provides valuable insights for leveraging deepfake technology responsibly and lays the foundation for future research on its implications for consumer behavior and advertising in the fashion industry.

Key-words: Deepfake; Deepfake-generated Advertising; Artificial Intelligence; Advertisement; Advertising; Fashion Industry

Classification JEL: M31 Marketing & M37 Advertising

Table of Contents

Introduction	1
Chapter 1: Literature Review	3
1.1. Fashion Advertisement	3
1.1.1. Key Theories and Findings in Fashion Advertising	3
1.1.2. The Future of Fashion Advertisement	5
1.2. Evolution of Advertisement Manipulation	6
1.2.1. Generation 1.0. Analog Manipulation	7
1.2.2. Generation 2.0. Digital Manipulation	7
1.2.3. Generation 3.0. Synthetic Manipulation	8
1.3. Introduction to Deepfake Technology	9
1.3.1. Deepfakes – an Illusion or a New Hope for Marketing	12
1.3.1.1. Deepfake in Advertising and Marketing	15
1.4. Trust and Perceptions of Authenticity	16
1.5. Ad Avoidance	17
1.5.1. Ad Avoidance in the Age of Deepfake Advertisement	18
Chapter 2: Conceptual Model and Hypothesis of Investigation	19
Chapter 3: Research Methodology	25
3.1. Construct Measurement	25
3.2. The Experiment	26
3.3. Experimental Design	27
3.4. Pre-test	28
3.5. Data Collection and Procedures	28
Chapter 4: Results of Findings	29
4.1. Demographic Description	29
4.2. Data Analysis	30

4.2.1. Assessment of Measurement Model - Outer Model	30
4.2.2. Assessment of Structural Model - Inner Model	32
4.2.3. Mediation Analysis	34
4.2.4. Analysis of Model 1 – Deepfake	35
4.2.5. Analysis of Model 2 – Non-Deepfake	38
Chapter 5: Discussion of Findings	41
5.1. Perceived Verisimilitude and Ad Avoidance	41
5.2. Perceived Verisimilitude and Awareness of Falsity	42
5.3. Perceived Creativity and Ad Avoidance	42
5.4. Perceived Creativity and Awareness of Falsity	43
5.5. Perceived Relevance and Ad Avoidance	43
5.6. Awareness of Falsity and Ad Avoidance	44
5.7. Trust in Brand and Perceived Relevance and Ad Avoidance	44
5.8. Effects of Increased Manipulation	45
Chapter 6: Conclusion	47
References	50
Annexes	65
Annex A: Viral Deepfake	65
Annex B: Measurement Scales	66
Annex C: Questionnaire Visual Stimuli	67
Annex D: Indicator Loadings Complete Model (with PV3)	67
Annex E: Indicator Loadings Complete Model (without PV3)	67
Annex F: Indicator Loadings Deepfake Model (with PV3)	67
Annex G: Indicator Loadings Deepfake Model (without PV3)	68
Annex H: Indicator Loadings Non-Deepfake Model (with PV3)	68
Annex I: Indicator Loadings Non-Deepfake Model (without PV3)	68

Annex J: Construct Reliability and Convergent Validity	69
Annex K: Variance Inflation Factor (VIF) Complete Model	70
Annex L: Variance Inflation Factor (VIF) Deepfake Model	70
Annex M: Variance Inflation Factor (VIF) Non-Deepfake Model	70
Annex N: Complete Model Discriminant Validity: Fornell-Larker Criterion	70
Annex O: Deepfake Model Discriminant Validity: Fornell-Larker Criterion	71
Annex P: Non-Deepfake Model Discriminant Validity: Fornell-Larker Criterion	71

Tables Index

Table 1.1: Key and emerging theories in fashion advertising. Adapted from Taylor & Costello (2017)	4
Table 1.2: Key findings in fashion advertising. Adapted from Taylor & Costello (2017)	4
Table 1.3: Generations of Manipulation in Advertising. Adapted from Campbell, Plangger, Sands & Kietzmann (2022)	7
Table 1.4: Types and examples of deepfakes business applications. Adapted from Kietzmann, Lee, McCarthy & Kietzmann (2020)	15
Table 3.1: Measurement Scale	25
Table 4.1: PLS-SEM Bootstrapping Results – Complete dataset	33
Table 4.2: PLS-SEM Bootstrapping Results Deepfake dataset	36
Table 4.3: PLS-SEM Bootstrapping Results Non-Deepfake dataset	39

Figures Index

Figure 1.1: Visualization of the Taxonomy of Deepfakes. Adapted from Masood, Nawaz, Malik, Javed, Irtaza & Malik (2022) And Firc, Malika & Hanáček (2023)	10
Figure 2.1: Conceptual Framework Model, Based on the Theoretical Model “A framework for consumer response to manipulated advertising” (Campbell, Plangger, Sands & Kietzmann, 2022)	20

Introduction

With the rapid evolution of AI technology, distinguishing between authentic and synthetic content has become more challenging, both humans and machines have difficulty recognizing deepfakes (Firc, Malinka & Hanáček, 2023). This is primarily due to AI's ability to convincingly replicate human abilities, such as generating realistic text, images, and even audio and video content. One prominent area of synthetic media that has gained significant attention is deepfakes, which involve manipulating and creating hyper-realistic imagery (Westerlund, 2019). This phenomenon is referred to as "deepfake," with "deep" denoting the utilization of deep learning neural networks (NN), and "fake" signifying its deviation from the authenticity of the original input (Silva, et al., 2022). They are often portrayed in the media as a "phantom menace", despite their growing relevance and potential in various fields, including marketing (Westerlund, 2019). This growing significance highlights the need to understand the benefits and potential risks of deepfakes, particularly within the field of marketing theory and practice, where their use could revolutionize communication strategies.

Accordingly, researchers worldwide are actively exploring different facets of deepfakes and working towards making significant breakthroughs in the field (Zachary, 2020). The deepfakes research has predominantly focused on improving the algorithms, developing methods to detect them, and analyzing the broader societal implications (Eberl, Kühn & Wolbring, 2022). However, there has been relatively little investigation into the marketing implications of this technology. Despite the attention given to the ethical and technical aspects, the intersection of deepfakes and consumer behavior, particularly in advertising, remains underexplored. One industry where the potential impact of deepfakes could be especially profound is fashion, which relies heavily on visual representation. Yet, how deepfake technology might affect consumer trust, perceptions of authenticity, and behaviors like ad avoidance in fashion advertising have not been thoroughly examined.

Deepfake advertisement is an imminent reality for the future reshaping of the advertising industry (Campbell, Plangger, Sands, Kietzmann, Bates, 2022). As the industry moves away from traditional content creation, which relies on analog and digital tools, it evolves into a more sophisticated territory known as "synthetic advertising." This emerging form of advertising, a subset of manipulated advertising, operates on a highly advanced level (Floridi, 2018; Karnouskos, 2020; Kietzmann et al., 2020).

This research seeks to address the existing gap by understanding how deepfake advertisements can impact and influence consumer behavior, particularly in terms of ad avoidance within the fashion industry. The research will also investigate how consumers respond to deepfake advertisements compared to traditional advertisements, investigate the strategies consumers use to deal with these ads, and evaluate how deepfake technology can impact perceptions of brand trust and authenticity.

Research questions guiding this investigation include:

RQ1: To what extent might highly manipulated advertising lead consumers to avoid ads as an information source?

RQ2: How do perceived verisimilitude, creativity, and relevance influence consumer reactions to deepfake advertisements?

RQ3: Do consumers view deepfake ads as less authentic compared to non-deepfake ones?

The findings from this study have the potential to deepen the understanding of the role of deepfakes in shaping marketing strategies, particularly within the fashion industry. Research will examine the impact of deepfakes on consumer behavior, including ad avoidance and perceptions of authenticity, to provide insights for brands interested in using this technology for marketing. The study aims to offer a balanced perspective on leveraging technological innovation while maintaining consumer trust.

To achieve the research's objectives, this paper is divided into six different sections. Chapter 1 provides an in-depth review of the existing literature on deepfake technology, advertising, and consumer behavior. Chapter 2 establishes the theoretical framework and presents the hypothesis that will guide the investigation. Chapter 3 outlines the research methodology, offering a detailed explanation of the data collection process, sampling techniques, and methods of data analysis employed. Chapter 4 presents and analyses the study's findings. Chapter 5 discusses the broader implications of these findings within the fashion industry. Finally, Chapter 6 concludes by summarizing the key insights delivered from the research, acknowledging its limitations and offering future research recommendations.

Chapter 1: Literature Review

1.1. Fashion Advertisement

The advertising setting of the fashion industry has seen significant advancements, which has encouraged fashion brands to adopt new and innovative approaches to showcase their brand image and effectively connect with their audience (Segal, 2023). This means that fashion brands must create visually appealing, specifically tailored, and captivating ads that can gain an advantage in the severely competitive market (Segal, 2023).

Fashion marketing is a subset of marketing within the fashion industry, encompassing advertising campaigns and promotional initiatives with a particular emphasis on specific customer segments (Bhasin, 2019). This definition captures the utilization of personalized promotional techniques that place customers and potential customers at the center (Easey, 2009).

Given the fundamentally visual character of the fashion industry, it is understandable that fashion advertising depends heavily on visual components (Santaella, Summers, & Belleau, 2012). From the beginning of its existence, fashion advertising has focused on presenting models with little or no text (Phillips & McQuarrie, 2011).

There has been a significant change in the past few years, with more fashion brands committing a large percentage of their advertising budget to online platforms (Strugatz, 2014). Although most research focuses on print and, to a lesser extent, social media advertising, a growing body of innovative studies investigate the intersection of fashion and social media advertising (e.g., Chu, Kamal, & Kim, 2013; Kamal, Chu, & Pedram, 2013; Kim & Ko, 2010, 2012).

1.1.1. Key Theories and Findings in Fashion Advertising

Research on fashion advertising has been approached using various theoretical frameworks to explore key subjects in the field. These include studies on the efficacy of fashion advertising, customer reactions to such advertising, and the identification of critical consumer categories. In this section, the most noted and commonly used ideas used by experts in the field of fashion advertising will be explained, as shown in Table 1.1.

Theory	Explanation	Authors
Involvement	Involvement pertains to the perceived relevance of a focal object based on an individual's needs, values, and interests. Within the fashion context, identified four types of involvement: product, purchase decision, advertising, and consumption.	Zaichkowsky (1985, 1994), O'Cass (2000), Auty and Elliott (1998), Cervellon (2012), Lee and Burns (2014)
Elaboration likelihood mode	Consumers with high involvement are more inclined to process advertisements using the central route to persuasion while those with a lower level of involvement would process information via the peripheral route to persuasion.	Petty et al. (1983), Petty and Cacioppo (1981), Santaella et al. (2012), Santaella, Summers, and Kuttruff (2014), Lee and Burns (2014)
Diffusion innovation theory	Opinion leaders in the fashion industry have significant influence, making them crucial targets for advertising. These leaders go beyond their purchasing power and act as stimulus in the diffusion process.	Summers, 1970), Vernet (2004), Goldsmith et al. (1993), Harben and Kim (2008), Janssen and Paas (2014)
Theory of social comparison	Social comparison theory proposes that humans evaluate themselves by comparison with others. Often studied in relation to the effects of idealized models in fashion advertising.	Festinger's (1954), Richins (1991), Hogg et al. (1999), Kamal et al. (2013)
Theory of narrative transportation	Narrative transportation represents a distinct path to persuasion where an individual is carried away by a story.	Phillips & McQuarrie (2010), Phillips and McQuarrie (2011), Green and Brock's (2000), Barry and Phillips (2016)

Table 1.1: Key and emerging theories in fashion advertising. Adapted from Taylor & Costello (2017)

Although the theories mentioned above are considered traditional, they remain widely relevant and serve as the foundation for much of the contemporary discourse in the field. In recent years, innovative theories have emerged that offer a fresh perspective on the traditional approaches. These new theories represent a novel format that has the potential to revolutionize the field and contribute to the advancement of knowledge in this domain. (Taylor & Costello, 2017).

Topic	Key Findings
Effectiveness Issues	- Although visual imagery is significant in fashion advertising, research indicate that a mix of text and pictures generates the best levels of consumer interest (Phillips & McQuarrie, 2011; Santaella et al., 2012). - Emotion is key in fashion advertising. Emotional vs informative appeals in fashion commercials positively influence customers' sentiments towards the company (Lee & Burns, 2014). - Consumers can interact with fashion commercials in a variety of ways, including acting, feeling, transporting, and immersing themselves (Barry & Phillips, 2016; Phillips & McQuarrie, 2010). This effect has been demonstrated with both male and female samples, suggesting that both genders interpret fashion commercials in the same way (Costello, 2017).
Model Issues	- Very thin models may lead to harmful effects for some women who compare themselves to these idealized portrayals (Borland & Akram, 2007; Dittmar & Howard, 2004; Halliwell & Dittmar, 2004; Hogg et al., 1999; Martin & Xavier, 2010; Murphy & Jackson, 2011). - Although customers do not appreciate extremely underweight or overweight models, fashion advertisements using relatively slim models are the most effective (Aagerup, 2011; Jackson & Ross, 1998; Janssen & Paas, 2014).
Segmentation	Due to their purchasing power, older customers are an important fashion sector (Borland & Akram, 2007). This group prefers models that are close in age to themselves; yet consumers may see cognitive age as more relevant than chronological age (Kozar, 2010; Kozar & Lynn Damhorst, 2008; Zurcher Wray & Nelson Hodges, 2008).

	- Content analyses in Europe in the 1990s discovered that apparel advertisements were more often localized than standardized across nations (Seitz, 1998; Seitz & Johar, 1993), but other research has suggested that various fashion values and lifestyles are shared across cultures and can be used in more global advertisements. - Similar advertising themes might reach fashion consumers in countries with similar values, lifestyles, and fashion leadership attributes (Goldsmith et al., 1993; Ko et al., 2007, 2012; Sarabia-Sanchez et al., 2012; Vernet, 2004).
Social Media Advertising	Social media may help fashion businesses, particularly premium ones, create brand equity and increase buy intent (Chu, Kamal, & Kim, 2013; Kim & Ko, 2010, 2012).
Controversial Advertising	Controversial fashion advertising (commercial with sex, violence, or strong political content, for example) appears to be ineffectual and results in poor consumer ratings of the ad (Andersson, Hedelin, Nilsson, & Welander, 2004; Harben & Kim, 2008). Brands that use these strategies risk obtaining poor news and perhaps losing business partners due to contentious advertisements (Andersson et al., 2004; Harben & Kim, 2008).

Table 1.2: Key findings in fashion advertising. Adapted from Taylor & Costello (2017)

The current setting of fashion advertising heavily relies on visual storytelling, emotional engagement, and carefully crafted content that resonates with diverse consumer segments. Research has shown that a mix of visual imagery and text generates the highest consumer interest, and emotional appeals are key to positively influencing brand perception (Phillips & McQuarrie, 2011; Lee & Burns, 2014). Additionally, models and representation in advertising have a significant impact on how consumers interact with and respond to fashion brands, with issues of body image and inclusivity playing a critical role (Borland & Akram, 2007).

As we move toward the future of fashion advertising, these foundational theories and findings will be enhanced by emerging technologies like AI, which promises to further revolutionize how brands engage with consumers, personalize content, and optimize their strategies for success (Rathore, 2019).

1.1.2. The Future of Fashion Advertisement

The fashion advertising field is a dynamical and essential aspect of the fashion industry, playing a pivotal role in its success. With the rise of AI as a disruptive technology, new opportunities have emerged for fashion marketers and advertisers to explore (Rathore, 2019). AI-powered technologies can optimize marketing strategies, customize customer experiences, and enhance overall satisfaction. Integrating AI technology into the fashion industry is expected to become increasingly common in the coming years. (Rathore, 2019).

The fashion industry's adoption of AI is not merely a matter of operational convenience, but rather a new approach that mandates a paradigm shift (Rathore, 2019).

In recent years, advertising has undergone a significant turn from content created and altered using only analog and digital tools to what is now referred to as synthetic advertising (Campbell, Plangger, sands & Kietzmann, 2022). Synthetic advertising is a highly advanced form of manipulated advertising generated or edited through artificial and automatic production and modification of data (Campbell, Plangger, sands & Kietzmann, 2022). This process relies on AI algorithms such as deepfakes and generative adversarial networks (GANs) to automatically generate content that depicts a highly convincing yet fabricated version of reality (Floridi, 2018; Karnouskos, 2020; Kietzmann et al., 2020).

In the contemporary setting, specialized companies are particularly involved in the creation of AI-driven promotional models employing deepfake technology and GANs. These cutting-edge technologies enable the customization of the physical attributes of fashion models and the outfits they wear to create promotional campaigns that do not require the use of actual models (Whittaker, Kietzmann, Kietzmann & Dabirian, 2020).

Deepfakes are among the most used AI tools in synthetic advertising. They employ machine-learning algorithms to generate realistic images, videos, and audio that can be applied to manipulate the audience's perceptions. While synthetic advertising provides new opportunities for businesses to create highly engaging and persuasive content, it also raises concerns about the possible misuse of AI (Whittaker, Kietzmann, Kietzmann & Dabirian, 2020).

The emerging partnership between the fashion industry and AI algorithms holds significant potential for achieving sustained progress and improvements in the years to come.

1.2. Evolution of Advertisement Manipulation

To shape brand perceptions, advertisers commonly employ tactics to influence how consumers perceive their brands, producing content that appeals to them both emotionally and logically (Pawle & Cooper, 2006). In that sense, ad manipulation techniques are employed throughout the advertising process, from pre-production (wardrobe and makeup choices to specialized lighting and camera lenses used in production) to post-production (image or recording retouching) (Rust & Oliver, 1994). These findings imply that advertisements have historically depicted an artistic version of reality.

Nonetheless, the evolution of this manipulation has undergone significant changes over time, a summary of which is presented in Table 1.3.

Generation	1.0. Analog	2.0. Digital	3.0. Synthetic
Sample Tools	Makeup, lights, camera lenses, physical editing	Computer-generated imagery (CGI), photoshop, Instagram filters, etc.	Deepfakes, Generative Adversarial Networks (GANs)
Agency	Human activities (manual)	Human-computer interactions (assisted)	Artificial intelligence (AI) and machine learning (ML) techniques (automated)
Targeting Channel	Generic, mass focused TV, radio, print	Micro and macro segments, TV, print, online	Hyperpersonalized, online

Table 1.3: Generations of Manipulation in Advertising. Adapted from Campbell, Plangger, Sands & Kietzmann (2022)

Numerous instances showcase the gradual improvement of advertising manipulation techniques, which can be classified into three eras: 1.0 Analog; 2.0 Digital; and 3.0 Synthetic. It is crucial to understand that these eras do not signify a linear progression, in which one technique replaces another. Instead, they represent evolutionary stages that frequently coexist and combine to achieve groundbreaking outcomes (Campbell, Plangger, Sands, Kietzmann, 2022).

1.2.1. Generation 1.0. Analog Manipulation

Analog manipulation entails employing numerous tools and techniques to improve the quality of content in photographs, videos, or audios. Specialized artists used physical tools like makeup, lighting equipment, airbrushes, paintbrushes, and dyes to refine the content and eliminate any flaws or imperfections (McDonald & Scott, 2007).

During the pre-production phase, analog manipulation is used to generate ideal production conditions, whereas post-production involves the cutting of audio or video on magnetic tape or the retouching of photographic negatives using paint, ink, or airbrushing techniques. However, analog manipulation can be time-consuming and labor-intensive, posing some challenges and limitations (Cambell et al., 2022).

1.2.2. Generation 2.0. Digital Manipulation

The use of digital manipulation arises from the practice of computer-based tools to modify and create content. With recent editing software, creating and editing has become more accessible to advertising professionals (Dewey, 2015).

The first digital filters allowed for the retouching, augmentation, or alteration of pictures and videos, and similar functionalities are now integrated into popular smartphone apps like Instagram and TikTok. More sophisticated digital tools, including computer animation and computer-generated imagery (CGI), enable advertisers to go beyond mere editing and use complex techniques, such as incorporating new digital elements into advertisements. These digital tools accelerate both the quality and quantity of feasible content manipulations (Campbell et al., 2022).

Although analog tools remain relevant in advertising, many analog techniques have now been digitalized. This has resulted in a transition away from primarily depending on human artistry and towards the employment of specialized computer programs alongside with human operators, resulting in a more collaborative approach (Campbell et al., 2022).

1.2.3. Generation 3.0. Synthetic Manipulation

Synthetic manipulation involves the autonomous alteration or generation of content through the utilization of AI algorithms. Such algorithms enable the creation of content in a synthetic manner, as demonstrated by GANs, and allow for the seamless editing of existing content through technologies such as deepfakes (Floridi, 2018; Karmnouskos, 2020; Kietzman et al., 2020). In the case of deepfakes, the process entails the substitution of attributes, such as facial features, voice, skin tone, gender, and fashion details from a targeted source (Floridi, 2018; Karmnouskos, 2020; Kietzman et al., 2020). GAN-generated deepfakes can produce entirely original and synthetic media (Whittaker et al., 2020). These applications include non-existing models and fashion designs, photorealistic anime characters, portraits, album covers, facial ageing or deageing transformations, gender-swapping, image generation from textual descriptions and many more (Antipov, Baccouche, & Dugelay, 2017; Reed, et al., 2016).

The usage of synthetic manipulation techniques presents new prospects for generating advertisements at significantly lower costs compared to traditional methods. However, advertisers are still in the nascent stages of comprehending the impact of this advanced manipulation on the efficacy of their advertisements (Campbell et al., 2022).

While ad manipulation can be customized to a certain degree, synthetic manipulation can achieve a level of hyper-personalization and individualization that is unmatched. This means that all content can be dynamically tailored in real-time to meet the needs and preferences of each customer, leveraging data gleaned from sources

such as social media interactions, retail sensors, or loyalty programs (Campbell et al., 2020; Kietzmann et al., 2020; Schelenz, Segal & Gal, 2020).

The era of traditional analog and digital content creation and modification is undergoing a significant transition and transformation, with the rapid rise of synthetic media. As these technologies continue to advance, it is expected that a broader range of applications will continue to emerge, challenging the prevailing perceptions of reality and authenticity (Whittaker, Kietzmann, Kietzmann & Dabrian, 2020).

1.3. Introduction to Deepfake Technology

The term "deepfake" is an amalgamation of the lexemes "deep" and "fake". The term "deep" inside the description of deepfake underlines the significance of deep learning, while the term "fake" underlines the nature of the output as a simulated reality (Kiliç & Kahraman, 2023).

Deepfake is a form of synthetic media that challenges the line between authenticity and realism. This technology is made possible by the junction of AI technology, Machine Learning, Deep Learning and Neural Networking which is a combination of algorithms such as GANs and Autoencoders. This enables the generation of numerous content forms, such as merging, replacing, overlaying, or combining different elements. The term 'deepfake' refers to the domain within a larger field of synthetic media. While capable of producing convincing yet fake videos, images, and audio, this technology can be employed for both harmless and malicious purposes (Bateman, 2020; Maras & Alexandrou, 2019).

Since its origin, researchers have actively studied deepfake technology, attempting to identify and categorize its various types. By doing so, we can gain a better understanding and analysis of this complex phenomenon. As part of this effort, researchers created a taxonomy of deepfakes that categorizes deepfake-generated media into two primary domains: facial (image and video) and speech. This taxonomy of deepfakes is depicted in Figure 1.1.

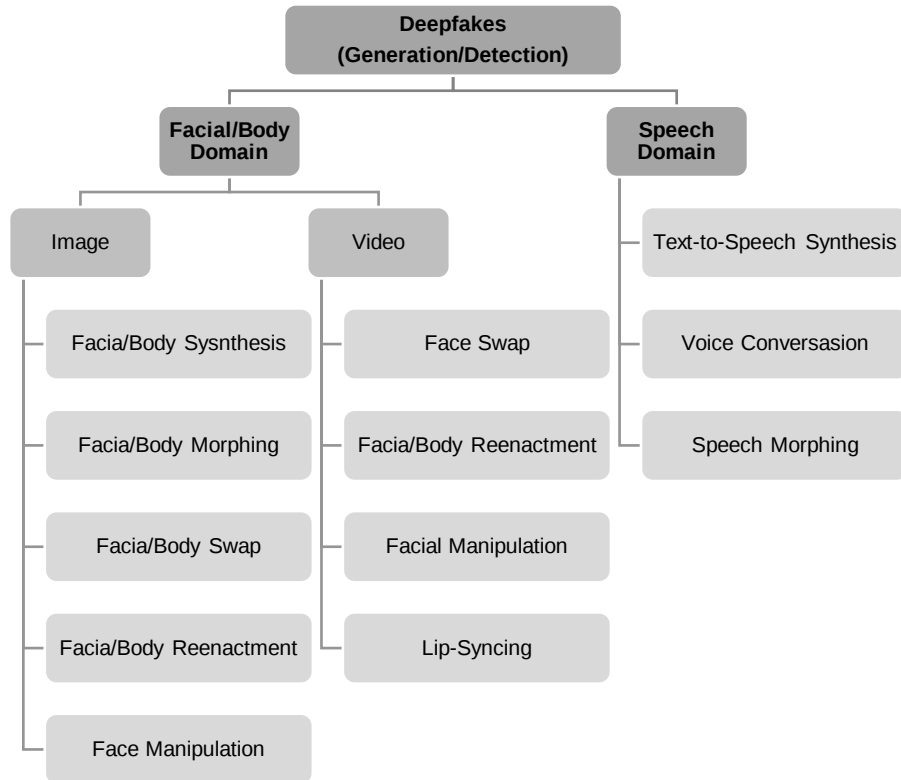


Figure 1.1: Visualization of the Taxonomy of Deepfakes. Adapted from Masood, Nawaz, Malik, Javed, Irtaza & Malik (2022) And Firc, Malika & Hanáček (2023)

Formerly, Farid et al. (2019, pp.4-6) examined the diverse manifestations assumed by deepfakes and effectively classified them into four distinct categories. First, they identified "face replacement" or "face swapping," which entails substituting the visage of one individual (the source) with that of another (the target). Second, the authors explored "face re-enactment," which involves the manipulation of facial attributes such as mouth movement and eye expressions, among others. Third, they investigated "face generation," which encompasses the creation of entirely novel faces by utilizing the extensive capabilities offered by Generative Adversarial Networks. Finally, the researchers researched into "speech synthesis," the process of modifying an individual's discourse in terms of cadence and intonation or generating an entirely original speech.

Recent studies have made notable strides in advancing our comprehension of deepfakes, leading to more sophisticated classification systems. As shown in Figure 1.1, the Taxonomy of Deepfakes has been visualized to provide an insightful depiction of this refined categorization. This improved taxonomy highlights two main areas in which deepfakes are applied: the "facial domain" and "speech domain". The "facial/body domain" is subcategorized into various techniques, including focal methods like:

1. **Face/Body Synthesis (image):** This technique involves a process of synthesizing non-existing faces/bodies based on learned high-level attributes, such as pose or identity, by generating images that do not exist (Karras, Laine, Aittala, Hellsten, Lehtinen & Aila, 2020).
2. **Face/Body Morphing (image):** This technique involves a process characterized by its seamless transition and is frequently employed to portray the metamorphosis of one individual into another (Ferrara, Franco & Maltoni, 2014).
3. **Face/Body Swap (image and video):** This technique involves the transfer of a face/body from a source photograph to a face/body in a target photograph. The desired outcome of this process is to achieve a realistic and unedited appearance (Firc, Malinka & Hanacek, 2023).
4. **Facial/Body Reenactment (image and video):** This technique involves a photo/video realistic facial/body reanimation employed to animate the face/body of a target video using expressions from a source actor. Originally proposed to enhance the visual component of a digital assistant scenario, these methods play a crucial role in generating synchronized audio and visually realistic human motions (Thies, Elgharib, Tewari, Theobalt & Nießner, 2020).
5. **Face Manipulation (image and video):** This technique involves modifying a specific facial area in an image or video while keeping the target's identity intact. This method allows for the addition or removal of features such as facial hair, glasses, and other attributes, as well as the alteration or transfer of expressions, lighting, or pose within the target's head region (Huijstee, Boheemen, Nierlin, Jahnel, Karaboga, Martin, Kool & Gerritsen, 2021).
6. **Lip-Syncing (video):** This technique involves coherence between the visual representation of mouth movements and the corresponding audio content in the synthesized video (Suwajanakorn, Seitz & Kemelmacher-Shlizerman., 2017).

On the other hand, focusing on the “speech domain” and its subdivision within this taxonomy, it encompasses focal techniques such as:

1. **Text-to-Speech Synthesis:** This technique transforms written text into spoken words. The aim is to generate synthesized speech that is not only highly intelligible, but also perceptually indistinguishable from human speech (Taylor, 2009), achieving a level of speech synthesis that is both natural sounding and comparable to human-produced speech (Tabet & Mohamed, 2011).
2. **Voice Conversation:** This technique involves modifying the vocal characteristics of a given speech from a source speaker to resemble those of a target speaker

(Machado & Queiroz, 2010 and Qian, Zhang, Chang, Cox & Hasegawa-Johnson, 2019). Unlike text-to-speech (TTS) systems, voice conversations don't rely on written text and can alter specific aspects of speech, such as tone, cadence, or pitch (Qian, Zhang, Chang, Cox & Hasegawa-Johnson, 2020).

3. **Speech Morphing:** This technique enables seamless transformation from one signal to another, resulting in the creation of a new signal with an intermediate timbre (Cano, Loscos, Bonada, Boer & Serra, 2000). The goal is to achieve a cohesive and harmonious output by smoothly merging the characteristics of the original signals (Pfitzinger, 2004).

1.3.1. Deepfakes – an Illusion or a New Hope for Marketing

As previously mentioned, deepfakes are a form of synthetic media that uses sophisticated technology to generate visual content that appears genuine, despite being fabricated. By utilizing AI to manipulate elements such as facial expressions, voice, and physical characteristics of an individual, deepfakes can create a sense of realism that can easily mislead people into believing that they are authentic. This distinguishes it from other types of fake imagery (Whittaker et al., 2020). However, it is essential to note that deepfake technology has potential applications beyond manipulation and deceit. As a result, researchers are increasingly exploring its possibilities, and numerous academic projects are currently underway to deepen our understanding of this technology (Killiç & Kahraman, 2023).

According to Kaplan and Haenlein (2019), deepfakes demonstrate greater emphasis on cognitive intelligence than other advanced forms of intelligence, such as emotional or social intelligence. Additionally, Letheren et al. (2019), The Proactive-Interactive-Passive (PIP) concept suggests that deepfakes exhibit proactive behavior by independently utilizing AI abilities to create synthetic representations based on input data. It is worth noting that, while human intervention is necessary in the creation of deepfakes, the AI-powered assistant acts on behalf of individuals (Letheren et al., 2019).

The phenomenon of deepfakes is increasing in sophistication and will eventually be imperceptible to the untrained eye (Maras & Alexandrou, 2019). The two main factors driving their spread through social media are their increasing accessibility and believability, as deepfakes become easier to produce but also more difficult to distinguish from authentic media due to their increasing sophistication (Kietzmann et al., 2020). With that, the uses of deepfakes can be broadly categorized into two categories, namely, malicious, and beneficial. The malicious uses of deepfakes include the creation of fake

news, fraud, identity theft, and others. Conversely, beneficial applications of deepfakes find their place in entertainment, education, healthcare, and others. Although the beneficial applications of deepfakes have proven to be revolutionary, malicious uses pose a significant threat to society. (Kwork & Koh, 2020). According to Kwork and Koh (2020), the use of deepfake technology has resulted in more harm than good.

Deepfakes have been described as a potentially dangerous phenomenon, exerting a negative influence on both consumers and businesses. The abundance of manipulated media raises the risk of impairing individuals' perception of reality and their sensitivity towards it. Such a scenario may have far-reaching implications, given the potential to distort public opinion and undermine trust in essential institutions (Elitaş, 2022). The statement suggested is grounded on the growing sophistication of deepfake technology, coupled with a significant reduction in the barriers that impede their creation (Whittaker, Letheren & Mulcahy, 2021). The intersection of these characteristics with the present setting of digitally documented life is of special significance. These malicious applications include but are not limited to, extortion, intimidation, sabotage, harassment, defamation, revenge porn distribution, identity theft, and cyberbullying (Kietzmann, Lee, McCarthy, & Kietzmann, 2020; Chesney & Citron, 2019). With the emergence of Deepfakes and GANs, deceitful media is becoming a significant threat to the reliability of online sources (Weikmann & Lecheler, 2023). In the absence of effective detection technology in verifying the authenticity of visual and auditory media, even genuine images, videos, or audio recordings can be discredited under such circumstances (Whittaker, Kietzmann, Kietzmann & Dabirian, 2019).

In the present context, several instances of deepfake visual content generated by AI have become viral on social media platforms. The technology used to create these fake videos and images is so advanced that it has managed to deceive millions of people worldwide. In Annex A it can be observed a few recent examples of deepfake imagery that have gained widespread attention.

From an optimistic perspective, deepfake technology has the potential to serve as a powerful tool for creating compelling and relatable content, thereby fostering stronger connections between brands and consumers (Whittaker, Letheren & Mulcahy, 2021). Furthermore, the capabilities of deepfake technology can assist the implementation of highly personalized advertising strategies, allowing for the detailed customization of commercial messaging to different target demographics within a specific campaign. This level of tailored personalization holds promise for driving increased sales and fortifying brand reputation. However, it is imperative to approach

this advancement with caution to prevent the potential emergence of negative consequences, such as intensified consumer vigilance, concerns regarding privacy intrusion, and increased susceptibility (Campbell, 2023).

In a broader sense, deepfake technology provides endless possibilities for personalized media creation. More specifically, deepfake technology allows easing of language barriers, thereby enhancing the cross-cultural video content distribution that would typically require supplementary subtitles. Furthermore, deepfakes enable those who have lost their voice due to medical conditions to regain their ability to communicate. Leveraging similar deep learning principles employed in the creation of video deepfakes, this application extends the scope of deepfake utility (Whittaker, Kietzmann, Kietzmann & Dabrian, 2020).

Another example of the limitless potential of deepfake technology is its application in the film industry. This technology offers substantial advantages, especially in the context of de-aging actors, this process is comparable to the expenses associated with Computer-Generated Imagery (CGI) effects. This technological innovation allows for more cost-effective and realistic rejuvenation of actors (Whittaker, Kietzmann, Kietzmann & Dabrian, 2020). Altered advertisements may offer enhanced entertainment, excitement, or engagement compared to non-synthetic advertisements developed within equivalent budgetary constraints. It is anticipated that persuasive advertising will be deemed acceptable if it delivers superior overall value (Campbell, Plangger, Sands & Kietzmann, 2022).

An emerging trend in the industry involves the specialized application of GANs by certain companies to construct AI-driven promotional models. This technology enables businesses to intricately tailor the physical attributes of fashion models and their attire for the development of personalized promotional campaigns, eliminating the need for human models (Whittaker, Kietzmann, Kietzmann, & Dabrian, 2020).

A substantial growth in the amount of deepfake variants has occurred in the current scene. This section seeks to clarify their possible uses. Kietzmann et al. (2019) have methodically catalogued these differences and thoroughly explained their commercial applications, as seen in Table 1.4. The key takeaways from their study are briefly captured in the following language and tabulated for clarity.

Type	Description	Business application
Image deepfakes	Face/Body-Swapping	Consumers can try on cosmetics, eyeglasses, hairstyles, or clothes virtually
Video and Image Deepfake	Face/Body Morphing	Video game players can insert their faces onto their favorite characters
	Face/Body Reenactment	In a video presentation, business leaders and athletes can easily hide any physical issues they might have.
Video Deepfake	Face Swapping	Face-swapped video can be used to put the leading actor's face onto the body of a stunt double for more realistic-looking action shots in movies
Video and Audio Deepfake	Lip-Syncing	Ads and instructional videos can be 'translated' into other languages using the same voice used in the original recording
Audio Deepfakes	Voice-Swapping	The voice of an audiobook narration can sound younger, older, male, or female and with different dialects or accents to take on different characters
	Text-to-Speech	Misspoken words or script changes in a voiceover can be quickly replaced without the need for re-recording.

Table 1.4: Types and examples of deepfakes business applications. Adapted from Kietzmann, Lee, McCarthy & Kietzmann (2020)

This innovative technology can revolutionize the marketing world, with boundless potential for personalized media creation.

1.3.1.1. Deepfake in Advertising and Marketing

The advent of deepfake technology has opened doors for innovative marketing strategies that can leverage the popularity of celebrities to enhance the outreach and recall of advertisements. By seamlessly integrating the image and voice of renowned individuals into ads, advertisers can create highly engaging and persuasive promotional content that resonates with their target audience. This technology has the potential to revolutionize the advertising industry by enabling brands to tap into the aspirational value and emotional appeal of celebrities and create campaigns that are not only aesthetically pleasing but also highly effective in driving consumer behavior. Thus, deepfake technology presents a unique opportunity for businesses to create impactful marketing campaigns that stand out in a crowded marketplace and leave a lasting impression on their customers (Kiliç & Kahraman, 2023).

On June 6th, 2021, Balenciaga presented its Spring/Summer 2022 virtual fashion show, named "Clones." The exhibition showcased a single model, Eliza Douglas, and investigated into society's perception of technological reality in the post-digital era. The "clones" in the show were created using a range of techniques, including deepfake, real-time game engine, and traditional visual effects. Notably, the brand's innovative use of deepfake technology in its promotional efforts emphasizes its dedication to adopting cutting-edge marketing strategies to enhance brand recognition and drive customer engagement.

In March 2023, Balenciaga, the luxury fashion brand, experienced a significant breakthrough. Videos featuring the brand, created using innovative deepfake technology, garnered widespread attention on the internet. Balenciaga strategically harnessed the immense popularity of iconic franchises such as Harry Potter, Breaking Bad, and Star Wars, among others, to effectively promote its brand, as shown in Annex A. According to Jennings (2023), utilizing YouTube ads to promote Balenciaga holds considerable potential for generating significant revenue.

Deepfake technology has been employed in various instances to produce highly convincing and manipulated content. One such instance involved digitally dressing the Pope in a Balenciaga ensemble, as shown in Annex A. This use of technology sparked extensive conversations and debates, and the fabricated content effectively deceived millions of viewers.

The incorporation of deepfake technology into advertising is still being studied. One potential benefit of this technology is the ability to create compelling and stimulating media through AI (Jennings, 2023). Deepfake technology has the potential to shape the future of advertising. By making deepfake advertisements more authentic and engaging, we could improve the effectiveness of marketing strategies. This could result in increased brand recognition, higher sales, and greater customer satisfaction.

1.4. Trust and Perceptions of Authenticity

With the increasing use of deepfake technology in the advertising and fashion industry, worries about authenticity and trust have increased. As brands adopt synthetic media, the potential for perceived deception intensifies, which could lead to a further decrease in trust, particularly in industries that depend heavily on visual and emotional engagement, such as fashion.

Building brand trust involves a combination of good intentions and strong capabilities (Ballester, 2004; Ballester & Alemán, 2001). Brands gain trust by demonstrating these characteristics through organizational values that reflect transparency, responsibility, and ethical behavior (Morhart et al., 2015; Schallehn, Burmann, & Riley, 2014). This reflects a larger definition of brand authenticity, which emphasizes consistency between a brand's communicated values and its actions (Morhart et al., 2015; Schallehn, Burmann, & Riley, 2014)

The concept of authenticity is connected to brand trust, indicating that authenticity plays a significant role in enhancing brand trust. (Eggers et al., 2013; Hon & Grunig, 1999; Napoli, Dickinson, Beverland, & Farrelly, 2014; Schallehn et al., 2014). Authentic brands foster trust establishing credibility and reliability over time (Molleda, 2010; Molleda & Jain, 2013; Morhart et al., 2015).

In an era of increased consumer skepticism, authenticity solves declining trust (Bruhn et al., 2012). When consumers perceive a brand as authentic, they are more likely to form trust-based relationships, leading to greater brand loyalty and advocacy. This trust is the foundation for successful brand-consumer relationships and is essential for cultivating long-term relationships with customers (Ballester & Alemán, 2001; Fournier, 1998; Morgan & Hunt, 1994).

However, as deepfake technology becomes more dominant in advertising, it confuses the dynamics of trust and authenticity. While deepfake technology presents unique challenges, it also emphasizes the need to reinforce brand authenticity.

1.5. Ad Avoidance

Over the past two decades, the advertising industry has undergone significant transformation driven by advancements in digital technologies, resulting in the establishment of a digital market (Sharma et al., 2022). This new advertising approach is continually evolving and has become universal in modern marketing strategies (Dodoo & Wen, 2019).

Consequently, the average consumer exposure to brand communication has reached an unparalleled level. This tendency has been primarily attributed to digitalization, which has revolutionized the advertising industry by providing efficient and effective marketing strategies, thus increasing the industry's overall growth (Lee & Cho, 2020; Tudoran, 2019).

Nevertheless, the increasing popularity of digital advertising has led to an increase in user exposure to excessive advertising. Such exposure has been demonstrated to evoke unfavorable sentiments in users and can lead them to avoid such messages (Ferreira & Barbosa, 2017; Sharma et al., 2022).

In the digital field, ad avoidance refers to efforts attempted to limit or eliminate exposure to digital advertising (Kelly et al., 2020). With changing media consumption

habits, people are turning towards ad-blocking tools such as Adblock or opting for paid platform services such as YouTube Premium to reduce or eliminate the presence of ads they encounter (Edelman, 2020).

Prior research has explored into the various factors that influence ad avoidance, including irritation, intrusiveness, and skepticism. However, the underlying antecedents of this phenomenon have yet to be established. Identifying the factors that provoke ad avoidance is critical (Çelik, Çam & Koseoglu, 2022).

1.5.1. Ad Avoidance in the Age of Deepfake Advertisement

Deepfakes have the potential to be influential and engaging tools for sharing information; however, they also present a significant risk of undermining public trust in factual content and in organizations and brands (Chesney and Citron, 2019). This technological advancement, combined with GANs, represents a progression in the dissemination of misinformation and fabricated fake news reports, further emphasizing the importance of reinforcing trust in online information (Whittaker, Kietzmann, Kietzmann & Dabrian, 2020).

Chapter 2: Conceptual Model and Hypothesis of Investigation

This chapter is dedicated to presenting the conceptual framework model and hypotheses central to this investigation's conclusions, providing a clear and concise understanding of the theoretical framework and the underlying hypotheses, which will guide the subsequent analysis and discussion.

The research will focus on analyzing consumer responses to ad manipulation and how it can lead to ad avoidance, a phenomenon that is becoming increasingly predominant in today's digital setting (Celik, Çam, & Köseoglu, 2022). The framework shown in Figure 2.1 is built on the existing framework proposed by Campbell, Plangger, Sands and Kietzmann (2022) a framework for consumer response to manipulated advertising. The framework begins with “manipulation sophistication” in synthetic advertisements. This relates to the degree of manipulation elements incorporated into these advertisements, which can distort reality and influence the perception of the image's authenticity. The increased sophistication of manipulation techniques can influence both verisimilitude (the extent to which an advertisement is perceived as genuine and truthful) and the perceived creativity intrinsic to an advertising (Campbell, Plangger, Snads & Kietzmann, 2022). Moreover, the increased modification sophistication can influence the perceived relevance since it reflects the users' perceptions of utility and usefulness (Dodoo & Wen, 2021).

Both the elements of verisimilitude and creativity hold significant potential to amplify the persuasiveness of an advertisement. However, they simultaneously possess the capacity to trigger an increased awareness of the advertisement's falsity, potentially leading to a decline in its persuasive impact and possibly resulting in ad avoidance (Campbell, Plangger, Snads & Kietzmann, 2022).

Moreover, it is relevant to consider the level of trust in the fashion brand as an additional factor that may exercise a direct influence, potentially shaping the phenomenon of ad avoidance among consumers.

The hypotheses that follow will undergo testing through the methodology presented in Chapter 3. The objective is to comprehend the correlation between the levels of manipulation in synthetic advertising and ad avoidance and to examine if the overall level of manipulation sophistication impacts the behaviors of ad avoidance.

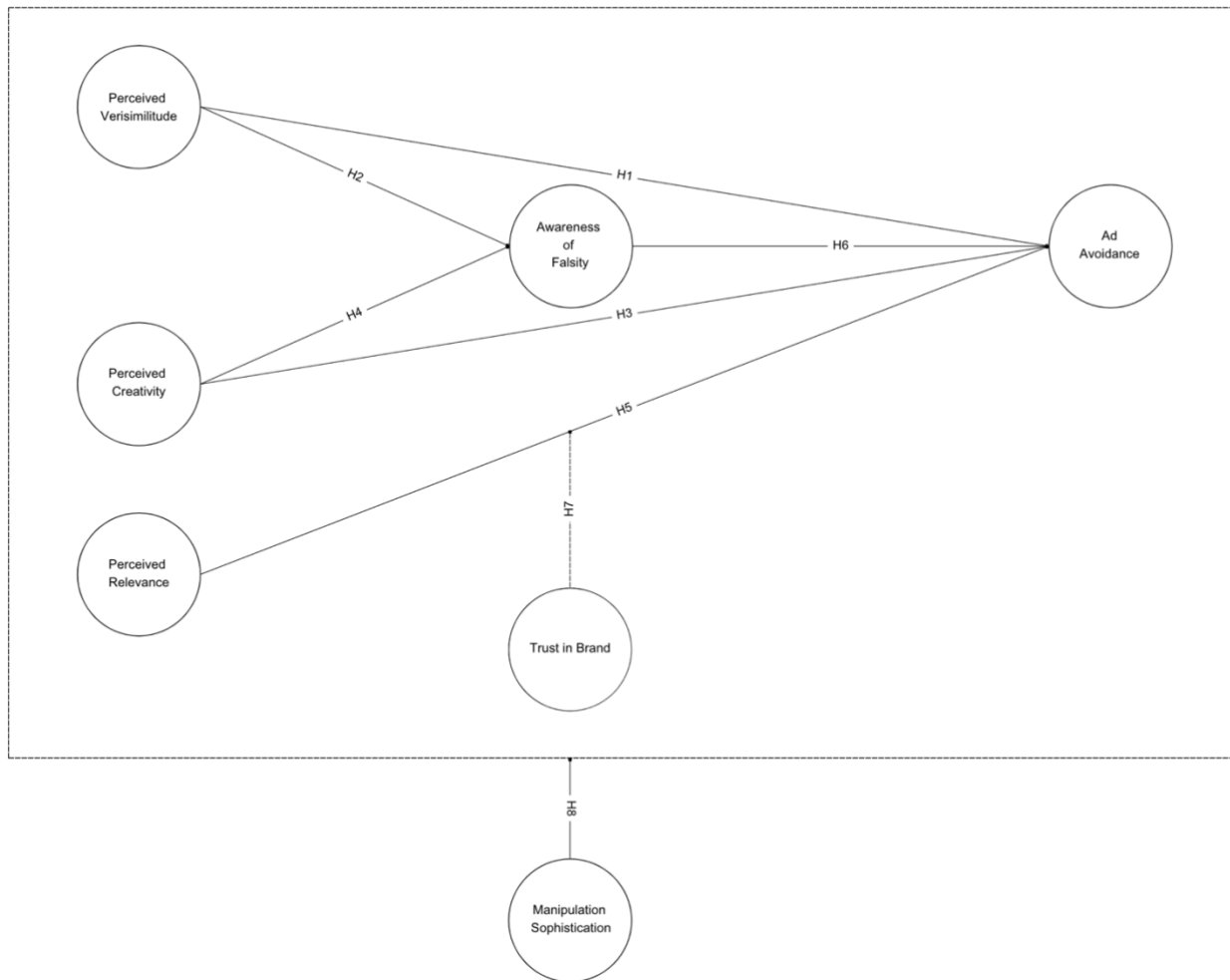


Figure 2.1: Conceptual Framework Model, Based on the Theoretical Model “A framework for consumer response to manipulated advertising” (Campbell, Plangger, Sands & Kietzmann, 2022)

Effects of Increased Verisimilitude

The concept of verisimilitude, which indicates a resemblance to truth (Fine, 2019), is crucial in advertising, particularly in the settings of product placement and narrative. The verisimilitude of an advertisement is determined by the consumer's judgement of its authenticity and commitment to reality (Campbell, Plangger, sands & Kietzmann, 2022).

Verisimilitude refers to the degree to which a manipulated advertisement delivers a sense of honesty, realism, or persuasiveness to the customer. When viewers are unable to tell whether the content has been manipulated, they perceive it as an authentic reflection of reality, increasing the effectiveness of ad manipulation. A higher perceived level of verisimilitude leads consumers to engage with the ad as if it were a true depiction, triggering established persuasive mechanisms that drive advertising engagement (Campbell, Plangger, Sands & Kietzmann, 2022).

Manipulated advertising with a higher degree of perceived verisimilitude is likely to persuade customers without hesitation (Campbell, Plangger, Sands & Kietzmann, 2022).

Thus, it is hypothesized:

H1: Greater perceived verisimilitude of an advertisement decreases ad avoidance.

When modern technologies generate a synthetic reality that is nearly indistinguishable from what's real, detecting manipulation becomes less likely. In certain circumstances, a customer's capacity to recognize ad manipulation is weakened, resulting in a decrease in awareness of falsity, or the degree to which a consumer believes an ad to be false. (Campbell, Plangger, Sands & Kietzmann, 2022).

Thus, it is hypothesized:

H2: Greater perceived verisimilitude of an advertisement decreases awareness of falsity, which in turn decreases ad avoidance.

Effects of Increased Creativity

Creativity in advertising is characterized by its freshness, unexpectedness, and newness (Kim, Han, & Yoon, 2010; Koslow, Sasser, & Riordan, 2003; Sheinin, Varki, & Ashley, 2011; Smith et al., 2007).

Researchers in marketing argue that increased creativity in advertisements helps to overcome consumer barriers, captures attention effectively, stimulates favorable responses, and strengthens attitudes towards the brand (Marra, 1990; Ogilvy, 1983; Rosengren, Dahlén, & Modig, 2013; Zinkhan, 1993). A widely held belief in the advertising industry is that creativity is a necessary component for ad effectiveness (Kover, Stephen, & Goldberg, 1995), with some marketers explicitly linking ad creativity to its effectiveness (Kover, 1995).

Increased perceptions of creativity are expected to make ads more successful (Campbell, Plangger, Sands & Kietzmann, 2022).

Thus, it is hypothesized:

H3: Greater perceived creativity of an advertisement decreases ad avoidance.

While it is usually true that deceptive advertising elicits negative responses, there are exceptions to this rule. Consumers may be willing to ignore manipulated advertising featuring high levels of creativity because of the inherent value obtained from such originality (Campbell, Plangger, Sands & Kietzmann, 2022). Altered advertisements may provide more amusement, excitement, or engagement than non-synthetic advertisements generated under the same cost limitations.

While increased creativity in advertising typically produces positive results, there is a subtle component in which excessive creativity may raise consumer awareness of an advertisement's falsity (Campbell, Plangger, Sands & Kietzmann, 2022).

Thus, it is hypothesized:

H4: Greater perceived creativity of an advertisement decreases awareness of ad falsity, which in turn decreases ad avoidance.

Effects of Increased Relevance

Relevance encompasses not only the appropriateness or alignment of an advertisement with a brand's strategy and positioning but, crucially, its utility and pertinence to the needs and preferences of consumers (Ang, Lee, & Leong, 2007; El-Murad & West, 2004). Relevant advertisements demonstrate agreement with a brand's strategic objectives and, more importantly, provide usefulness and pertinence to customers. Ad customization is a practice that commonly improves persuasive efficacy in influencing or reinforcing attitudes, intentions, and behaviors (Aguirre et al., 2015; Mukherjee, Smith, & Turri, 2018; Tong, Luo, & Xu, 2020).

Dodoo & Wen (2019) and Kelly et al. (2020) identify ad relevancy as a critical antecedent. Previous studies have shown that an increase in ad relevance correlates with a decrease in ad avoidance (Brinson & Britt, 2021; Dodoo & Wen, 2021; Jung, 2017; Kelly et al., 2010; Li et al., 2020). Increased perceptions of relevance are expected to make ads more successful (Campbell, Plangger, Sands & Kietzmann, 2022).

Thus, it is hypothesized:

H5: Greater perceived relevance of an advertisement decreases ad avoidance.

Effects of Increased Awareness on Ad Falsity

There are various reasons for consumers' typically negative reactions when they are aware of an ad's falsity. Consumers are more sensitive to disinformation in ads as a method of self-protection or dealing with persuasion attempts (Friestad & Wright, 1994).

The significance of authenticity is emphasized in domains like brand expansions (Spiggle, Nguyen, & Caravella, 2012) and social media (Audrezet, De Kerviler, & Guidry Moulard, 2020). When consumers believe advertisements to be fake or unauthentic, their interest in the offer decreases (Spielmann and Orth 2020).

Thus, it is hypothesized:

H6: Greater awareness of ad falsity increases ad avoidance.

Effects of Trust in Brand on Ad Falsity

Brand trust entails a readiness to accept risks while relying on the brand's promise of value. It is characterized by feelings of confidence and security. Moreover, it cannot exist without the possibility of error. When brand trust is established, consumers are more likely to engage with the brand's message, even when they are confronted with elements of uncertainty or potential misinformation (Ballester, 2011).

As a result, greater trust in a certain fashion brand can enhance perceived relevance and reduce the negative impact of ad avoidance.

Thus, it is hypothesized:

H7 (Moderator): Greater trust in a particular fashion brand decreases the negative effect of ad avoidance.

Effects of Increased Manipulation

Campbell, Plangger, Sands & Kietzmann (2022), refer to the term "manipulation sophistication" as the enhancement and refinement that comes from the process of generating or altering content within the context of advertising. This complexity can be achieved using various approaches, in this case, synthetic processing.

Based on existing research by Friestad and Wright (1994), it has been found that consumers might respond in varied ways to an advertisement if they are conscious of the fact that it has been altered or manipulated. This indicates that consumer awareness of manipulation in advertising can influence their reactions and responses.

This hypothesis aims to investigate whether the use of deepfake visual stimuli or non-deepfake stimuli, meaning different levels of manipulation, affects ad avoidance and consumer behavior in the context of deepfake advertising.

Thus, it is hypothesized:

H8: Greater manipulation sophistication of an advertisement decreases ad avoidance.

Chapter 3: Research Methodology

This research aims to comprehensively analyze the potential impact of deepfake advertisements on consumer behavior and attitudes, focusing particularly on ad avoidance in the fashion industry. By shedding light on this emerging trend and its implications for the future of advertising, we hope to provide valuable insights for industry professionals. Our findings from previous chapters inform this study's exploration of the subject.

The primary objective of this chapter is to elucidate the research methodology employed to test the hypothesis theorized in Chapter 2. To gain a more profound comprehension of the subject, a quantitative research methodology was adopted. This involved the collection of data from a broader sample and its analysis to identify patterns and arrive to conclusions (Malhotra & Birks, 2007). Consequently, a questionnaire was deemed the most suitable quantitative research method to test the hypotheses and address the research questions.

3.1. Construct Measurement

The constructs in this study were developed to assess key variables affecting consumer behavior in the context of deepfake advertising in the fashion industry. The primary constructs measured include Perceived Verisimilitude, Perceived Creativity, Perceived Relevance, Awareness of Falsity, Trust in Brand, and Ad Avoidance. These constructs form the foundation of the conceptual model outlined in Chapter 2 and were essential for testing the hypotheses related to how deepfake technology influences consumer responses to deepfake advertisements. The measurement of these constructs was based on validated scales from prior research, with modifications tailored to the specific context of this study, as shown in Table 3.1.

Most of the constructs in this study were measured using a 7-point Likert scale of agreement, where participants rated their level of agreement from 1 (strongly disagree) to 7 (strongly agree). This scale was chosen for its ability to capture more nuanced responses, allowing for a broader range of consumer perceptions. Constructs such as Perceived Verisimilitude, Perceived Creativity, and Perceived Relevance required a higher degree of sensitivity in measurement, as these are subjective assessments that can vary significantly across individuals. By using a 7-point scale, the study aimed to avoid central tendency bias, ensuring that participants could express a wide spectrum of opinions, which in turn helps to measure their reactions more accurately to deepfake

advertisements. This method supports a more detailed analysis of how different levels of manipulation in ads impact consumer behavior and ad avoidance. However, two constructs, Awareness of Falsity and Trust in Brand, were originally adapted from Nijhuis (2018) and Ballester (2011) using a 5-point Likert scale. Despite this, it is important to note that when developing the final questionnaire for this study, all constructs, including Awareness of Falsity and Trust in Brand, were standardized to a 7-point Likert scale for consistency. This adjustment ensures uniformity across all measured variables, allowing for easier comparison and analysis. The decision to apply the 7-point scale universally throughout the questionnaire was made to maintain consistency in data interpretation and to align with the more complex constructs that required greater precision in capturing participant responses.

Constructs	Type of Scale	Source
Perceived Verisimilitude	7-Point Likert Scale of agreement (1 = Strongly Disagree to 7 = Strongly Agree)	Campbell & Reiman (2022)
Perceived Creativity	7-Point Likert Scale of agreement (1 = Strongly Disagree to 7 = Strongly Agree)	Yang (2006)
Perceived Relevance	7-Point Likert Scale of agreement (1 = Strongly Disagree to 7 = Strongly Agree)	Yang (2006)
Awareness of Falsity	5-Point Likert Scale of agreement (1 = Strongly Disagree to 5 = Strongly Agree)	Nijhuis (2018)
Trust in Brand	5-Point Likert Scale of agreement (1 = Strongly Disagree to 5 = Strongly Agree)	Ballester (2011)
Ad Avoidance	7-Point Likert Scale of agreement (1 = Strongly Disagree to 7 = Strongly Agree)	Youn & Kim (2019)

Table 3.1: Measurement Scale

3.2. The Experiment

The data in question was sourced from two surveys that were constructed via Qualtrics, a widely recognized survey creation software, and subsequently broadcasted via Prolific, a reputable online survey platform.

Prior to the beginning of the questionnaires, the participants were notified that the surveys would require approximately two minutes to complete. Furthermore, it was emphasized that all responses would remain anonymous, ensuring the confidentiality of the participants' identities.

In terms of sampling limitations, the experiment was limited to those who are literate in English; have Internet access; can use Qualtrics; and consented to participate in the experiment.

3.3. Experimental Design

Given the intricate nature of the deepfake technology phenomenon, a quantitative research methodology was reasoned appropriate to verify hypotheses and address research questions. Specifically, two surveys were conducted to produce additional empirical data, which could serve as another source for validating and supporting the conceptual framework model proposed in Figure 2.1.

Concerning the structure and design of the questionnaire, these surveys contained the same questions and structure, differing only in the visual stimuli shown in Annex B. Both studies used two different visual stimuli. One stimulus was a standard BALENCIAGA runway – Balenciaga Winter 23 Collection - video with no image manipulation. The other stimulus was a video from the clone spring collection, where BALENCIAGA used deepfake technology to raise awareness about its use. The collection was accompanied by a statement that explained the use of this technology.

“We see our world through a filter – perfected, polished, conformed, photoshopped. We no longer decipher between unedited and altered, genuine and counterfeit, tangible and conceptual, fact and fiction, fake and deepfake. Technology created alternate realities and identities, a world of digital clones.” – BALENCIAGA

The deliberate use of different levels of image manipulation in the study aims to cover a range of responses and perceptions. This will allow us to thoroughly investigate how different levels of manipulation affect ad avoidance and consumer behavior in the context of deepfake advertising.

The experiment was structured into three distinct sections, with participants required to complete all questions within each section before advancing to the subsequent part of the experiment. The initial section served as a concise introduction to the research. In the second section, participants were presented with a video, either a non-deepfake or a deepfake, and were prompted to respond to inquiries using a 7-point Likert scale of agreement. The Likert scale encompassed responses ranging from "Strongly Agree" to "Strongly Disagree" and was derived from established studies probing the same constructs. The questions were carefully selected and adapted from reputable sources, ensuring alignment with the study's context while preserving the essence of the original material. Detailed insight into the adaptations is delineated in

Annex B. The third and final section asks three general questions regarding their personal information (gender, age and monthly income).

3.4. Pre-test

Prior to publishing the questionnaires, it was determined that a pilot test was necessary. The purpose of the pilot test was to evaluate whether any modifications or adjustments were needed before their actual application. The assessment aimed to identify any potential misunderstandings regarding the content or questions in the questionnaires.

3.5. Data Collection and Procedures

The survey was carefully designed to be inclusive by not imposing any age, geographic, or gender restrictions. This approach assumed that the deepfake and non-deepfake phenomenon would be comprehensible to nearly every individual, regardless of their demographic characteristics.

The quantitative approach was chosen to ensure a structured and systematic examination of the research objectives and to simplify the generation of statistically significant results. The sample size of $N = 268$ participants (all of which random sampling) total was deemed adequate for achieving the study's goals and for acquiring a substantial dataset for in-depth analysis, with 137 contributing to the deepfake survey and 131 to the non-deepfake survey

The data in question was obtained from two surveys (one survey focused on a deepfake video stimuli, while the other focused on a non-deepfake stimuli) created using Qualtrics and distributed through Prolific on May 4th and 5th. Respondents were randomly chosen for each survey, and they were not informed about the presence of two separate surveys.

Chapter 4: Results of Findings

The main goal of this chapter is to align the research findings with theoretical insights to determine the impact of deepfake advertisements on ad avoidance and consumer behavior in the fashion industry.

The experiment's results may have been influenced by a variety of factors, including but not limited to user experience, the content of the video displayed, and information overload. We must consider these possible elements to gain a comprehensive understanding of the outcomes. Doing so will allow us to make informed decisions based on the results and to identify areas for improvement in future experiments.

This section seeks to gain a deeper understanding of the nature of the gathered data. It begins with a demographic overview before delving into an exploration of the collected data with the conceptual model of the study and the interplay of its constructs in both the inner and outer models.

4.1. Demographic Description

In this investigation, the sample size comprises a total of 268 participants. Survey One consists of a representative sample of $n = 137$ (survey with deepfake visual stimuli), while Survey Two encompasses a representative sample of $n = 131$ (survey with no deepfake visual stimuli).

In the analyzed sample consisting of 268 participants, the gender distribution was as follows: 142 participants (52.78%) identified as female, 123 participants (46.10%) identified as male, and 3 participants (1.12%) opted not to disclose their gender. In terms of age distribution, most participants, 112 (41.79%), were aged 18 to 25 years, followed by 87 (32.46%) in the 26 to 35 age range, indicating a predominance of individuals from Generation Z, with Millennials following closely. The outcome was as anticipated, given that the experiment was conducted digitally and distributed via Prolific. Lastly, with regards to monthly income, most participants, 88 (32.84%), reported a monthly individual income between 820€ and 1999€, with 68 participants (25.37%) reporting an income between 2000€ and 4999€.

4.2. Data Analysis

The data obtained from the administered questionnaires was subjected to thorough analysis using the Partial Least Squares (SmartPLS) estimation method. This approach was chosen due to its suitability for handling complex structural models containing high-order constructs, as highlighted by Hair et al. (2019). Additionally, the method is particularly well-suited for testing a theoretical framework from a predictive perspective, as emphasized by Hair et al. (2019). Given the complexity of the proposed model and the need for high-order measurement, the decision to use PLS software was thoroughly justified.

Upon reaching the required sample size, the analysis of the PLS-SEM conceptual model can proceed, which comprises two main parts. Firstly, the outer model (measurement model) illustrates the relationships between the constructs and the indicator variables. Secondly, the inner model (structural model) explores the direct connections between constructs (Hair et al., 2021). The subsequent subchapters give insight into the PLS-SEM outer and inner models using the PLS algorithm and bootstrapping calculation techniques.

4.2.1. Assessment of Measurement Model - Outer Model

The investigation's outer model encompasses four main aspects of the conceptual model, including internal consistency reliability, convergent validity, discriminant validity, and multicollinearity.

In the outer model, the first step involves assessing the reliability of the indicators to determine how much of the indicator variance is explained by the construct (Hair et al., 2021). While all indicators have been previously tested and used in other studies, it is important to consider that some may not be suitable for this analysis.

Indicator loadings that are above 0.7 are recommended because they are an adequate indicator of reliability and indicators with an outer loading below 0.4 should be eliminated (Hair et al., 2021). In terms of the reliability of the items within the constructs, all items, except for one, demonstrated loadings exceeding 0.7, thus signifying their appropriateness and reliability. The item that demonstrated a loading indicator below 0.7 was eliminated (PV3 – “The ad resembles daily life tasks”). The same patterns can be observed when analyzing the complete dataset, which includes responses from both the survey with deepfake-generated visual stimuli and the survey without deepfake-

generated content, as well as when looking at the separate data groups, as shown in Annexes from D to I.

The study assessed convergent validity by calculating the Average Variance Extracted (AVE). The results showed that all constructs had AVE values higher than the threshold of 0.5 (Hair et al., 2021), with values ranging from 0.670 to 0.917, as indicated in Annex J. This indicates that the complete conceptual model is reliable. The same can be observed when analyzing the separate data groups, with values ranging from 0.661 to 0.916 – for the deepfake group and 0.686 to 0.919 – for the non-deepfake group.

Moreover, the study involved the calculation of Cronbach's alpha and Composite Reliability (rho c and rho a) values to assess construct reliability, also known as internal consistency. Cronbach's alpha values, ranged from 0.878 to 0.940, all of which surpass the commonly accepted lower limit of 0.7 (Hair et al., 2010), signifying the reliability of all constructs. Furthermore, the calculation of Composite reliability, recognized as a more precise measure than Cronbach's alpha (Fornell & Larcker, 1981; Loureiro & Kaufmann, 2016), was undertaken. All Composite reliability values (including rho c and rho a) ranged from 0.910 to 0.957 and 0.894 to 0.943, respectively, exceeding both the lower limit of 0.7 and the stricter threshold of 0.8, thus confirming the reliability of all constructs (Nunnally, 1978; Loureiro & Kaufmann, 2016). It is important to note the same patterns can be observed when analyzing the complete dataset, which includes responses from both the survey with deepfake-generated visual stimuli and the survey without deepfake-generated content, as well as when looking at the separate data groups. All values are available in Annex J.

To evaluate the distinctiveness of the first-order constructs, the Fornell-Larcker criterion was applied. This criterion states that the square root of each construct's average variance extracted (AVE) should exceed its correlations with other constructs (Fornell & Larcker, 1981). Based on the results shown in Annex J, all constructs fulfil these conditions, indicating appropriate distinctiveness.

Finally, it is essential to evaluate the collinearity of indicators in the model, as higher correlations can lead to increased standard errors. The Variance Inflation Factor (VIF) is studied for this purpose. In this model, the VIF ranges from 5.284 (PR3) to 1.814 (AA3) as indicated in Annex K. The general rule is that VIF values should not exceed 10 (Henseler et al., 2009), while more conservative recommendations suggest values below 5 or 3.3 (Kock & Lynn, 2012). When applying the most conservative criteria, it can be determined that the results provide no evidence of multicollinearity issues.

4.2.2. Assessment of Structural Model - Inner Model

The structural model, also referred to as the inner model, is designed to elucidate the relationships among variables and to disclose the outcomes of hypothesis tests (Hair et al., 2011). To study the relationships between the variables and test the hypothesis's validity, the structural model was measured and is presented in Figure 4.1. Consequently, this model is used to appraise the path coefficients derived from the PLS algorithm computation and to comprehend the significance of the paths linking the latent constructs (Hair et al., 2011).

Instigating with the model's fitness, it is worth noting that a Standardized Root Mean Square Residual (SRMR) index within the range of 0 to 0.08 indicates a good fit (Hu & Bentler, 1999). The model demonstrates a strong fit with an SRMR of 0.061 saturated model and 0.064 estimated model. In addition to the outer model, it is essential to evaluate the VIF for the inner model, considering the multicollinearity of latent variables.

It is also important to assess the predictive capabilities of the model. This assessment has three parts: first, the coefficient of determination (the R-squared); second, path coefficients and third, the bootstrapping. R-squared values should be between 0 and 1 (Chin, 1999). The results show weak values for awareness of falsity (0.265) and moderate values for ad avoidance (0.479).

The Stone-Geisser's Q-squared value is used to assess the predictive relevance of the model (Geisser, 1974; Stone, 1974). A Q-squared value greater than 0 indicates that the model demonstrates predictive accuracy for the respective endogenous construct. In this study, the Q-squared values obtained are as follows: Awareness of Falsity (0.253) and Ad Avoidance (0.445). Since both values exceed 0, it can be concluded that the model possesses sufficient predictive relevance.

A complementary analysis conducted is the calculation of the effect size of each construct using F-square. The results indicate varying effect sizes. According to Cohen (1988, p. 414), F-Square is the change in R-Square when an exogenous variable is removed from the model, an F-square value between 0.02 and 0.15 signifies a small effect size, between 0.15 and 0.35 represents a medium effect size, and values exceeding 0.35 indicate a large effect size. In this study, Perceived Creativity → Awareness of Falsity relationship (0.008) and Awareness of Falsity → Ad Avoidance (0.004) shows no significant effect, as the value falls below 0.02. The relationships

Perceived Verisimilitude → Ad Avoidance (0.028), Perceived Creativity → Ad Avoidance (0.063), Perceived Relevance → Ad Avoidance (0.124), and Trust in Brand → Perceived Relevance → Ad Avoidance (0.023) all exhibit small effect sizes. However, Perceived Verisimilitude → Awareness of Falsity relationship (0.281) demonstrates a medium effect size.

When assessing the interplay between two constructs within the inner model, it is imperative to examine three fundamental components: the path coefficients (β) of the constructs, their respective p-values, and their t-values.

Based on the analysis of the p-values and t-values, it is evident that two relationships within the data lack statistical significance, meaning their p-values are lower than 0.05. The relationships Perceived Creativity → Awareness of Falsity ($\beta = -0.079$; $p = 0.164$; $t = 1.391$) and Awareness of Falsity → Ad Avoidance ($\beta = -0.050$; $p = 0.330$; $t = 0.975$) exhibit p-values exceeding 0.05 and t-values below 1.96, meaning a lack of statistical significance in the model. The remaining relationships are statistically significant. After analyzing their path coefficients, it is evident that two relationships have a notably low impact between the constructs: Perceived Verisimilitude → Ad Avoidance ($\beta = -0.150$; $p = 0.011$; $t = 2.536$), Perceived Creativity → Ad Avoidance ($\beta = -0.232$; $p = 0.000$; $t = 3.690$) and Trust in Brand → Perceived Relevance → Ad Avoidance ($\beta = 0.093$; $p = 0.024$; $t = 2.258$). Similarly, the remaining relationships exhibit a moderate effect among their variables, such as Perceived Verisimilitude → Awareness of Falsity ($\beta = -0.483$; $p = 0.000$; $t = 8.122$), and Perceived Relevance → Ad Avoidance ($\beta = -0.430$; $p = 0.000$; $t = 5.323$).

Hypothesis 1 suggests that when an advertisement is perceived as more realistic, ad avoidance decreases. The results support the hypothesis, showing that when perceived verisimilitude is high, ad avoidance tends to decrease ($\beta = -0.150$; $p = 0.011$; $t = 2.536$). This finding indicates that consumers are less likely to avoid ads they perceive as highly realistic. Hypothesis 2 proposes that greater perceived verisimilitude decreases awareness of falsity, which in turn decreases ad avoidance. The results strongly support this hypothesis ($\beta = -0.573$; $p = 0.000$; $t = 8.656$), indicating that a higher degree of verisimilitude reduces consumers' awareness of falsity, thereby mitigating ad avoidance. Hypothesis 3 examines whether greater perceived creativity of an advertisement decreases ad avoidance. The results support the hypothesis, showing that creativity in advertisements can lead to lower ad avoidance ($\beta = -0.232$; $p = 0.000$; $t = 3.690$). Highly creative ads attract more attention and are less likely to be avoided. Hypothesis 4 suggests that greater perceived creativity decreases awareness of falsity, which in turn

decreases ad avoidance. The analysis shows that although perceived creativity can decrease awareness of falsity, the effect on ad avoidance is not statistically significant in this case ($\beta = -0.037$; $p = 0.578$; $t = 0.557$). Hypothesis 5 investigates the relationship between perceived relevance and ad avoidance. The results indicate that when advertisements are perceived as highly relevant, ad avoidance decreases ($\beta = -0.503$; $p = 0.000$; $t = 5.419$), confirming the importance of relevance in engaging consumers. Hypothesis 6 hypothesizes that greater awareness of falsity increases ad avoidance. However, the results do not support this, as the relationship between awareness of falsity and ad avoidance is not statistically significant ($\beta = 0.000$; $p = 0.998$; $t = 0.003$). This suggests that even if consumers recognize manipulation in a deepfake ad, it does not necessarily lead to higher ad avoidance. Hypothesis 7 evaluates if trust in a brand influences the relationship between perceived relevance and ad avoidance. The analysis shows that brand trust, when acting through perceived relevance, significantly reduces ad avoidance ($\beta = 0.093$; $p = 0.024$; $t = 2.258$).

Hypothesis	Relationship	Std β	p-value	t-value	f-squared	Decision
H1	PV $\rightarrow$ AA	- 0.150	0.011*	2.536	0.028	H1: supported
H2	PV $\rightarrow$ AF	- 0.483	0.000*	8.122	0.281	H2: supported
H3	PC $\rightarrow$ AA	- 0.232	0.000*	3.690	0.063	H3: supported
H4	PC $\rightarrow$ AF	- 0.079	0.164	1.391	0.008	H4: not supported
H5	PR $\rightarrow$ AA	- 0.430	0.000*	5.323	0.124	H5: supported
H6	AF $\rightarrow$ AA	- 0.050	0.330	0.975	0.004	H6: not supported
H7	TB $\rightarrow$ PR $\rightarrow$ AA	0.093	0.024*	2.258	0.023	H7: supported

Table 4.1: PLS-SEM Bootstrapping Results – Complete dataset

* $p < 0.05$ | Note: PV = Perceived Verisimilitude; PC = Perceived Creativity; PR = Perceived Relevance; AF = Awareness of Falsity; TB = Trust in Brand; AA = Ad Avoidance

4.2.3. Mediation Analysis

As a completion of the analysis of the results, it is essential to examine potential mediation effects. These effects entail the involvement of a third variable that serves an intermediary role in the relationship between dependent and independent variables (Cepeda-Carrion et al., 2018). A complete mediation occurs when the direct effect lacks significance, while the indirect effect does, signifying its presence only when confirmed through the mediator (Cepeda-Carrion et al., 2018).

While the mediation effect is statistically significant ($p = 0.024$), the small effect size ($\beta = 0.093$) suggests that the mediator Trust in Brand $\rightarrow$ Perceived Relevance $\rightarrow$ Ad Avoidance has a weak influence. So, while it may be statistically valid, it may not be considered a strong mediator due to the small size of the effect.

4.2.4. Analysis of Model 1 – Deepfake

In a previous phase, two surveys were conducted, each featuring distinct visual stimuli. These stimuli were categorized into two groups: "deepfake advertising" and "non-deepfake advertising." The "deepfake advertising" group, represented as Model 1, consisted of 137 participants, while the "non-deepfake advertising" group, indicated as Model 2, comprised 131 participants. Moving forward, the next step is to conduct a detailed analysis of the model within its individual data groups. Specifically, we will be focusing on the analysis of model 1, which is based on the survey involving deepfake visual stimuli. To ensure a comprehensive analysis, both the outer and inner models will be examined.

The reliability of the items within the constructs was evaluated, and it was found that all items showed loadings exceeding 0.7. This indicates that the items are appropriate and reliable, as demonstrated in Annex G.

In the study, the findings show that all constructs exhibited AVE values that surpassed the threshold of 0.5. The specific AVE values from Model 1 ranged from 0.661 to 0.916, indicating a strong level of variance explained by the constructs, as shown in Annex J.

The reliability of all constructs was assessed using Cronbach's alpha values, which ranged from 0.861 to 0.936, indicating strong internal consistency. Furthermore, the composite reliability values, including rho c and rho a, were examined, and found to range from 0.906 to 0.956 and 0.876 to 0.943, providing additional confirmation of the reliability of all constructs, as indicated in Annex J.

In model 1, the variance inflation factor (VIF) ranges from 5.236 for variable PR3 to 1.795 for variable AA3. Based on these results, there is no evidence of multicollinearity issues, indicating that the independent variables in the model are not highly correlated with each other. This can be observed in Annex L.

Upon analyzing the internal structure of model 1, it was found that the model demonstrated a strong fit. Specifically, the saturated model had an SRMR of 0.068, while the estimated model had an SRMR of 0.071. These findings indicate a robust fit for the model.

The R-squared value offers valuable insights into the impact of one factor's variance on another. In this instance, the findings indicate that awareness of falsity has a low impact (0.222), while ad avoidance has a moderate impact (0.410).

In this data group, the Q-squared values obtained are Awareness of Falsity (0.192) and Ad Avoidance (0.335). As both values are greater than 0, it can be concluded that the model demonstrates adequate predictive relevance.

An additional analysis performed involves calculating the effect size of each construct using F-square. In this model, the Awareness of Falsity → Ad Avoidance relationship (0.007) shows no significant effect, as it falls below the 0.02 threshold. The relationships Perceived Verisimilitude → Ad Avoidance (0.031), Perceived Verisimilitude → Awareness of Falsity (0.127), Perceived Creativity → Ad Avoidance (0.057), Perceived Creativity → Awareness of Falsity (0.42), Perceived Relevance → Ad Avoidance (0.087), and Trust in Brand → Perceived Relevance → Ad Avoidance (0.020) all reflect small effect sizes.

Upon analysis of the p-values and t-values of this model, it is apparent that three relationships within the data do not show statistical significance, differing from the results observed in the complete data group. Specifically, the relationship between the constructs: Perceived Verisimilitude → Ad Avoidance ($\beta = -0.164$; $p = 0.054$; $t = 1.925$), Awareness of Falsity → Ad Avoidance ($\beta = -0.072$; $p = 0.336$; $t = 0.961$), and Trust in Brand → Perceived Relevance → Ad Avoidance ($\beta = 0.109$; $p = 0.134$; $t = 1.499$) have p-values higher than 0.05 and t-values lower than 1.96. These findings suggest a lack of statistical significance in the mentioned relationships. Among the remaining relationships, the analysis reveals that some relationships are statistically significant. Upon examining the path coefficients, three relationships exhibit notably low impact between the constructs: Perceived Creativity → Ad Avoidance ($\beta = -0.240$; $p = 0.011$; $t = 2.534$), and Perceived Creativity → Awareness of Falsity ($\beta = -0.201$; $p = 0.036$; $t = 2.102$). Similarly, the remaining relationships demonstrate a low to moderate effect among their variables, such as Perceived Verisimilitude → Awareness of Falsity ($\beta = -0.349$; $p = 0.001$; $t = 3.474$), and Perceived Relevance → Ad Avoidance ($\beta = -0.391$; $p = 0.001$; $t = 3.191$).

Hypothesis 1 suggests that ad avoidance decreases as deepfake advertisements are perceived as more realistic. However, since $p\text{-value} > 0.05$, H1 on the deepfake model is not supported ($\beta = -0.164$; $p = 0.054$; $t = 1.925$). This indicates that when deepfake ads are perceived as authentic and lifelike, it doesn't mean consumers are less

likely to avoid them. Hypothesis 2 proposes that consumers' ability to detect falsity decreases as the perceived verisimilitude of deepfake advertisements increases. The results support this hypothesis, demonstrating that as verisimilitude increases, consumers' ability to detect falsity decreases ($\beta = -0.349$; $p = 0.001$; $t = 3.474$). This suggests that realistic deepfake ads make it harder for consumers to recognize manipulation. Hypothesis 3 examines whether greater perceived creativity in deepfake advertisements reduces ad avoidance. The results support this hypothesis, showing that creative deepfake ads lead to lower levels of ad avoidance ($\beta = -0.240$; $p = 0.011$; $t = 2.534$). Creative elements in the ads appear to capture attention and interest, thereby reducing avoidance behavior. Hypothesis 4 suggests that greater perceived creativity in an advertisement decreases awareness of ad falsity, subsequently reducing ad avoidance. The results support this hypothesis, showing that higher creativity in deepfake ads reduces consumers' awareness of falsity ($\beta = -0.201$; $p = 0.036$; $t = 2.102$). This indicates that creative deepfake ads make it difficult for consumers to recognize potential manipulation, thereby reducing their tendency to avoid the ad. Hypothesis 5 investigates the relationship between perceived relevance and ad avoidance. The results support this hypothesis, demonstrating that advertisements perceived as highly relevant to the consumer significantly reduce ad avoidance ($\beta = -0.391$; $p = 0.001$; $t = 3.191$), confirming that relevance plays a critical role in keeping consumers engaged and reducing their inclination to avoid the ad. Hypothesis 6 hypothesizes that greater awareness of falsity increases ad avoidance. However, the results do not support this hypothesis, as the relationship between awareness of falsity and ad avoidance is not statistically significant ($\beta = -0.072$; $p = 0.336$; $t = 0.961$). This suggests that even if consumers recognize manipulation in a deepfake ad, it does not necessarily lead to higher ad avoidance. Hypothesis 7 evaluates whether trust in the brand influences the relationship between perceived relevance and ad avoidance. The analysis shows that brand trust does not significantly moderate this relationship ($\beta = 0.109$; $p = 0.134$; $t = 1.499$), indicating that brand trust may not play a strong role in reducing ad avoidance in this deepfake model.

Hypothesis	Relationship	Std β	p-value	t-value	f-squared	Decision
H1	PV $\rightarrow$ AA	-0.164	0.054	1.925	0.031	H1: not supported
H2	PV $\rightarrow$ AF	-0.349	0.001*	3.474	0.127	H2: supported
H3	PC $\rightarrow$ AA	-0.240	0.011*	2.534	0.057	H3: supported
H4	PC $\rightarrow$ AF	-0.201	0.036*	2.102	0.042	H4: supported
H5	PR $\rightarrow$ AA	-0.391	0.001*	3.191	0.087	H5: supported
H6	AF $\rightarrow$ AA	-0.072	0.336	0.961	0.007	H6: not supported
H7	TB $\rightarrow$ PR $\rightarrow$ AA	0.109	0.134	1.499	0.020	H7: not supported

Table 4.2: PLS-SEM Bootstrapping Results Deepfake dataset

4.2.5. Analysis of Model 2 – Non-Deepfake

The concluding section of the data analysis chapter focuses on dissecting model 2, which is derived from a survey that does not employ deepfake visual stimuli. To ensure a comprehensive analysis, both the external and internal models will be subjected to examination.

The reliability of the items comprising the constructs underwent thorough evaluation, and it was determined that all items displayed loadings above 0.7. This observation verifies the suitability and reliability of the items, as indicated in Annex I.

The study's analysis reveals that all constructs demonstrated AVE values surpassing the 0.5 threshold. Specifically, in Model 2, the AVE values ranged from 0.686 to 0.919, indicating a substantial degree of variance accounted for by the constructs, as denoted in Annex J.

The study's measures were evaluated for reliability using Cronbach's alpha values, which showed strong internal consistency ranging from 0.886 to 0.947. Additionally, composite reliability values (ρ_c and ρ_a) were examined, further confirming the reliability of the measures with values ranging from 0.916 to 0.959 and 0.901 to 0.951, respectively, as shown in Annex J.

The variance inflation factor (VIF) for model 2 ranges from 6.735 for variable PR3 to 1.827 for variable AA3. These VIF values suggest that there is no evidence of multicollinearity issues, indicating that the independent variables in the model are not highly correlated with each other. This can be observed in Annex M.

Upon analyzing the internal structure of model 2, it was evident that the model displayed a robust fit. Specifically, the saturated model boasted an SRMR of 0.068, whereas the estimated model showcased an SRMR of 0.072. These findings unequivocally signify a robust fit for the model.

The R-squared value results illustrate that the awareness of falsity has a relatively low impact (0.344), while ad avoidance has a more significant effect (0.580).

Within model 2, the Q-squared values for Awareness of Falsity (0.325) and Ad Avoidance (0.541) were obtained. Given that both values are greater than 0, it can be determined that the model has strong predictive relevance.

An additional analysis performed involves calculating the effect size of each construct using F-square. In this model, the Perceived Verisimilitude → Ad Avoidance (0.0016), Perceived Creativity → Awareness of Falsity (0.002), and Awareness of Falsity → Ad Avoidance (0.000) relationships show no significant effect, as they fall below the 0.02 limit. Perceived Creativity → Ad Avoidance (0.053), Perceived Relevance → Ad Avoidance (0.216), and Trust in Brand → Ad Avoidance (0.029), and Trust in Brand → Perceived Relevance → Ad Avoidance (0.025) all reflect small effect sizes. In contrast, the relationships Perceived Verisimilitude → Awareness of Falsity (0.443), exhibit a large effect size.

Reviewing the p-values and t-values, it is evident that four relationships in the data lack statistical significance, and these results differ from those observed in the complete data set. Specifically, the relationships between the following constructs lack statistical significance: Perceived Verisimilitude → Ad Avoidance ($\beta = -0.107$; $p = 0.238$; $t = 1.180$), Perceived Creativity → Awareness of Falsity ($\beta = -0.037$; $p = 0.578$; $t = 0.557$), Awareness of falsity → Ad Avoidance ($\beta = 0.000$; $p = 0.998$; $t = 0.003$), and Trust in Brand and Perceived Relevance → Ad Avoidance ($\beta = 0.072$; $p = 0.066$; $t = 1.835$). All these relationships display p-values exceeding 0.05 and t-values below 1.96. These results indicate that the mentioned relationships do not exhibit statistical significance. The remaining relationships are statistically significant. Upon analyzing their path coefficients, it is evident that one of the remaining relationships has a notably low impact between the constructs: Perceived Creativity → Ad Avoidance ($\beta = -0.197$; $p = 0.014$; $t = 2.456$). Similarly, the remaining relationships exhibit a moderate effect among their variables, such as Perceived Verisimilitude → Awareness of Falsity ($\beta = -0.573$; $p = 0.000$; $t = 8.656$), and Perceived Relevance → Ad Avoidance ($\beta = -0.503$; $p = 0.000$; $t = 5.419$).

Hypothesis 1 suggests that higher perceived realism reduces ad avoidance. However, the results do not support this hypothesis, as the relationship between awareness of falsity and ad avoidance is not statistically significant ($\beta = -0.107$; $p = 0.238$; $t = 1.180$), suggesting that even if consumers perceived a non-deepfake ad as authentic, it does not lead to ad avoidance. Hypothesis 2 suggests that greater perceived verisimilitude decreases awareness of falsity. The results supported the hypothesis, indicating that as something seems more realistic, people are less able to detect that it's false ($\beta = -0.573$; $p = 0.000$; $t = 8.656$). This supports the idea that more realistic non-deepfake ads make it harder for people to recognize that they're being manipulated. H3 evaluates whether greater perceived creativity decreases ad avoidance. The results

supported the hypothesis, demonstrating that non-deepfake ads with creative content led to decreased levels of ad avoidance ($\beta = -0.197$; $p = 0.014$; $t = 2.456$). It appears that creative elements in advertisements attract attention and maintain consumer engagement, thereby reducing avoidance behavior. Hypothesis 4 suggests that when an advertisement is perceived as more creative, it leads to lower awareness of ad falsity, which in turn reduces ad avoidance. The results do not support this hypothesis, as the relationship between perceived creativity and awareness of falsity is not statistically significant ($\beta = -0.037$; $p = 0.578$; $t = 0.557$), suggesting that creativity does not have the same falsity-masking effect in the non-deepfake advertisements. Hypothesis 5 inspects the relationship between perceived relevance and ad avoidance. The results supported the hypothesis, indicating that ads considered relevant to consumers' interests notably decrease ad avoidance ($\beta = -0.503$; $p = 0.000$; $t = 5.419$). This emphasizes the crucial role of ad relevance in capturing consumers' interest and decreasing their inclination to avoid the ad. Hypothesis 6 hypothesizes that greater awareness of falsity increases ad avoidance. However, the results do not support this hypothesis ($\beta = 0.00$; $p = 0.998$; $t = 0.003$), suggesting that even if consumers recognize manipulation in a non-deepfake ad, it does not significantly lead to ad avoidance. Hypothesis 7 evaluates whether trust in the brand influences the relationship between perceived relevance and ad avoidance. The analysis showed that brand trust does not significantly moderate this relationship ($\beta = 0.072$; $p = 0.066$; $t = 1.835$). This suggests that brand trust may not play a strong role in reducing ad avoidance in the context of non-deepfake advertisements.

Hypothesis	Relationship	Std β	p-value	t-value	f-squared	Decision
H1	PV $\rightarrow$ AA	-0.107	0.238	1.180	0.0016	H1: not supported
H2	PV $\rightarrow$ AF	-0.573	0.000*	8.656	0.443	H2: supported
H3	PC $\rightarrow$ AA	-0.197	0.014*	2.456	0.053	H3: supported
H4	PC $\rightarrow$ AF	-0.037	0.578	0.557	0.002	H4: not supported
H5	PR $\rightarrow$ AA	-0.503	0.000*	5.419	0.216	H5: supported
H6	AF $\rightarrow$ AA	0.000	0.998	0.003	0.000	H6: not supported
H7	TB $\rightarrow$ PR $\rightarrow$ AA	0.072	0.066	1.835	0.025	H7: not supported

Table 4.3: PLS-SEM Bootstrapping Results Non-Deepfake dataset

Based on the data analysis, it has been established that both Model 1 and Model 2 demonstrate independent support and validation. Furthermore, when these models are combined, they exhibit substantial support and validation. It is important to note that although individual comparisons of the models within specific data groups may not significantly support certain hypotheses, the overall analysis of the complete data set confirms the validation and support for the model.

Chapter 5: Discussion of Findings

Although the model is supported, it is important to note that not all the developed hypotheses demonstrate satisfactory levels of reliability or validity, and therefore being classified as not supported. Conversely, all discriminants show validity and collinearity in the outer-model analysis. The results of the inner model indicate that the conceptual model also fits well and shows no collinearity issues.

To fully understand the inner and outer models, it is important to carefully examine the connections between each component.

5.1. Perceived Verisimilitude and Ad Avoidance

Research indicates that consumers actively avoid advertisements that they perceive as manipulated (Baek and Morimoto, 2012). Additionally, when consumers engage with such advertisements, their recognition of manipulation can obstruct their involvement in the ad's narrative (Kim, Ratneshwar, et al., 2017), causing them to focus on the manipulation rather than the core message and imagery of the ad (Dessart and Pitardi, 2019; Russell, 2002; Russell and Russell, 2009).

The analysis of Hypothesis 1 demonstrates that perceived verisimilitude influences the audience's response to advertisements. When ads are perceived as highly realistic, they tend to evoke fewer negative reactions, such as ad avoidance, as consumers who process highly realistic ads more positively, often become more immersed in the narrative (Campbell et al., 2022). The results of the complete model support this hypothesis, indicating that perceived verisimilitude plays an essential role in reducing ad avoidance.

However, upon analyzing the individual data sets, neither the deepfake nor the non-deepfake model exhibited significant results. When deepfake technology is well-executed and convincing, it does not necessarily mean that consumers will be less likely to avoid deepfake advertisements. This implies that a high level of realism in deepfake advertisements has the potential to deceive viewers into perceiving the content as authentic. On the other hand, for non-deepfake advertisements, the lack of significance can suggest that traditional advertisements may already be expected to appear real, thereby making verisimilitude less of a determining factor.

5.2. Perceived Verisimilitude and Awareness of Falsity

Consumers generally prefer authentic content over fake content (Beverland, Lindgreen, and Vink, 2008; Becker, Wiegand, and Reinartz, 2019; Stern, 1994), leading to increased skepticism towards manipulated ads (Obermiller, Spangenberg, and MacLachlan, 2005) and a decreased likelihood of acceptance even if the messages are understood (Spielmann and Orth, 2020).

The analysis of Hypothesis 2 demonstrated that perceived verisimilitude can affect people's ability to detect falsehood across all models, indicating that a higher degree of verisimilitude reduces consumers' awareness of falsity, thereby mitigating ad avoidance. This implies that when advertisements are perceived as highly realistic, consumers are less likely to detect any manipulation, leading to a reduction in their awareness of falsity. When analyzing the deepfake model it demonstrates a strong effect, as the perceived verisimilitude decreases viewers' awareness of falsity, making them less critical and more trusting of the content. The same trend was observed in the non-deepfake model; however, the effect was more evident. This could be explained because traditional advertisements are not often viewed with skepticism, so the focus remains on how authentic the ad seems.

5.3. Perceived Creativity and Ad Avoidance

Marketing experts believe that highly creative advertisements are more effective in overcoming consumers' barriers, capturing their attention, stimulating positive responses, and reinforcing their attitudes toward the promoted brand (Marra 1990; Ogilvy 1983; Rosengren, Dahlen, and Modig 2013; Zinkhan 1993). According to Rosengren et al. (2020), the primary advantage of creativity in advertising lies in its capacity to make ads enjoyable and appealing. Creative ads are more likely to capture attention and positively influence brand attitudes. This aligns with theories suggesting creativity enhances advertising effectiveness by making ads more enjoyable and attention-grabbing (Rosengren et al., 2013).

Creativity plays an important role in reducing ad avoidance in all models, with a slightly stronger effect in the deepfake model. When analyzing the deepfake model, creativity appears to increase engagement, possibly because deepfake allows for highly innovative and unexpected visual effects, which can easily capture viewers' attention. The ability to manipulate and generate a new reality opens new possibilities that non-deepfake advertisements cannot explore as extensively. In the context of the non-

deepfake model, while creativity still decreases ad avoidance its effect is slightly less pronounced. Traditional advertisements may rely more on more common creative strategies that, while effective, do not possess the uniqueness and disruptive power of deepfakes.

5.4. Perceived Creativity and Awareness of Falsity

Although greater advertisement creativity generally has a positive impact, it may lead viewers to question the authenticity of the advertisement. The growing availability of sophisticated manipulation tools enables advertisers to produce highly creative content that distorts the boundaries between reality and fiction, or even fabricate entirely new scenarios. (Campbell, Plangger, Sands & Kietzmann, 2022).

While creativity in advertising generally reduces ad avoidance, overly creative or manipulative content, especially those relating deepfake technologies, can increase skepticism and increase the awareness of falsity. This hypothesis was not supported in the complete dataset and non-deepfake model but was supported in the deepfake model, indicating that creativity can mask falsity in deepfake ads but not in non-deepfake ones. In the deepfake model, creativity not only engages consumers but also distracts them from realizing that an advertisement can be manipulated. This implies that when deepfake technology is combined with high creativity, consumers may be less likely to inspect the ad for signs of falsity. In contrast, creativity does not have the same falsity-masking in the non-deepfake model. This indicates that while traditional advertisements can be creative, consumers are more familiar with their elements and are less likely to be deceived by them.

5.5. Perceived Relevance and Ad Avoidance

Research suggests that when people believe an advertisement is closely aligned with their interests or needs, they are less inclined to actively avoid or disregard it. This implies that perceived relevance of an advertisement can have a notable impact on how consumers perceive and engage with advertising content (Brinson & Britt, 2021; Dodoo & Wen, 2021; Jung, 2017; Kelly et al., 2010; Li et al., 2020).

Perceived relevance has a significant influence on whether consumers tend to avoid certain advertisements. When ads are considered relevant to their personal needs or interests, ad avoidance decreases. This hypothesis was supported across all models. Deepfake advertisements can benefit greatly from relevance, as consumers are more

forgiving of manipulated content if they perceive it as personally meaningful. Even when consumers are aware that an ad is manipulated, if the content feels relevant to their needs or interests, they are less likely to avoid it. Non-deepfake advertisements also show strong support for relevance reducing ad avoidance, even more than deepfake ads. In non-manipulated environments, consumers expect ads to be relevant and aligned with their preferences. When this expectation is met, ad avoidance decreases significantly.

5.6. Awareness of Falsity and Ad Avoidance

The connection between awareness of falsity and ad avoidance is rooted in consumers' sensitivity to misinformation and authenticity in advertising. Consumers are often cautious of persuasive efforts, especially when they perceive an ad to be misleading or not genuine. This caution can lead to defensive reactions. According to Friestad and Wright (1994), this sensitivity to deception can prompt negative responses, causing individuals to disengage from the ad. Authenticity, as emphasized in situations such as brand expansions (Spiggle, Nguyen, & Caravella, 2012) and social media (Audrezet, De Kerviler, & Moulard, 2020), plays a significant role in shaping consumer perceptions. When ads are viewed as fake, consumers are less likely to accept the offer (Spiellmann and Orth, 2020). As highlighted by Çelik, Çam, and Koseoglu (2022), identifying the factors driving ad avoidance is crucial, and awareness of falsity may be a key trigger that influences this behavior.

The Hypothesis 6 was not supported across all models. This lack of significant findings suggests that consumers might be tolerant of a certain level of manipulation, especially in deepfake ads, if other factors such as relevance and creativity are strong. Even when consumers know that they are seeing fake or altered content, they may not necessarily avoid the ad if it provides value in other ways. For non-deepfake ads, realizing that the content is fake may not matter because people expect these ads to be genuine. The fact that knowing the content is fake doesn't stop people from avoiding non-deepfake ads might show that consumers are used to traditional advertising methods.

5.7. Trust in Brand and Perceived Relevance and Ad Avoidance

The relationship between brand trust, perceived relevance, and ad avoidance suggests that trust in a brand can mitigate consumers' tendency to avoid ads, particularly when they find the advertisements relevant. Brand trust involves a willingness to accept risks

based on the belief that the brand will deliver value. This trust fosters feelings of assurance and security, as consumers form expectations that the brand is dependable and trustworthy. When such trust is established, consumers are more likely to engage with the brand's messages, even if they encounter elements of uncertainty or potential misinformation (Ballester, 2011).

The level of trust in a brand can positively impact the perceived relevance, making consumers less likely to avoid ads from trusted sources. This hypothesis was supported for the complete model, but not for the deepfake and non-deepfake models. When consumers trust a brand, they may be more willing to engage with their ads, even if they detect manipulation or falsity. This finding suggests that brand trust acts as a barrier, helping consumers stay engaged even in manipulated environments. Conversely, when analysing the individual models, trust did not have a significant effect on reducing ad avoidance.

5.8. Effects of Increased Manipulation

Existing research suggests that consumers may react differently to an advertisement if they are aware that it has been manipulated (Friestad & Wright, 1994). However, in certain contexts such as clothing or personalized products, consumers may be very receptive to synthetic ads. It may also be the case that some consumers simply are not concerned about ad falsity (Campbell, Plangger, Sands & Kietzmann, 2022).

The findings support the idea that higher levels of manipulation sophistication, utilizing deepfake technology, can enhance the viewers perceived verisimilitude, creativity, and relevance of an advertisement, which in turn can decrease ad avoidance. Specifically, the study demonstrated that when consumers perceive advertisements as creative and relevant to their needs, they are more likely to engage with them, even if they are aware of manipulation or falsity.

Consumers may be less inclined to avoid advertisements that they perceive as realistic, creative, and tailored to their preferences, thus offsetting any skepticism towards manipulation. This effect was strong in both models, but more evident in the deepfake model. However, it is important to note that while, in this case, greater manipulation sophistication can generally lower ad avoidance, the findings suggest that there is a limit to this effect. When consumers perceive the manipulation as overly deceptive or unethical, this can lead to an increase in ad avoidance, highlighting the importance of balance in manipulation techniques. Moreover, it's crucial to acknowledge

that the impact of this effect may vary depending on the specific stimuli presented to viewers, as it has the potential to influence their perceptions significantly.

Chapter 6: Conclusion

As deepfake technology grows in popularity in numerous sectors, its impact on advertising, particularly in the fashion industry, deserves a thorough examination. This study added to the understanding of how deepfake-generated advertisements can affect consumer behavior, ad avoidance, and perceptions of authenticity. Through its investigation into the connection of AI, synthetic media, and advertisement, the research has opened new possibilities for theoretical exploration and practical application, highlighting the need for a careful balance between innovation and ethical considerations in advertising.

The research provides a better understanding on how consumers react to deepfake advertisements. Contrary to initial expectations, the presence of deepfake technology does not necessarily result in increased ad avoidance. Consumers seem less likely to avoid advertisements that they perceive as creative or personally relevant, even if they are aware of manipulation. Nevertheless, when deepfakes are seen as overly deceptive or manipulative, they cause consumer skepticism and increase the awareness of falsity, leading to higher ad avoidance.

The study found that three main factors significantly impact how consumers react to deepfake advertisements: perceived verisimilitude, perceived creativity, and perceived relevance. Advertisements that effectively mimic reality without raising authenticity concerns can capture the audience's attention. Using deepfake technology in innovative ways to create ads can reduce ad avoidance. Most importantly, ads that match consumer needs and interests are not only better received but also less likely to be avoided. This highlights the importance of a thoughtful approach to using deepfake technology in advertising strategies, with a focus on personalization and creative storytelling. However, the study revealed a duality in consumer skepticism towards deepfake generated content. While participants exhibited a general untrust, it did not equivalently translate into negative outcomes. When deepfakes were employed creatively and aligned with consumer interests, the skepticism was often outweighed by the perceived value of the content.

For the fashion industry, these findings present both opportunities and challenges. When executed thoughtfully, deepfakes offer a powerful tool for creating innovative, engaging content that can significantly enhance consumer interaction with fashion brands. However, the use of this technology can evoke skepticism and increase ad avoidance if consumers perceive the manipulation as deceptive or irrelevant. The

research emphasizes that creativity and relevance are crucial factors in mitigating ad avoidance, suggesting that well-designed deepfake advertisements can be highly effective when meaningfully connected to consumer needs. Based on these insights, fashion industry professionals looking to leverage deepfake technology in their advertising efforts should prioritize transparency by openly communicating the use of synthetic media to maintain consumer trust. Focus on harnessing deepfake technology to create compelling narratives that resonate deeply with your audience. Personalization should be a key consideration, tailoring content to align closely with consumer preferences and needs, thereby mitigating risks associated with perceptions of manipulation. Carefully weigh the ethical implications of deepfake usage to ensure long-term consumer trust as technology continues to advance.

This study lays the base for understanding the implications of deepfake advertisements on ad avoidance and consumer behavior within the fashion industry. However, as this is an emerging area of research, several recommendations can be made to guide future studies and industry practices. First, the development of ethical frameworks and regulations for the use of deepfake technology in advertising is essential. These frameworks should prioritize transparency, consumer consent, and authenticity in synthetic media, addressing potential misuse and fostering trust. Future research should explore how such guidelines could be effectively implemented across different industries while aligning with consumer expectations and regulatory requirements. While the potential risks of deepfakes are evident, their creative applications in advertising present significant opportunities. Future research should explore how this technology can be coupled to produce engaging, personalized, and culturally relevant advertisements that resonate with diverse audiences. Investigating these positive applications will allow brands to enhance consumer engagement while maintaining ethical standards. By addressing these areas, future research can deepen the understanding of deepfake technology's role in advertising, enabling brands to navigate its challenges while leveraging opportunities. These efforts will contribute to the responsible and innovative use of synthetic media in shaping consumer behavior and trust in the digital age.

While this study provides valuable insights into the use of deepfake technology in advertising, it is not without limitations. One of the primary limitations is the relatively small sample size. The study focused primarily on a younger, digitally savvy demographic, which may have skewed the results toward more favorable perceptions of deepfake technology. Another limitation of this research is its dependence on an

experimental design conducted in an online setting. While the experiment provided controlled conditions for measuring consumer responses, it may not have fully captured the complexities of real-world interactions with deepfake advertisements. Another important point to consider is the potential impact of the visual stimuli used. It's crucial to acknowledge that using a different variety of videos could potentially lead to different outcomes in the study.

As deepfake technology continues to evolve, its impact on advertising, particularly in the fashion industry, will positively grow. This study provides a foundation for understanding the complex interplay between synthetic media, consumer behavior, and advertising effectiveness. By balancing innovation with ethical considerations and focusing on creativity, relevance, and transparency, fashion brands can harness the power of deepfake technology to create compelling and engaging advertisements that resonate with consumers in the digital age. The future of fashion advertising lies in the thoughtful application of these powerful tools, creating experiences that captivate and inspire while maintaining the trust and loyalty of an increasingly discerning customer base.

References

- Aguirre, E., D. Mahr, D. Grewal, K. de Ruyter, and M. Wetzels. 2015. Unravelling the personalization paradox: The effect of information collection and trust-building strategies on online advertisement effectiveness. *Journal of Retailing* 91 (1):34–49. doi:10.1016/j.jretai.2014.09.005
- Ang, S. H., Lee, Y. H., & Leong, S. M. (2007). The ad creativity cube: Conceptualization and initial validation. *Journal of the Academy of Marketing Science* 35 (2): 220–32. doi:10.1007/s11747-007-0042-4
- Antipov, G., Baccouche, M., & Dugelay, J.-L. (2017). Face aging with conditional generative adversarial networks. 2017 IEEE International Conference on Image Processing (ICIP), IEEE, Beijing, China, 2089–2093. doi: 10.1109/ICIP.2017.8296650
- Audrezet, A., Kerviler, G., & Moulard, J. (2020). Authenticity under threat: When social media influencers need to go beyond self-presentation. *Journal of Business Research* 117:557–69. doi:10.1016/j.jbusres.2018.07.008
- Aufderheide, P. (1993). *Media Literacy: A report of the national leadership conference on media literacy*. Washington, DC: Aspen Institute
- Bachmair, B. & Bazalgette, C. (2007). The European charter for media literacy: meaning and potential. *Research in Comparative and International Education*, 2(1), 80-87
- Baek, H., & Morimoto, M. (2012). Stay away from me. *Journal of Advertising* 41 (1):59–76. doi:10.2753/JOA0091-3367410105
- Ballester, E. (2011). *Development and Validation of a Brand Trust Scale*. University of Murcia
- Ballester, E., & Alemán, J. (2001). Brand trust in the context of consumer loyalty. *European Journal of Marketing*, 35(11/12), 1238–1258
- Ballester, E. (2004). Applicability of a brand trust scale across product categories: A multigroup invariance analysis. *European Journal of Marketing*, 38(5/6), 573–592
- Bateman, J. (2020). Deepfakes and synthetic media in the financial system: assessing threat scenarios. <http://www.jstor.org/stable/resrep25783.1>

- Becker, M., Wiegand, N., & Reinartz W. J. (2019). Does it pay to be real? Understanding authenticity in TV advertising. *Journal of Marketing* 83 (1):24–50. doi:10.1177/0022242918815880
- Beverland, M. B., Lindgreen, A., & Vink, M. W. (2008). Projecting authenticity through advertising: Consumer judgments of advertisers' claims. *Journal of Advertising* 37 (1):5–15. doi:10.2753/JOA0091-3367370101
- Bhasin, H. (2019, June 3). What is Fashion Marketing?. *Marketing91*. <https://www.marketing91.com/what-is-fashion-marketing/>
- Birtwistle, G., & Moore, C. M. (2007). Fashion clothing—Where does it all end up? *International Journal of Retail & Distribution Management*, 35(3), 210-216
- Bogicevic, V., Seo, S., Kandampully J. A., Liu, S. Q., & Rudd, N. A. (2019). Virtual reality presence as a preamble of tourism experience: The role of mental imagery. *Tourism Management* 74:55–64. doi:10.1016/j.tourman.2019.02.009
- Brinson, N. H., & Britt, B. C. (2021). Reactance and turbulence: Examining the cognitive and affective antecedents of ad blocking. *Journal of Research in Interactive Marketing*, 15(4), 549–570. <https://doi.org/10.1108/JRIM-04-2020-0083>
- Bruhn, M., Schoenmüller, V., Schäfer, D., & Heinrich, D. (2012). Brand authenticity: Towards a deeper understanding of its conceptualization and measurement. *Advances in Consumer Research*, 40,567–576
- Campbell, C., Plangger, K., Sands, S., & Kietzmann, J. (2022). Preparing for an Era of Deepfakes and AI-Generated Ads: A Framework for Understanding Responses to Manipulated Advertising. *Journal of Advertising*, 51:1, 22-38, DOI: 10.1080/00913367.2021.1909515
- Campbell, C., Plangger, K., Sands, S., & Kietzmann, J., Bates, K. (2022). How Deepfakes and Artificial Intelligence Could Reshape the Advertising Industry: The Coming Reality of AI Fakes and Their Potential Impact on Consumer Behavior. *Journal of Advertising Research*
- Campbell, C., Sands, S., Ferraro, C, Tsao, H.-Y J., & Mavrommatis, A. (2020). From data to action: How marketers can leverage AI. *Business Horizons* 63 (2):227–43. doi:10.1016/j.bushor.2019.12.002

- Cano, P., Loscos, A., Bonada, J., Boer, M., & Serra, X. (2000). Voice morphing system for impersonating in karaoke applications. Audiovisual Institute, Pompeu Fabra University
- Çelik, F., Çam, M., & Koseoglu, M. (2022). Ad Avoidance in the Digital Context: A Systematic Literature Review and Research Agenda. *International Journal of Consumer Studies*
- Chang, Y. S., & Yang, C. (2013). Why do we blog? From the perspectives of technology acceptance and media choice factors. *Behavior and Information Technology*, 32(4), 371–386. <https://doi.org/10.1080/0144929X.2012.656326>
- Chesney, R., & Citron, D. (2019). Deepfakes and the new disinformation war: The coming age of post-truth geopolitics. *Foreign Affairs*, vol. 98, no. 1, pp. 147– 152.
- De Bogota' CDC. (2021). The state of fashion 2021
- Dessart, L., & Pitardi, V. (2019). How stories generate consumer engagement: An exploratory study. *Journal of Business Research* 104:183–95. doi:10.1016/j.jbusres.2019. 06.045
- Dewey, C. (2015). How 25 years of photoshop changed the way we see reality. The Washington Post. <https://www.washingtonpost.com/news/the-intersect/wp/2015/02/19/how-25-years-of-photoshop-changed-the-way-we-see-reality/>
- Dodoo, N. A., & Wen, J. (2019). A path to mitigating SNS ad avoidance: Tailoring messages to individual personality traits. *Journal of Interactive Advertising*, 19(2), 116–132. <https://doi.org/10.1080/15252019.2019.1573159>
- Dodoo, N. A., & Wen, J. (2021). Weakening the avoidance bug: The impact of personality traits in ad avoidance on social networking sites. *Journal of Marketing Communications*, 27(5), 457–480. <https://doi.org/10.1080/13527266.2020.1720267>
- Easey, M. (2008). *Fashion Marketing* (3rd ed.). Wiley
- Eberl, A., kühn, J., & Wolbring, T. (2022). Using deepfakes for experiments in the social sciences – A pilot study. *Frontiers in Sociology*. 7:907199
- Edelman (2020). Trust barometer special report: Brand trust in 2020. <https://www.edelman.com/research/brand-trust-2020>

- Eggers, F., O'Dwyer, M., Kraus, S., Vallaster, C., & Güldenber, S. (2013). The impact of brand authenticity on brand trust and SME growth: A CEO perspective. *Journal of World Business*, 48(3), 340–348
- Elitaş, T. (2022). Digital Manipulation 'Deepfake' Technology and the Credibility of What Is Not. *Hatay Mustafa Kemal University Journal of Institute of Social Sciences*, 19(49), 113-128
- El-Murad, J., & West, D. C. (2004). The definition and measurement of creativity: What do we know?. *Journal of Advertising Research* 44 (2):188–201. doi:10.1017/S0021849904040097
- Farid, H., Davies, A., Webb, L., Wolf, C., Hwang, T., Zucconi, A., & Lyu, S. (2019). Deepfakes and audio-visual disinformation. London: Centre for Data Ethics and Innovation.
- Ferrara, M., Franco, A., & Maltoni, D. (2014). The magic passport. *IEEE International Joint Conference on Biometrics*. 1–7, 2014, <https://doi.org/10.1109/BTAS.2014.6996240>
- Ferreira, F., & Barbosa, B. (2017). Consumers' attitude toward Facebook advertising. *International Journal of Electronic Marketing and Retailing*, 8(1), 45–57
- Fine, V. (2019). Verisimilitude and truthmaking. *Erkenntnis*. <https://doi.org/10.1007/s10670-019-00152-z>
- Firc, A. (2021). Applicability of Deepfakes in the Field of Cyber Security. Brno University of Technology, Faculty of Information Technology, Brno. Supervisor Mgr. (Kamil Malinka, Ph.D). Master's Thesis
- Firc, A., Malinka, K., & Hanáček, P. (2023). Deepfakes as a threat to a speaker and facial recognition: An overview of tools and attack vectors. Brno University of Technology, Božetěchova 2, Brno, 612 00, Czech Republic
- Floridi, L. (2018). Artificial intelligence, deepfakes and a future of ectypes. *Philosophy & Technology* 31 (3): 317–21. doi:10.1007/s13347-018-0325-3
- Fornell, C., & Larcker, D. F. (1981). Evaluating Structural Equation Models with Unobservable Variables and Measurement Error. *Journal of Marketing Research*, 18(1), 39. <https://doi.org/10.2307/3151312>

- Fournier, S. (1998). Consumers and their brands: Developing relationship theory in consumer research. *Journal of Consumer Research*, 24(4), 343–353
- Frade, J. L. H., de Oliveira, J. H. C., & Giralddi, J. d. M. E. (2021). Advertising in streaming video: An integrative literature review and research agenda. *Telecommunications Policy*, 45(9), 102186. <https://doi.org/10.1016/j.telpol.2021.102186>
- Friestad, M., & Wright, P. (1994). The persuasion knowledge model: How people cope with persuasion attempts. *Journal of Consumer Research* 21 (1):1–31. doi:10.1086/209380
- Geisser, S. (1974). A Predictive Approach to the Random Effects Model. 61(1): 101-107. <https://doi.org/10.1093/biomet/61.1.101>
- Goodfellow, I.J., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A. & Bengio, Y. (2014). Generative Adversarial Nets. NIPS 2014
- Green, M. C., & Brock, T. C. (2002). In the mind's eye: Transportation-imagery model of narrative persuasion. In *Narrative impact: Social and cognitive foundations*, ed. M. C. Green, J. J. Strange, and T. C. Brock, 315–41. Lawrence Erlbaum Associates Publishers
- Gu, Y., He, M., Nagano, K., & Li, H. (2019). *Protecting World Leaders Against Deep Fakes*. University of Southern California/USC Institute for Creative Technologies Los Angeles CA, USA
- Hair, J. F., Hult, G. T. M., Ringle, C., Sarstedt, M., Danks, N., & Ray, S. (2021). Partial least squares structural equation modeling (PLS-SEM) using R: A workbook. Springer, 1–208
- Hair, J. F., Ringle, C. M., & Sarstedt, M. (2011). PLS-SEM: Indeed a Silver Bullet. *Journal of Marketing. Theory and Practice*, 19(2), 139–152. <https://doi.org/10.2753/MTP1069-6679190202>
- Hair, J. F., Risher, J. J., Sarstedt, M., & Ringle, C. M. (2019). When to use and how to report the results of PLS-SEM. *European Business Review*, 31(1), 2–24. <https://doi.org/10.1108/EBR-11-2018-0203>
- Henseler, J., Ringle, C. M., & Sinkovics, R. R. (2009). The use of partial least squares path modeling in international marketing. *Advances in International Marketing*, 20, 277–319. [https://doi.org/10.1108/S1474-7979\(2009\)0000020014](https://doi.org/10.1108/S1474-7979(2009)0000020014)

- Henseler, J., Ringle, C.M. & Sarstedt, M. (2015). A new criterion for assessing discriminant validity in variance-based structural equation modeling. *J. of the Acad. Mark. Sci.* 43, 115–135 <https://doi.org/10.1007/s11747-014-0403-8>
- Huijstee, M., Boheemen, P., Das, D., Nierling, L., Jahnel, J., Karaboga, M., Martin, F., Kool, L., & Gerritsen, J. (2021). Tackling Deepfakes in European Policy. <https://doi.org/10.2861/325063>
- Hon, L. C., & Grunig, J. E. (1999). Guidelines for measuring relationships in public relations. Gaines, FL:Institute for public relations
- IBM. (2023). What is an autocoder? <https://www.ibm.com/topics/autoencoder>
- Jennings, R. (2023). AI art freaks me out. So I tried to make some. <https://www.vox.com/culture/23678708/ai-art-balenciaga-harry-potter-midjourney-eleven-labs>
- Jolls, T. (2008). Literacy for the 21st century: an overview & orientation guide to media literacy education. Center for Media Literacy. Retrieved on 2 January 2016 from <http://medialit.org/medialitkit.html>
- Jung, A.R. (2017). The influence of perceived ad relevance on social media advertising: An empirical examination of a mediating role of privacy concern. *Computers in Human Behavior*, 70, 303–309. <https://doi.org/10.1016/j.chb.2017.01.008>
- Karnouskos, S. (2020). Artificial intelligence in digital media: The era of deepfakes. *IEEE Transactions on Technology and Society* 1 (3):138–47. doi:10.1109/TTS.2020.3001312
- Karras, T., Laine, S., Aittala, M., Hellsten, J., Lehtinen, J., & Aila, T. (2020). Proceedings of the IEEE/CVF. Conference on Computer Vision and Pattern Recognition (CVPR), pp. 8110-8119
- Kelly, L., Kerr, G., & Drennan, J. (2010). Avoidance of advertising in social networking sites: The teenage perspective. *Journal of Interactive Advertising*, 10(2), 16–27. <https://doi.org/10.1080/15252019.2010.10722167>
- Kelly, L., Kerr, G., & Drennan, J. (2020). Triggers of engagement and avoidance: Applying approach-avoid theory. *Journal of Marketing Communications*, 26(5), 488–508. <https://doi.org/10.1080/13527266.2018.1531053>

- Kietzmann, J., Lee, L. W., McCarthy, I. P., & Kietzmann, T. C. (2020). Deepfakes: Trick or treat?. *Business Horizons* 63 (2):135–46. doi:10.1016/j.bushor.2019.11.006
- Kiliç, B., & Kahraman, M. E. (2023). Current Usage Areas of Deepfake Applications with Artificial Intelligence Technology. *İletişim ve Toplum Araştırmaları Dergisi*, 3(2), 301-332
- Kim, B. H., Han, S., & Yoon, S. (2010). Advertising creativity in Korea. *Journal of Advertising* 39 (2):93–108. doi:10. 2753/JOA0091-3367390207
- Kim, E., Ratneshwar S., & Thorson, E. (2017). Why narrative ads work: An integrated process explanation. *Journal of Advertising* 46 (2):283–96. doi:10.1080/00913367.2016.1268984
- Kover, A. J. (1995). Copywriters' implicit theories of communication: An exploration. *Journal of Consumer Research* 21 (4):596–611. doi:10.1086/209421
- Kover, A. J., Stephen, M., & Goldberg, W. L. J. (1995). Creativity vs. effectiveness? An integrating classification for advertising. *Journal of Advertising Research* 35 (6): 29–41
- Koslow, S., Sasser, S. L., & Riordan, E. A. (2003). What is creative to whom and why?: Perceptions in advertising agencies. *Journal of Advertising Research* 43 (1):96–110. doi:10.2501/JAR-43-1-96-110
- Kwok, A.O.J. & Koh, S.G.M. (2021) Deepfake: a social construction of technology perspective, *Current Issues in Tourism*, Vol. 24 No. 13
- Lane, VR. (2000). The impact of ad repetition and ad content on consumer perceptions of incongruent extensions. *Journal of Marketing* 64 (2):80–91. doi:10.1509/jmkg.64.2. 80.17996
- Lee, H., & Cho, C.-H. (2020). Digital advertising: Present and prospects. *International Journal of Advertising*, 39(3), 332–341. <https://doi.org/10.1080/02650487.2019.1642015>
- Li, H. (2019). Special section introduction: Artificial intelligence and advertising. *Journal of Advertising* 48 (4): 333–7. doi:10.1080/00913367.2019.1654947
- Li, M., & Zhang, J. (2014). Fashion marketing in China: An empirical analysis. *Fashion and Textiles*, 1(1), 3

- Livingstone, S. (2003). The changing nature and uses of media literacy. MEDIA@LSE Electronic Working Papers. 4
- Logan, K. (2013). And now a word from our sponsor: Do consumers perceive advertising on traditional television and online streaming video differently?. *Journal of Marketing Communications*, 19(4), 258–276
- Loureiro, S. M. C., & Kaufmann, H. R. (2016). Luxury values as drivers for affective commitment: The case of luxury car tribes. *Cogent Business & Management*, 3(1), 1171192. <https://doi.org/10.1080/23311975.2016.1171192>
- Machado, A., & Queiroz, M. (2010). Voice Conversion: A Critical Survey. <https://doi.org/10.5281/ZENODO.849853>
- Malhotra, N., Birks, D. (2007). *Marketing Research: an applied approach* (3rd European Edition), Harlow, UK. Pearson Education
- Maras, M. H., & Alexandrou, A. (2019). Determining the authenticity of video evidence in the age of artificial intelligence and in the wake of Deepfake videos. *The International Journal of Evidence & Proof*, 23(3), 255–262
- Marra, J. L. (1990). *Advertising creativity: Techniques for generating ideas*. Prentice Hall
- Masood, M., Nawaz, M., Malik, K., Javed, A., Irtaza, A., & Malik, H. (2021). Deepfakes generation and detection: state-of-the-art, open challenges, countermeasures, and way forward. 53(4), 3974–4026. <https://doi.org/10.1007/s10489-022-03766-z>
- McDonald, C., & Scott, J. (2007). A brief history of advertising. In *The Sage Handbook of Advertising*. ed. G. J. Tellis and T. Ambler, 17–34. Sage
- Molleda, J, C. (2010). Authenticity and the construct's dimensions in public relations and communication research. *Journal of Communication Management*, 14(3), 223–236
- Molleda, J.-C., & Jain, R. (2013). Testing a perceived authenticity index with triangulation research: The case of Xcaret in Mexico. *International Journal of Strategic Communication*, 7(1), 1–20
- Morgan, R. M., & Hunt, S. D. (1994). The commitment-trust theory of relationship marketing. *Journal of Marketing*, 58(3), 20–38

- Morhart, F., Malär, L., Guèvremont, A., Girardin, F., & Grohmann, B. (2015). Brand authenticity: An integrative framework and measurement scale. *Journal of Consumer Psychology*, 25(2), 200–218
- Mukherjee, A., Smith, R. J., & Turri, A. M. (2018). The smartness paradox: The moderating effect of brand quality reputation on consumers' reactions to RFID-based smart fitting rooms. *Journal of Business Research* 92: 290–9. doi:10.1016/j.jbusres.2018.07.057
- Napoli, J., Dickinson, S. J., Beverland, M. B., & Farrelly, F. (2014). Measuring consumer-based brand authenticity. *Journal of Business Research*, 67(6), 1090–1098
- Navarro, V., Strelet, D., Carrubi, D., & Gomez, H. (2023). Media or Information Literacy as Variables for Citizen Participation in Public Decision-Making? A Bibliometric Overview. Department of Business Organization, Universitat Politècnica de València, Spain. Sustainable Technology and Entrepreneurship
- Nunnally, J. C. (1978). *Psychometric Theory*. (2nd ed.). New York: McGraw-Hill
- Ogilvy, D. (1983). *Ogilvy on advertising*. New York, NY: Crown Publishing
- O'Mahony, N., Campbell, S., Carvalho, A., Harapanahalli, S., Hernandez, G. V., Krpalkova, L., & Walsh, J. (2020). Deep learning vs. traditional computer vision. In *Advances in Computer Vision: Proceedings of the 2019, Computer Vision Conference (CVC)*, Volume 11 (pp. 128-144). Springer International Publishing
- Paul, Olympia A. (2021). *Deepfakes Generated by Generative Adversarial Networks*. Honors College Theses. 671
- Pawle, J., & Cooper, P. (2006). Measuring emotion - Love marks, the future beyond brands. *Journal of Advertising Research* 46 (1):38–48. doi:10.2501/ S0021849906060053
- Pfitzinger, H. (2004). Unsupervised speech morphing between utterances of any speakers. *Proceedings of the 10th Australian Int. Conf. Speech Sci. Technol.*, pp. 545–550
- Pitta, D. A., & Katsanis, L. P. (1995). Understanding brand equity for successful brand extension. *Journal of Consumer Marketing*, 12(4), 51-64
- Potter, J. (2018). *Media Literacy*. Library of Congress Cataloging-in-Publication Data. https://books.google.pt/books?hl=en&lr=&id=P6pzDwAAQBAJ&oi=fnd&pg=PT21&dq=Effects+of+Media+Literacy+on+Ad+Falsity&ots=QggRdmv0jU&sig=dLWA_YZt0cqWxA

[wfGD5RUVNzX3Q&redir_esc=y#v=onepage&q=Effects%20of%20Media%20Literacy%20on%20Ad%20Falsity&f=false](https://arxiv.org/abs/2007.04836)

- Qian, K., Zhang, Y., Chang, S., Cox, D., & Hasegawa-Johnson, M. (2020). Unsupervised speech decomposition via triple information bottleneck. *Proceedings of the 37th International Conference on Machine Learning (ICML'20)*, p. 11. JMLR.org, Article 726
- Qian, K., Zhang, Y., Chang, S., Yang, X & Hasegawa-Johnson, M. (2019). AutoVC: zero-shot voice style transfer with only autoencoder loss. Kamalika Chaudhuri, Ruslan Salakhutdinov (Eds.), *Proceedings of the 36th Int. Conf. Mach. Learn. (Proc.Mach. Learn. Res., 97*, pp. 5210–5219. PMLR
- Qin, X., and Jiang, Z. (2019). The impact of AI on the advertising process: The Chinese experience. *Journal of Advertising*
- Rathore, B. (2019). International Peer Reviewed. Refereed Multidisciplinary Journal (EIPRMJ), ISSN: 2319-5045 Volume 8, Issue 2, July-December, 2019, Impact Factor: 5.679, Available online at: www.eduzonejournal.com
- Reed, S., Zeynep, A., Xinchun, Y., Lajanugen, L., Bernt, S., & Honglak, L. (2016). Generative Adversarial Text to Image Synthesis. arXiv preprint arXiv:1605.05396
- Rosengren, S., Dahlen, M., & Modig, E. (2013). Think outside the ad: Can advertising creativity benefit more than the advertiser?. *Journal of Advertising* 42 (4):320–30. doi:10.1080/00913367.2013.795122
- Russell, C. A. (2002). Investigating the effectiveness of product placements in television shows: The role of modality and plot connection congruence on brand memory and attitude. *Journal of Consumer Research* 29 (3):306–18. doi:10.1086/344432
- Russell, C. A., & Russell, D. W. (2009). Alcohol messages in Prime-Time Television Series. *The Journal of Consumer Affairs* 43 (1):108–28. doi:10.1111/j.1745-6606.2008.01129.x
- Rust, T., & Oliver, W. (1994). The Death of Advertising. *Journal of Advertising* 23 (4):71–7. doi:10.1080/00913367.1994.10673460
- Samuel, A. (1959). Some studies in machine learning using the game of checkers. *IBM J Res Dev.* 3(3):210–229.

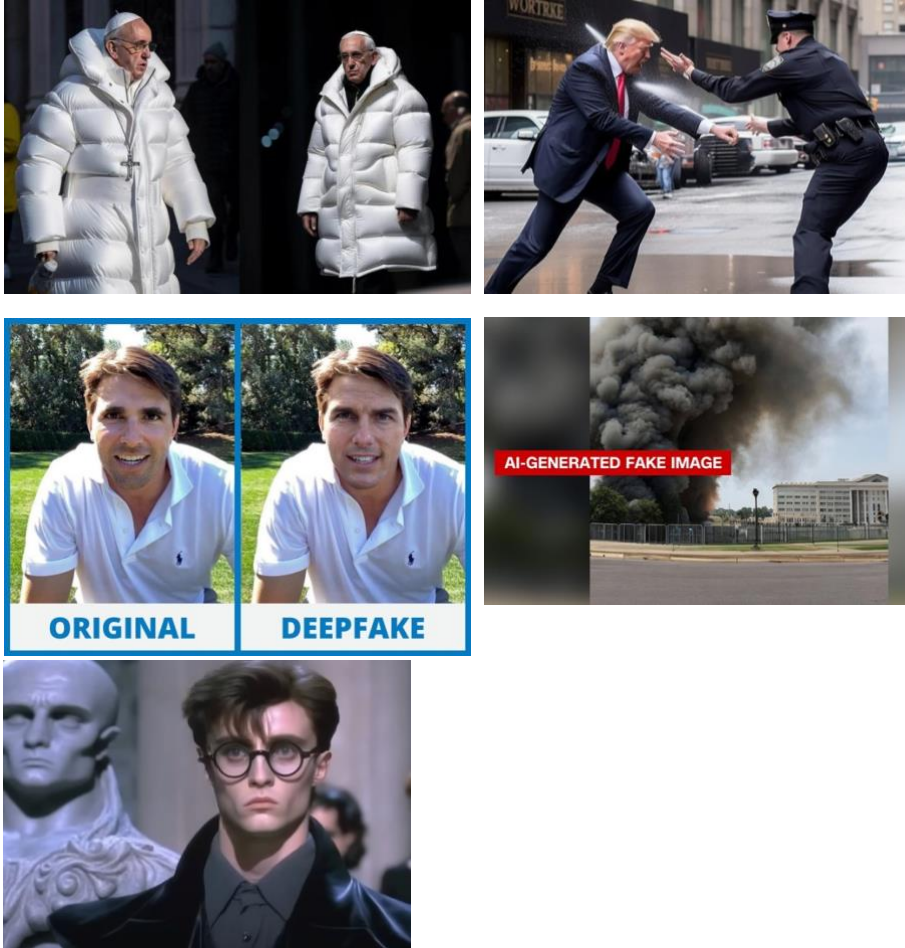
- Saxe, A., Nelli, S., & Summerfield, C. (2021). If Deep Learning Is the Answer, What Is the Question?. *Nature Reviews Neuroscience*, 22(1), 55-67
- Schallehn, M., Burmann, C., & Riley, N. (2014). Brand authenticity: Model development and empirical testing. *Journal of Product & Brand Management*, 23(3), 192–199
- Segal, D. (2023). How Digital Marketing has Impacted the Fashion Industry. LinkedIn
- Schelenz, L., Segal, A., & Gal, K. (2020). Best Practices for Transparency in Machine Generated Personalization. arXiv preprint arXiv:2004.00935
- Sharma, A., Dwivedi, R., Mariani, M. M., & Islam, T. (2022). Investigating the Effect of Advertising Irritation on Digital Advertising Effectiveness: A Moderated Mediation Model. *Technological Forecasting and Social Change*, 180, 121731. <https://doi.org/10.1016/j.techfore.2022.121731>
- Schallehn, M., Burmann, C., & Riley, N. (2014). Brand authenticity: Model development and empirical testing. *Journal of Product & Brand Management*, 23(3), 192–199
- Sheinin, D. A., Varki, S., & Ashley, C. (2011). The differential effect of ad novelty and message usefulness on brand judgments. *Journal of Advertising* 40 (3):5–18. doi:10.2753/JOA0091-3367400301
- Silva, S., Bethany, M., Votto, A., Carft, I., Beebe, N., Najafirad, P. (2022). Deepfake forensics analysis: An explainable hierarchical ensemble of weakly supervised models. *Forensic Science International: Synergy*
- Silver, A. (2009). A European approach to media literacy: moving toward an inclusive knowledge society. In D. Frau-Meigs & J. Torrent (Eds.), *Mapping media education policies in the world: Visions, programmes and challenges* (pp. 11-13). New York: UNAlliance of Civilizations
- Silverblatt, A., Smith, A., Miller, D., Smith, J. & Brown, N. (2014). *Media literacy: Keys to interpreting media messages*. Santa Barbara: Praeger
- Smith, R. E., MacKenzie, S. B., Yang, X., Buchholz, L. M., & Darley W. K. (2007). Modeling the determinants and effects of creativity in advertising. *Marketing Science* 26 (6):819–33. doi:10.1287/mksc.1070.0272

- Spielmann, N., & Orth, U. R. (2020). Can advertisers overcome consumer qualms with virtual reality?: Increasing operational transparency through self-guided 360-degree tours. *Journal of Advertising Research*
- Spiggle, S., Nguyen, H. T., & Caravella, M. (2012). More than fit: Brand extension authenticity. *Journal of Marketing Research* 49 (6):967–83. doi:10.1509/jmr.11. 0015
- Stern, B. (1994). Authenticity and the textual persona: Postmodern paradoxes in advertising narrative. *International Journal of Research in Marketing* 11 (4): 387–400. doi:10.1016/0167-8116(94)90014-0
- Stone, M. (1974). Cross-Validatory Choice and Assessment of Statistical Predictions. *Journal of the Royal Statistical Society*, 36(2), 111–147. <https://doi.org/10.1111/j.2517-6161.1974.tb00994>
- Suwajanakorn, S., Seitz, S., & Kemelmacher-Shlizerman, I. (2017). Synthesizing Obama: Learning Lip Sync From Audio. *ACM Trans Graph* 36:95–108. <https://doi.org/10.1145/3072959.3073640>
- Tabet, Y., & Mohamed, B. (2011). Speech synthesis techniques. A survey, in: *Int. Workshop Syst. Signal Process. Appl, WOSSPA*, pp.67–70. <https://doi.org/10.1109/WOSSPA.2011.5931414>
- Taylor, C., Costello, J. (2017). What do we know about Fashion Advertising? A Review of the Literature and Suggested Research Directions. *Journal of Global Fashion Marketing*. Vol. 8, No. 1, 1-20
- Taylor, P. (2009). *Text-to-Speech Synthesis*. Cambridge University Press. <https://doi.org/10.1017/CBO9780511816338>
- Thies, J., Elgharib, M., Tewari, A., Theobalt, C., & Nießner, M. (2020). Neural Voice Puppetry: Audio-Driven Facial Reenactment. arXiv: 1912.05566 [cs.CV]
- Tong, S., Luo, X., & Xu, B. (2020). Personalized mobile marketing strategies. *Journal of the Academy of Marketing Science* 48 (1):64–78. doi:10.1007/s11747-019- 00693-3
- Tudoran, A. A. (2019). Why do internet consumers block ads? New evidence from consumer opinion mining and sentiment analysis. 29(1), 144–166. <https://doi.org/10.1108/IntR-06-2017- 0221>

- Van Laer, T., Feiereisen S., & Visconti, L. M. (2019). Storytelling in the digital era: A meta-analysis of relevant moderators of the narrative transportation effect. *Journal of Business Research* 96:135–46. doi:10.1016/j.jbusres. 2018.10.053
- Westerlund, M. (2019). The emergence of deepfake technology: A review. *Technology Innovation Management Review*, 9(11), 39–52. <https://doi.org/10.22215/timreview/1282>
- Whittaker, L. W., Kietzmann, T. C., Kietzmann, J., & Dabirian, A. (2020). All around me are synthetic faces: The Mad World of AI-generated Media. *IT Professional*, 22(5), 90–99
- Whittaker, L., Letheren, K., & Mulcahy, R. (2021). The Rise of Deepfakes: A Conceptual Framework and Research Agenda for Marketing. *Australasian Marketing Journal*. Vol. 29(3), 204-214
- Zachary, G. (2020). Digital manipulation and the future of electoral democracy in the US. *Digital Manipulation and the Future of Electoral Democracy in the US*, 1(2), 104-112
- Zinkhan, G. M. (1993). Creativity in advertising. *Journal of Advertising* 22 (2):1–3. doi:10.1080/00913367.1993.10673390
- Zito, P. (2023). Deepfake Advertising: Less or More Value?. Jelly. Agência de Marketing Digital Portugal. Colin Campbell statement. <https://jelly.pt/en/news-en/deepfake-advertising-less-or-more-value/>

Annexes

Annex A: Viral Deepfake



Harry Potter by Balenciaga:

https://www.youtube.com/watch?time_continue=2&v=iE39q-IKOzA&embeds_referring_euri=https%3A%2F%2Fwww.google.com%2Fsearch%3Fq%3Dbelenciaga%2Bharry%2Bpotter%26rlz%3D1C5CHFA_enPT971PT971%26oq%3Dbelenciaga%2Bharry%2Bpotter%26gs_lcrp%3DEgZjaHJvbWUy&source_ve_path=MjM4NTE

Annex B: Measurement Scales

Constructs	Adapted Item	Type of Scale	Source
Perceived Verisimilitude	PV1 - The ad seems realistic. PV2 - The ad resembles daily life tasks. PV3 - The ad represents common, everyday situations.	7-Point Likert Scale of agreement (1 = Strongly Disagree to 7 = Strongly Agree)	Campbell & Reiman (2022)
Perceived Creativity	PC1 - In general, the ad was creative. PC2 - In general, the ad was innovative. PC3 - In general, the ad was clever.	7-Point Likert Scale of agreement (1 = Strongly Disagree to 7 = Strongly Agree)	Yang (2006)
Perceived Relevance	PR1 - The ad was relevant to me. PR2 - The ad spoke to my concerns. PR3 - The ad fits my needs well. PR4 - The ad was important to me. PR5 - The ad was related to something important to me.	7-Point Likert Scale of agreement (1 = Strongly Disagree to 7 = Strongly Agree)	Yang (2006)
Awareness of Falsity	AF1 - I think the ad was fake news. AF2 - The ad sounded doubtful.	5-Point Likert Scale of agreement (1 = Strongly Disagree to 5 = Strongly Agree)	Nijhuis (2018)
Trust in Brand	TB1- With BALENCIAGA I obtain what I look for in an ad. TB2 - BALENCIAGA is a brand name that meets my expectations. TB3 - I feel confidence in BALENCIAG brand name. TB4 - BALENCIAGA is not constant in satisfying my needs. TB5 - BALENCIAGA would be honest and sincere in addressing my concerns. TB6 - BALENCIAGA would make any effort to satisfy me. TB7 - BALENCIAG would be interested in my satisfaction.	5-Point Likert Scale of agreement (1 = Strongly Disagree to 5 = Strongly Agree)	Ballester (2011)
Ad Avoidance	AA1 - I would ignore this ad. AA2 - I would not pay attention to this ad. AA3 - I gloss over this kind of ad. AA4 - I block this kind of ad. AA5 - I click the "hide" option to block this kind of ad.	7-Point Likert Scale of agreement (1 = Strongly Disagree to 7 = Strongly Agree)	Youn & Kim (2019)

Annex C: Questionnaire Visual Stimuli

Deepfake Visual Stimuli:

https://www.youtube.com/watch?v=O2XVFT7ep6M&t=37s&ab_channel=Balenciaga

Non-Deepfake Visual Stimuli:

https://www.youtube.com/watch?v=Zy2mZrYYPWI&ab_channel=Balenciaga

Annex D: Indicator Loadings Complete Model (with PV3)

Perceived Verisimilitude	PV 1 = 0.911	PV2 = 0.917	PV3 = 0.551				
Perceived Creativity	PC1 = 0.933	PC2 = 0.944	PC3 = 0.927				
Perceived Relevance	PR1 = 0.877	PR2 = 0.944	PR3 = 0.930	PR4 = 0.922	PR5 = 0.881		
Awareness of Falsity	AF1 = 0.969	AF2 = 0.955					
Trust in Brand	TB1 = 0.770	TB2 = 0.863	TB3 = 0.833	TB4 = 0.870	TB5 = 0.853	TB6 = 0.861	TB7 = 0.837
Ad Avoidance	AA1 = 0.890	AA2 = 0.861	AA3 = 0.784	AA4 = 0.757	AA5 = 0.791		

Annex E: Indicator Loadings Complete Model (without PV3)

Perceived Verisimilitude	PV 1 = 0.940	PV2 = 0.953					
Perceived Creativity	PC1 = 0.933	PC2 = 0.944	PC3 = 0.927				
Perceived Relevance	PR1 = 0.877	PR2 = 0.881	PR3 = 0.930	PR4 = 0.922	PR5 = 0.881		
Awareness of Falsity	AF1 = 0.961	AF2 = 0.954					
Trust in Brand	TB1 = 0.770	TB2 = 0.863	TB3 = 0.833	TB4 = 0.870	TB5 = 0.853	TB6 = 0.861	TB7 = 0.837
Ad Avoidance	AA1 = 0.889	AA2 = 0.860	AA3 = 0.783	AA4 = 0.759	AA5 = 0.793		

Annex F: Indicator Loadings Deepfake Model (with PV3)

Perceived Verisimilitude	PV 1 = 0.905	PV2 = 0.898	PV3 = 0.561				
Perceived Creativity	PC1 = 0.921	PC2 = 0.930	PC3 = 0.924				
Perceived Relevance	PR1 = 0.882	PR2 = 0.851	PR3 = 0.931	PR4 = 0.911	PR5 = 0.887		
Awareness of Falsity	AF1 = 0.960	AF2 = 0.954					
Trust in Brand	TB1 = 0.775	TB2 = 0.860	TB3 = 0.836	TB4 = 0.856	TB5 = 0.867	TB6 = 0.890	TB7 = 0.858
Ad Avoidance	AA1 = 0.891	AA2 = 0.854	AA3 = 0.777	AA4 = 0.741	AA5 = 0.791		

Annex G: Indicator Loadings Deepfake Model (without PV3)

Perceived Verisimilitude	PV 1 = 0.927	PV2 = 0.947					
Perceived Creativity	PC1 = 0.921	PC2 = 0.930	PC3 = 0.924				
Perceived Relevance	PR1 = 0.882	PR2 = 0.951	PR3 = 0.931	PR4 = 0.911	PR5 = 0.887		
Awareness of Falsity	AF1 = 0.960	AF2 = 0.953					
Trust in Brand	TB1 = 0.775	TB2 = 0.861	TB3 = 0.836	TB4 = 0.856	TB5 = 0.867	TB6 = 0.889	TB7 = 0.858
Ad Avoidance	AA1 = 0.890	AA2 = 0.853	AA3 = 0.777	AA4 = 0.744	AA5 = 0.793		

Annex H: Indicator Loadings Non-Deepfake Model (with PV3)

Perceived Verisimilitude	PV 1 = 0.918	PV2 = 0.937	PV3 = 0.547				
Perceived Creativity	PC1 = 0.936	PC2 = 0.950	PC3 = 0.920				
Perceived Relevance	PR1 = 0.869	PR2 = 0.928	PR3 = 0.927	PR4 = 0.940	PR5 = 0.874		
Awareness of Falsity	AF1 = 0.961	AF2 = 0.956					
Trust in Brand	TB1 = 0.762	TB2 = 0.865	TB3 = 0.825	TB4 = 0.886	TB5 = 0.844	TB6 = 0.823	TB7 = 0.807
Ad Avoidance	AA1 = 0.888	AA2 = 0.870	AA3 = 0.783	AA4 = 0.782	AA5 = 0.813		

Annex I: Indicator Loadings Non-Deepfake Model (without PV3)

Perceived Verisimilitude	PV 1 = 0.953	PV2 = 0.960					
Perceived Creativity	PC1 = 0.936	PC2 = 0.950	PC3 = 0.920				
Perceived Relevance	PR1 = 0.869	PR2 = 0.928	PR3 = 0.927	PR4 = 0.940	PR5 = 0.874		
Awareness of Falsity	AF1 = 0.961	AF2 = 0.956					
Trust in Brand	TB1 = 0.762	TB2 = 0.865	TB3 = 0.825	TB4 = 0.886	TB5 = 0.844	TB6 = 0.823	TB7 = 0.807
Ad Avoidance	AA1 = 0.887	AA2 = 0.870	AA3 = 0.783	AA4 = 0.782	AA5 = 0.814		

Annex J: Construct Reliability and Convergent Validity

	Cronbach's Alpha	Composite Reliability (rho_c)	Composite Reliability (rho_a)	Average Variance Extracted (AVE)
Perceived Verisimilitude	0.885	0.945	0.894	0.896
Perceived Creativity	0.928	0.954	0.935	0.873
Perceived Relevance	0.940	0.954	0.943	0.807
Awareness of Falsity	0.910	0.957	0.912	0.917
Trust in Brand	0.931	0.944	0.932	0.708
Ad Avoidance	0.878	0.910	0.898	0.670

	Cronbach's Alpha – group deepfake	Composite Reliability (rho_c) – group deepfake	Composite Reliability (rho_a) – group deepfake	Average Variance Extracted (AVE) – group deepfake
Perceived Verisimilitude	0.861	0.935	0.876	0.877
Perceived Creativity	0.916	0.947	0.930	0.856
Perceived Relevance	0.936	0.952	0.943	0.797
Awareness of Falsity	0.908	0.956	0.910	0.916
Trust in Brand	0.936	0.948	0.939	0.722
Ad Avoidance	0.872	0.906	0.899	0.661

	Cronbach's Alpha – group non-deepfake	Composite Reliability (rho_c) – group non- deepfake	Composite Reliability (rho_a) – group non-deepfake	Average Variance Extracted (AVE) – group non-deepfake
Perceived Verisimilitude	0.907	0.955	0.910	0.915
Perceived Creativity	0.929	0.955	0.930	0.875
Perceived Relevance	0.947	0.959	0.951	0.825
Awareness of Falsity	0.912	0.958	0.914	0.919
Trust in Brand	0.925	0.940	0.927	0.691
Ad Avoidance	0.886	0.916	0.901	0.686

Annex K: Variance Inflation Factor (VIF) Complete Model

PV1	PV2	PC1	PC2	PC3	PR1	PR2	PR3	PR4	PR5	AF1	AF2	TB x PR
2.698	2.698	4.130	4.580	3.024	2.974	3.415	5.284	4.470	3.200	3.292	3.292	1.000
TB1	TB2	TB3	TB4	TB5	TB6	TB7	AA1	AA2	AA3	AA4	AA5	
1.891	3.766	3.067	3.229	3.181	3.889	2.971	3.581	3.364	1.814	3.044	3.234	

Annex L: Variance Inflation Factor (VIF) Deepfake Model

PV1	PV2	PC1	PC2	PC3	PR1	PR2	PR3	PR4	PR5	AF1	AF2	TB x PR
2.336	2.336	3.542	3.747	2.785	2.771	2.828	5.236	4.067	3.110	3.236	3.236	1.000
TB1	TB2	TB3	TB4	TB5	TB6	TB7	AA1	AA2	AA3	AA4	AA5	
1.830	3.752	3.225	3.111	3.385	4.411	3.151	3.355	3.074	1.795	2.334	2.561	

Annex M: Variance Inflation Factor (VIF) Non-Deepfake Model

PV1	PV2	PC1	PC2	PC3	PR1	PR2	PR3	PR4	PR5	AF1	AF2	TB x PR
3.210	3.210	4.217	4.982	3.021	3.486	6.202	6.735	6.548	3.954	3.371	3.371	1.000
TB1	TB2	TB3	TB4	TB5	TB6	TB7	AA1	AA2	AA3	AA4	AA5	
2.246	4.325	3.118	3.505	3.385197	4.523	2.783	3.963	3.981	1.827	4.902	5.128	

Annex N: Complete Model Discriminant Validity: Fornell-Larker Criterion

	AA	AF	PC	PR	PV	TB
AA	0.818					
AF	0.196	0.958				
PC	-0.551	-0.243	0.935			
PV	-0.610	-0.258	0.584	0.899		
PR	-0.376	-0.510	0.338	0.345	0.947	
TB	-0.526	-0.162	0.504	0.624	0.290	0.842

Annex O: Deepfake Model Discriminant Validity: Fornell-Larker Criterion

	AA	AF	PC	PR	PV	TB
AA	0.813					
AF	0.209	0.957				
PC	-0.512	-0.350	0.925			
PV	-0.551	-0.323	0.556	0.893		
PR	-0.406	-0.425	0.430	0.401	0.937	
TB	-0.488	-0.249	0.536	0.663	0.339	0.849

Annex P: Non-Deepfake Model Discriminant Validity: Fornell-Larker Criterion

	AA	AF	PC	PR	PV	TB
AA	0,828					
AF	0,217	0,959				
PC	-0,597	-0,231	0,936			
PV	-0,694	-0,211	0,618	0,908		
PR	-0,370	-0,586	0,339	0,314	0,956	
TB	-0,566	-0,087	0,453	0,568	0,265	0,831